

**Universidad Nacional Autónoma de Nicaragua
UNAN – León
Facultad de Ciencias y Tecnología
Departamento de Computación**



Tesis para Optar al Título de Ingeniero en Telemática

PROPUESTA DE PRÁCTICAS DE LABORATORIOS DE SWITCHING Y ROUTING PARA LA
CARRERA DE INGENIERÍA EN TELEMÁTICA DE LA UNAN - LEÓN

Elaborado por:

Br. Ruddy Otoniel Quiroz Vázquez.
Br. Franklin Ernesto Ramírez Medina.
Br. Yoel Francisco Rivera González.

Tutor:

MSc. Aldo René Martínez Delgadillo.

León, Nicaragua 23 de agosto de 2013

Le agradezco a Dios por haberme acompañado y guiado a lo largo de mi carrera, por ser mi fortaleza en los momentos de debilidad y por brindarme una vida llena de aprendizajes, experiencias y sobre todo felicidad.

Le doy gracias a mis padres Jorge y Nora por apoyarme en todo momento, por los valores que han inculcado, y por haberme dado la oportunidad de tener una excelente educación en el transcurso de mi vida. Sobre todo por ser un excelente ejemplo de vida a seguir.

A mis hermanos por ser parte de mi vida y representar la unidad familiar. A Edgardo y a Alberto por ser un ejemplo de desarrollo profesional a seguir, a Daymeris y Cleotilde por llenar mi vida de alegría y amor cuando más lo he necesitado

Les agradezco la confianza, apoyo y dedicación de tiempo a todos mis profesores

A Yoel y Ruddy por haber sido excelentes compañeros de tesis y amigos, por haberme tenido la paciencia necesaria y por motivarme a seguir adelante en los momentos de desesperación.

A mis amigos por confiar y creer en mí y haber hecho de mi etapa universitaria un trayecto de vivencia que nunca olvidare.

Franklin Ernesto Ramírez Medina.

En primer lugar quiero agradecer a Dios por haberme prestado la vida para estar culminando una etapa más, la cual fue de mucho trabajo, sacrificio, desánimos en un momento pero con su ayuda pude salir adelante y llegar hasta ésta fase.

A mis padres y hermanas que con tanto sacrificio lucharon día a día para que yo pudiese terminar mi universidad preparándome como profesional, apoyándome psicológicamente para que nunca abandonara mis estudios y siempre siguiera adelante a pesar de las dificultades.

Ruddy Otoniel Quiroz Vázquez

A Dios, a mi padre Francisco Rivera a quien tengo mucha admiración, a mi madre Guadalupe González por ser el motor que me hace seguir siempre viendo hacia adelante, a mis familiares, maestros y amigos.

Yoel Francisco Rivera González.

A todas las personas que harán uso de nuestro trabajo

ÍNDICE DE CONTENIDOS

CAPÍTULO N° 1: ASPECTOS INTRODUCTORIOS	1
1. INTRODUCCIÓN	2
2. ANTECEDENTES	3
3. DEFINICIÓN DE PROBLEMAS	4
4. JUSTIFICACIÓN.....	5
4.1 Originalidad	5
4.2 Alcance	5
4.3 Producto	5
5. OBJETIVOS.....	6
5.1 Objetivo general.....	6
5.2 Objetivos específicos.....	6
6. DISEÑO METODOLÓGICO	7
CAPÍTULO N° 2: DESARROLLO TEÓRICO.....	11
1. Aspectos generales de Redes de Computadores.....	12
1.1 Aspectos generales de Redes de Computadores	13
1.2 Modelo OSI y modelo TCP/IP	14
1.3 Protocolos de Internet	14
2. DIRECCIONAMIENTO IPv4	16
2.1 Direcciones IP	17
2.2 Formato de las direcciones IPv4.....	17
2.3 Máscara de subred	19
2.4 Encaminamiento IP (Enrutamiento)	20
2.5 Subnetting o Subneteo.....	20
3. VLANs.....	34
3.1 Tipos de VLANs	35
3.2 Tipos de Puertos en la VLANs	37
3.3 Interconexiones de Switch con VLANs y Puertos Trunk	37
3.4 Creación de VLANs	38
3.5 Configuración de VLANs	38
4. SPANNING TREE	40
4.1 Spanning Tree Tradicional	41
4.2 Prevención de bucles con el protocolo Spanning Tree	44
4.3 Comunicación Spanning Tree: Bridge Protocol Data Units (BPDU)	44
4.4 Elección del puente raíz	45
4.5 Eligiendo puertos raíz.....	47
4.6 Elección de puertos designado.....	48
4.7 Estados STP.....	51

4.8	Temporizadores STP	51
4.9	Cambio en la topología.....	52
4.10	Tipos de STP	52
4.11	Convergencia de enlaces redundantes	54
4.12	Protección ante BPDU inesperadas en puertos de puente	54
4.13	Protección ante la pérdida de BPDU	55
5.	ETHERNET CHANNEL	58
5.1	EtherChannel.....	59
6.	PPP.....	65
6.1	Definición de PPP	66
6.2	PPP en el Protocolo Stack.....	66
6.3	Protocolo de Control de Enlace (LCP).....	66
6.4	Protocolos de Control de Red (NCP)	66
6.5	Estructura de PPP	66
6.6	Protocolo de control de enlace	67
6.7	Formato de la trama.....	67
6.8	LCP enlace proceso de configuración.....	68
6.9	Establecimiento de comunicaciones PPP	69
6.10	Configurar la autenticación PPP	69
7.	HSRP	71
7.1	Definición de HSRP	72
7.2	HSRP-Ejemplo práctico	73
7.3	HSRP - Formato de paquetes.....	73
7.4	HSRP-Descripción campos.....	73
7.5	Configuración CISCO	75
8.	VRRP	77
8.1	Descripción del protocolo VRRP.....	78
8.2	Definiciones.....	79
8.3	Funcionamiento del protocolo VRRP	80
8.4	Prioridad VRRP	81
8.5	Modo de trabajo.....	81
8.6	Modo de autenticación	82
8.7	Timers VRRP	82
8.8	Formato de paquete VRRP	82
8.9	Tipos de paquetes VRRP.....	83
8.10	Aplicación de VRRP ejemplos	84
8.11	Reparto de carga	84
9.	FRAME RELAY	86
9.1	Definición de Frame Relay.....	87
9.2	Rentabilidad de Frame Relay.....	87
9.3	Flexibilidad de Frame Relay.....	87
9.4	Requisitos de Frame Relay vs líneas dedicadas	88

9.5	Funcionamiento de Frame Relay.....	89
9.6	Circuitos virtuales.....	90
9.7	DLCI	90
9.8	Mapa Frame Relay.....	92
9.9	Encapsulación Frame Relay	93
9.10	LMI.....	95
9.11	Formato de trama LMI	96
9.12	ARP inverso	96
10.	LISTAS DE CONTROL DE ACCESO	97
10.1	Definición de Listas de Control de Acceso (ACL)	98
10.2	Razones para el uso de ACLs	98
10.3	Reglas de Funcionamiento de ACL	99
10.4	Tipos de listas de acceso IP e IPX.....	100
10.5	Número de lista de acceso	101
10.6	Número de Puerto.....	101
10.7	Mascara de Wildcard.....	102
10.8	Configuración de las Listas de Control de Acceso (ACL).....	103
10.9	Filtros aplicados a terminales virtuales	104
10.10	Comandos especiales	104
10.11	Monitoreo de las listas	105
10.12	Tips de aplicación	105
10.13	Borrar las ACLs existentes	106
10.14	Ejemplo de usos de Listas de Control de Accesos (ALC).....	106
11.	TELEFONÍA SOBRE VoIP	108
11.1	Protocolo de voz sobre IP.....	109
11.2	Principales componentes de la VOIP.....	109
11.3	Funcionamiento de VoIP	109
11.4	Transporte de la Voz en redes IP	110
11.5	Clasificación de los protocolos sobre VoIP	113
12.	PROTOCOLOS DE ENRUTAMIENTO CON IPv4	121
12.1	RIP (Media Gateway Control Protocol)	122
12.2	OSPF (Media Gateway Control Protocol)	122
12.3	IS-IS.....	123
12.4	EIGRP (Media Gateway Control Protocol).....	123
13.	PROTOCOLO IPv6 (Direccionamiento).....	144
13.1	IPv6.....	145
14.	PROTOCOLOS DE ENRUTAMIENTO CON IPv6	158
14.1	OSPF3	159
14.2	IS-IS.....	162
14.3	RIPng.....	163
15.	POLÍTICAS DE SEGURIDAD EN REDES DE COMUNICACIONES	166
15.1	Protocolo SSH (Secure Shell)	167

15.2	Servidor RADIUS	169
15.3	IPSec	173
15.4	VPN (Virtual Private Network).....	175
16.	SERVICIOS Y GESTIÓN DE REDES.....	180
16.1	DNS	182
16.2	HTTP	189
16.3	FTP	195
16.4	Correo Electrónico.....	197
16.5	Gestión de Red	201
CAPÍTULO Nº 3: DESARROLLO PRÁCTICO		204
PRÁCTICA Nº 1: ASIGNACIÓN DE DIRECCIONES IPv4		205
PRÁCTICA Nº 2: SEGMENTACIÓN DE REDES A TRAVÉS DE VLANS ESTÁTICAS		210
PRÁCTICA Nº 3: SEGMENTACIÓN DE REDES A TRAVÉS DE VLANS DINÁMICAS.....		218
PRÁCTICA Nº 4: ALTA DISPONIBILIDAD A NIVEL DE ENLACE CON SPANNING TREE Y BALANCEO DE CARGA CON ETHERCHANNEL		226
PRÁCTICA Nº 5: ALTA DISPONIBILIDAD A NIVEL DE ENLACE CON SPANNING TREE Y SEGURIDAD CON EL PROTOCOLO PPP CON AUTHENTICACIÓN CHAP		239
PRÁCTICA Nº 6: BALANCEO DE CARGA HSRP		249
PRÁCTICA Nº 7: BALANCEO DE CARGA VRRP		255
PRÁCTICA Nº 8: FRAME RELAY E INTERVLANS		262
PRÁCTICA Nº 9: LISTAS DE CONTROL DE ACCESO		271
PRÁCTICA Nº 10: TELEFONÍA SOBRE IP (VoIP)		281
PRÁCTICA Nº 11: PROTOCOLOS DE ENRUTAMIENTO CON DIRECCIONAMIENTO IPv4		289
PRÁCTICA Nº 12: PROTOCOLOS DE ENRUTAMIENTO CON DIRECCIONAMIENTO IPv6		301
PRÁCTICA Nº 13: DISEÑO Y ANÁLISIS DE REDES WLAN		313
PRÁCTICA Nº 14: POLÍTICAS DE SEGURIDAD EN REDES DE COMUNICACIONES.....		321
PRÁCTICA Nº 15: GESTIÓN DE SERVICIOS DE APLICACIÓN		330
CAPÍTULO Nº 4: ASPECTOS FINALES.....		340
3.	CONCLUSIONES	341
4.	RECOMENDACIONES	342
5.	BIBLIOGRAFÍA	343

ÍNDICE DE FIGURAS

Figura 1. Ciclo de trabajo.....	7
Figura 2. Simuladores usados	8
Figura 3. Figura de muestra	9
Figura 4. Rango de direcciones IP clase A (Red, Sub Red, y Host)	18
Figura 5. Rango de direcciones IP clase B (Red, Sub Red y Host)	18
Figura 6. Rango de direcciones IP clase C (Red, Sub Red y Host)	19
Figura 7. Diseño tradicional de una red	36
Figura 8. Esquematización de la red en un edificio	36
Figura 9. Diseño de una red usando VLANs.....	36
Figura 10. Cambio en el diseño físico de una red usando VLANs.....	37
Figura 11. Puertos Trunk y Puertos Acces	38
Figura 12. Puente transparente con un Switch.....	42
Figura 13. Puente redundante con 2 Switches	42
Figura 14. Ejemplo de elección de puente raíz	46
Figura 15. Ejemplo de elección de puerto raíz	48
Figura 16. Ejemplo de elección de puerto designado.....	50
Figura 17. Aplicando protección STP	57
Figura 18. Funcionalidad sin implementación EtherChannel	60
Figura 19. Funcionalidad con implementación EtherChannel	60
Figura 20. Estructura del paquete HDLC y PPP.....	68
Figura 21. Configuración de PPP Chap.....	70
Figura 22. Ejemplo de HSRP	73
Figura 23. Formato de paquetes HSRP	73
Figura 24. Configuración Cisco de HSRP	75
Figura 25. Diagrama de red implementando VRRP con un router virtual	81
Figura 26. Formato del paquete VRRPv2.....	83
Figura 27. VRRP en modo master/backup	84
Figura 28. VRRP en modo de reparto de carga	85
Figura 29. Requisitos líneas dedicadas.....	88
Figura 30. Requisitos Frame Relay	88
Figura 31. Operación Frame Relay	90
Figura 32. Circuitos Virtuales.....	91
Figura 33. Importancia DLCI.....	91
Figura 34. Identificación de VC.....	92
Figura 35. Configuración de Frame Relay	92
Figura 36. Encapsulación Frame Relay.....	94
Figura 37. Trama Frame Relay.....	94
Figura 38. Formato de trama Frame Relay	94
Figura 39. Estadísticas LMI	95
Figura 40. Formato de Trama LMI.....	96
Figura 41. Ejemplo del modo de operación de las ACL	98
Figura 42. Ciclo de las ACL	99
Figura 43. Diagrama de muestra de las ACL	106
Figura 44. Funcionamiento de VoIP	109
Figura 45. Flujo de mensajes VoIP.....	114
Figura 46. Protocolos de señalización	115

Figura 47. Intercambio de mensajes con MCU	117
Figura 48. Intercambios de mensajes con el protocolo SIP	120
Figura 49. Comparación de EIGRP con IGRP	125
Figura 50. Identificación de sucesor	127
Figura 51. Instalación de varios sucesores	128
Figura 52. Identificación de suceso factible	129
Figura 53. Tabla de vecinos y topología	129
Figura 54. Establecimiento de relación en Routers vecinos	131
Figura 55. Tabla de vecinos y topología	132
Figura 56. Construcción de tabla topológica	133
Figura 57. Convergencia DUAL algoritmo EIGRP	135
Figura 58. Selección de sucesor factible en la ruta de D a B	137
Figura 59. Sucesor factible de D	138
Figura 60. Sucesor factible de E	139
Figura 61. Elección final de sucesor factible de C, D y E	140
Figura 62. Topología de configuración de EIGRP	141
Figura 63. Configuración manual en la interfaz	142
Figura 64. Auto resumen en redes no continuas	142
Figura 65. Formato de paquetes IPv6	146
Figura 66. Cabeceras de opciones para el destino	148
Figura 67. Cabecera IPv6	149
Figura 68. Divisiones de los campos de direcciones IPv6	151
Figura 69. Direcciones IPv6 UNICAST	151
Figura 70. Identificador de Interfaz UNICAST	151
Figura 71. Monodifusión (anycast)	152
Figura 72. Multidifusión IPv6 (Multicast)	152
Figura 73. Cabeceras de opciones salto a salto para el destino	154
Figura 74. Cabecera de fragmentación	155
Figura 75. Cabecera de encaminamiento genérica	155
Figura 76. Cabecera de encaminamiento tipo 0	155
Figura 77. Topología de muestra para implementación de OSPF3 con IPv6	160
Figura 78. Topología de muestra para implementación de RIPng con IPv6	163
Figura 79. Transporte de información vía SSH	167
Figura 80. Secuencia de acceso a la red.	172
Figura 81. Carga de seguridad encapsuladora	174
Figura 82. Encabezado de autenticación	174
Figura 83. Operación modo Túnel	175
Figura 84. Puerta de enlace VPN	177
Figura 85. Túnel VPN	177
Figura 86. Negociación de asociación de seguridad	178
Figura 87. Consultas recursivas para obtener las correspondencias de pc2.isi.edu	186
Figura 88. Consulta recursiva con un servidor de nombres intermedio	187
Figura 89. Combinación de consultas recursivas e iterativas	188
Figura 90. Formato de mensaje DNS.	189
Figura 91. Comportamiento de petición-respuesta de HTTP	190
Figura 92. Cálculo estimado de la petición de un archivo HTML	192
Figura 93. Formato general de un mensaje de petición	194
Figura 94. Formato general de un mensaje de respuesta	195

Figura 95. Transferencia de archivos entre sistemas locales y remotos	196
Figura 96. Conexiones de control y de datos	196
Figura 97. Perspectiva de alto nivel del sistema de correo electrónico	197
Figura 98. Ruddy envía un mensaje a Yoel.....	198
Figura 99. Cabecera típica de un mensaje SMTP	199
Figura 100. Componentes principales de una arquitectura de gestión de red	202
Figura 101. Topología de práctica de VLAN estática	212
Figura 102. Topología de práctica de VLAN dinámica	220
Figura 103. Topología de práctica de Spnning Tree & EtherChannel	229
Figura 104. Topología de práctica de Spanning Tree & PPP	241
Figura 105. Topología de práctica de HSRP	251
Figura 106. Topología de práctica de VRRP	257
Figura 107. Topología de práctica de Frame Relay & InterVLANS	264
Figura 108. Topología de práctica de Listas de Control de Acceso	273
Figura 109. Topología de práctica de Telefonía sobre VoIP	283
Figura 110. Topología de práctica de protocolos de enrutamiento con direccionamiento IPv4	291
Figura 111. Topología de práctica de protocolos de enrutamiento con direccionamiento IPv6	303
Figura 112. Topología de práctica de Diseño y análisis de una red WLAN	316
Figura 113. Topología de práctica de Políticas de seguridad en redes de comunicaciones	324
Figura 114. Topología de práctica de Gestión de servicios de aplicación.....	332

ÍNDICE DE CUADROS

Cuadro 1. Muestra de cuadro de datos de los equipos	10
Cuadro 2. Ejemplo de una dirección IP	17
Cuadro 3. Rango de direcciones IP según su tipo de clase	18
Cuadro 4. Mascara IP según su tipo de clase	19
Cuadro 5. Tabla de verdad de la compuerta AND	20
Cuadro 6. Tabla de conversión de binario a decimal	21
Cuadro 7. Obtención de la máscara de la red clase A	22
Cuadro 8. Obtención de subredes y cantidad de host por cada subred clase A	23
Cuadro 9. Obtención de la máscara de la red clase B	24
Cuadro 10. Obtención de subredes y cantidad de host por cada subred clase B	25
Cuadro 11. Contenido del mensaje de BPDU de configuración	45
Cuadro 12. Costo de ruta STP	47
Cuadro 13. Contenido de mensaje TCN BPDU	52
Cuadro 14. Distribución de tráfico en un EtherChannel de 2 enlaces	62
Cuadro 15. Tipos de métodos de balanceo de carga EtherChannel	63
Cuadro 16. Configuración de EtherChannel PAgP, LACP, ON	64
Cuadro 17. Estructura de PPP	67
Cuadro 18. Opciones de configuración LCP	69
Cuadro 19. Estado de un Router al enviar un mensaje	74
Cuadro 20. Costos de Frame Relay	89
Cuadro 21. Identificación de VC por cada equipo	92
Cuadro 22. Configuración de Frame Relay por cada equipo	93
Cuadro 23. Map Frame Relay	93
Cuadro 24. Identificadores LMI	96
Cuadro 25. Número de listas de acceso según su tipo	101
Cuadro 26. Número de puertos y protocolos	102
Cuadro 27. Distribución de los protocolos VoIP dentro del modelo OSI	110
Cuadro 28. Vecinos del proceso IP-EIGRP	126
Cuadro 29. Leyenda de convergencia DUAL algoritmo EIGRP	136
Cuadro 30. Leyenda de selección de sucesor factible en la ruta de D a B	137
Cuadro 31. Leyenda sucesor factible de D	138
Cuadro 32. Leyenda sucesor factible de E	139
Cuadro 33. Leyenda elección final de sucesor factible de C, D y E	140
Cuadro 34. Distribución de los tipos de direcciones IPv6	153
Cuadro 35. Comparaciones de direcciones IPv4 e IPv6	154
Cuadro 36. Tipo de paquetes del campo Code	171

**Propuesta de Prácticas de Laboratorios de Switching y Routing
para la Carrera de Ingeniería en Telemática de la UNAN – León.**

CAPÍTULO N° 1: ASPECTOS INTRODUCTORIOS

1. INTRODUCCIÓN

El presente documento tiene como propósito desarrollar prácticas de laboratorios de Switching y Routing para la carrera de Ingeniería en Telemática de la UNAN-León; en el cual se desarrollarán enunciados y abordarán temáticas teóricas y prácticas, permitiendo que los estudiantes puedan adquirir conocimientos teóricos-prácticos y estos sean puestos en práctica durante el desarrollo de los laboratorios relacionados al tema de Switching y Routing.

Se estarán abordando temas de importancia tales como segmentación de redes por medio de VLANs, protocolos y tecnologías que brindan seguridad en las capas de enlace y red, alta disponibilidad, balanceo de carga, conexiones de acceso remoto y conectividad inalámbrica.

El trabajo monográfico está compuesto por dos partes. La primera parte es donde se desarrollan los contenidos teóricos que serán necesarios para poder desarrollar y dar solución a las prácticas, al igual que el enunciado de cada una de ellas. La segunda parte contiene el desarrollo y solución de cada una de las prácticas propuestas.

2. ANTECEDENTES

El presente trabajo monográfico tiene como antecedente un trabajo de fin de carrera 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. El documento fue desarrollado con un total de 17 prácticas donde se abarcan las configuraciones de Switching y Routing básicas y sencillas utilizando la tecnología cisco. El trabajo fue dirigido por el profesor Aldo René Martínez Delgadillo, para el año 2008, en la Universidad Nacional Autónoma de Nicaragua-León (UNAN-León), Departamento de Computación, carrera de Ingeniería en Sistema de Información.

Teniendo como punto de referencia el trabajo antes mencionado, desarrollamos un nuevo documento en donde se proponen nuevas prácticas de laboratorios enfocado a la carrera de Ingeniería en Telemática de la UNAN-León, que cumpla con las necesidades educativas para los estudiantes y las exigencias que las empresas hoy en día demandan en el mercado laboral en el área de redes.

Es válido mencionar que existen documentos y sitios web con algunos ejercicios prácticos relacionado con este tema; sin embargo, no existe uno completo por parte de la UNAN-León con amplios contenidos teóricos-prácticos como los que se desarrollan en el presente.

3. DEFINICIÓN DE PROBLEMAS

El desarrollo de prácticas de laboratorios, es esencial para la formación académica de los estudiantes de la carrera de Ingeniería en Telemática de la UNAN-León. Sin embargo, la ausencia de un documento estándar con una secuencia de los contenidos abordados, dificulta algunas veces la realización de las mismas por parte de los estudiantes.

Los aspectos antes mencionados provocan que los docentes, en este momento, no posean un documento base actualizado para la asignación de prácticas de laboratorios relacionadas con Switching y Routing, y que los estudiantes no posean una documentación y/o bibliografía oficial para dar una correcta solución de las mismas.

Estos problemas provocan las siguientes interrogantes:

- ¿Será necesario desarrollar un plan de prácticas de laboratorio relacionadas con Switching y Routing 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 relacionado con Switching y Routing?
- ¿Qué temas deben ser abordados, donde estos a su vez, sean de importancia en el área de redes?

4. JUSTIFICACIÓN

Tomando como referencia los problemas antes expuestos, se vuelve una necesidad la creación de un documento que permita a los estudiantes realizar de manera eficiente, ordenada y secuencial el desarrollo de prácticas de laboratorios relacionadas con Switching y Routing para la carrera de Ingeniería en Telemática del Departamento de Computación de la UNAN-León.

Por este motivo hemos decidido:

- Proponer prácticas de laboratorio bien documentadas y enunciadas que aumenten gradualmente su complejidad, y afines a algunos escenarios de la vida real.

4.1 Originalidad

Existen trabajos y documentos que contienen prácticas a fines a estas áreas; sin embargo ninguno de ellos se adecúa a las prácticas de laboratorios relacionadas con Switching y Routing de la carrera de Ingeniería en Telemática, por lo que este trabajo ayudará de manera significativa a mejorar la calidad del proceso enseñanza-aprendizaje en la carrera.

4.2 Alcance

Este documento será de gran apoyo tanto para docentes y estudiantes.

- **Para los docentes**
 - Tendrán un documento que les permita asignar prácticas guiadas a los estudiantes con una secuencia y un nivel de complejidad uniforme al transcurso de las prácticas.
- **Para los estudiantes**
 - Comprender los temas y resolver los ejercicios prácticos propuestos de cada una las prácticas asignadas.

4.3 Producto

El documento presenta las características siguientes:

- **Guiado:** Describirá paso a paso la forma en la que se deberá realizar cada práctica y se explicará de forma breve y de manera sencilla el funcionamiento de los protocolos y tecnologías que se configuraran en el desarrollo de las prácticas.
- **Sencillo:** Debido a que el desarrollo de las prácticas y la documentación serán elaboradas de manera que los estudiantes puedan comprenderla fácilmente.
- **Secuencial:** Por el aumento de la complejidad de configuración de algunas prácticas, se propusieron las mismas de manera secuencial, para que el estudiante tenga una secuencia lógica de aprendizaje.

5. OBJETIVOS

5.1 Objetivo general

- Crear propuesta de prácticas de laboratorio de Switching y Routing para la carrera de Ingeniería en Telemática, Departamento de Computación de la UNAN – León.

5.2 Objetivos específicos

- Elaborar un documento con información teórica y práctica, que sirva de base a los estudiantes para el desarrollo de las prácticas de laboratorios.
- Definir un formato que registrará el enunciado de las prácticas de laboratorios a desarrollar, en base a la experiencia acumulada en el transcurso de la carrera.
- Enunciar las prácticas en orden secuencial lógico de acuerdo a la complejidad que presenta cada una ellas, abordando temas de nivel medio-avanzado, con contenidos como balanceo de carga, alta disponibilidad y seguridad.

6. DISEÑO METODOLÓGICO

La metodología del presente trabajo es **Teoría Fundamentada con un diseño sistemático**¹. Para poder cumplir con los objetivos de este trabajo monográfico, se siguió el siguiente esquema.

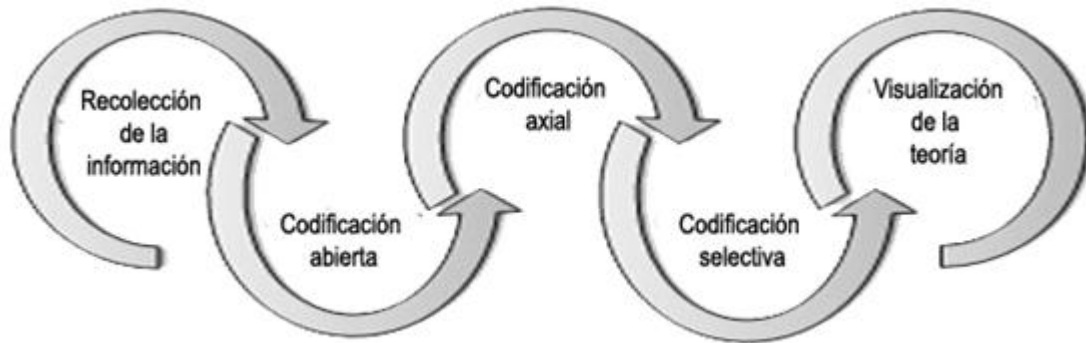


Figura 1. Ciclo de trabajo

➤ **Recolección de datos**

- Análisis y selección de los contenidos teóricos y prácticos que se estarán abordando.

➤ **Codificación abierta**

- Recolección y búsqueda de la información
- Selección de la información útil de acuerdo al desarrollo de los contenidos de cada tema.

➤ **Codificación axial**

- Organización de la información seleccionada: Es el punto donde la información es organizada según el nivel de complejidad que tiene cada uno de los temas a desarrollar en los aspectos teóricos, así como también en los aspectos prácticos. La secuencia de los contenidos teóricos es la siguiente:
 - Aspectos generales de Redes de Computadoras
 - Direccionamiento IPv4
 - VLANs
 - Spanning Tree
 - Ethernet Channel
 - PPP
 - HSRP
 - VRRP
 - Frame Relay
 - Listas de Control de Acceso
 - Telefonía sobre IP (VoIP)

¹ “Metodología de la Investigación”, Roberto Sampieri; 5 ed.

- Protocolos de enrutamiento IPv4
- Protocolo IPv4 (Direccionamiento)
- Protocolos de enrutamiento IPv6
- Políticas de seguridad en redes de comunicaciones
- Servicios y gestión de redes

La organización de las prácticas es la siguiente:

- Asignación de direcciones IPv4
- Segmentación de redes a través de VLANs estáticas
- Segmentación de redes a través de VLANs dinámica.
- Alta disponibilidad a nivel de enlace con Spanning Tree y balanceo de carga con EtherChannel
- Alta disponibilidad a nivel de enlace con Spanning Tree y seguridad con el protocolo PPP con autenticación CHAP.
- Balanceo de carga HSRP
- Balanceo de carga VRRP
- Frame Relay e InterVLANs
- Listas de Control de Acceso
- Telefonía sobre IP (VoIP)
- Protocolos de enrutamiento con direccionamiento IPv4
- Protocolos de enrutamiento con direccionamiento IPv6
- Diseños y análisis de Redes WLAN
- Políticas de seguridad en redes de comunicaciones
- Gestión y servicios de aplicación

➤ **Codificación Selectiva**

- Elección de emuladores a usar en el desarrollo de las prácticas, en donde para utilizar un simulador en una determinada práctica, se tomará en cuenta el soporte que tienen cada uno de ellos para las tecnologías y/o protocolos que serán usados en la práctica enunciada, siendo así la manera que se selecciona el simulador adecuado.



Figura 2. Simuladores usados

- Desarrollo del enunciado de las prácticas: El formato propuesto para enunciar cada una de las practicas propuestas es el siguiente:

Título

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 a rasgos generales 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 una lista detallada con las características de la computadora que se usará en la realización de las prácticas.
- **Software:** Detalla el simulador o entorno en que se desarrollará 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

La Figura 3 ilustra un ejemplo de cómo se muestra la topología correspondiente a cada práctica.

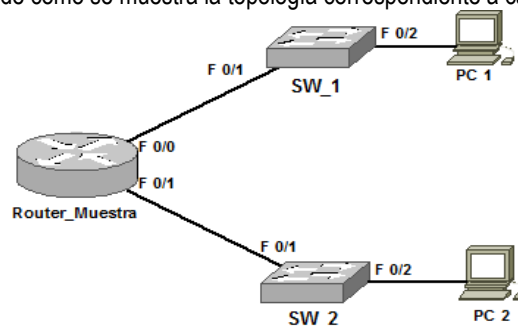


Figura 3. Figura de muestra

Funcionalidad

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

Resumen de comando

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

Datos de los equipos

Los datos necesarios serán mostrados en cuadros, similar al Cuadro 1:

Router

Nombre	Dirección IP	
	Interfaz Fa 0/0	Interfaz Fa 0/1
Router_Muestra	192.168.1.0/24	192.168.2.0/24

Cuadro 1. Muestra de cuadro de datos de los equipos

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 clases 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.

➤ Visualización de la teoría

Presentación de los documentos finales generados con los aspectos teóricos y enunciados de las prácticas, así como también la solución de cada una de las prácticas que han sido propuestas.

CAPÍTULO N° 2: DESARROLLO TEÓRICO

En este capítulo sentaremos las bases teóricas para la correcta comprensión del resto del documento. Serán abordados todos los aspectos teóricos que serán necesario para dar solución a las prácticas.

1. Aspectos generales de Redes de Computadores

1.1 Aspectos generales de Redes de Computadores

La estructura de las redes de computadores varía dependiendo de su tamaño, tecnología y topología. Las redes de propiedad privada dentro de un sólo edificio o instalaciones próximas entre sí, son llamadas redes de área local (*Local Area Network*, LAN). Se usan ampliamente para conectar computadores personales y estaciones de trabajo en oficinas de compañías y fábricas. Dos de las principales redes de difusión son las topologías de buses y anillos normadas por los estándares IEEE 802.3 y IEEE 802.5 respectivamente. Para abarcar grupos de oficinas corporativas, ciudades o redes privadas o públicas de mayor tamaño se utilizan las redes de área metropolitana (*Metropolitan Area Network*, MAN). Finalmente para conectar países o continentes se utilizan redes de área amplia (*Wide Area Network*, WAN), cuya función es principalmente interconectar redes usando líneas de transmisión y elementos de conmutación.

La capacidad de transmisión, el retardo, el costo y las facilidades de instalación y mantenimiento de una red quedan determinadas por el medio de transmisión que se utilice. Se pueden usar diferentes medios físicos para interconectar equipos como par trenzado, cable coaxial, fibra óptica y sistemas inalámbricos dependiendo de las características propias que se requieran de la red.

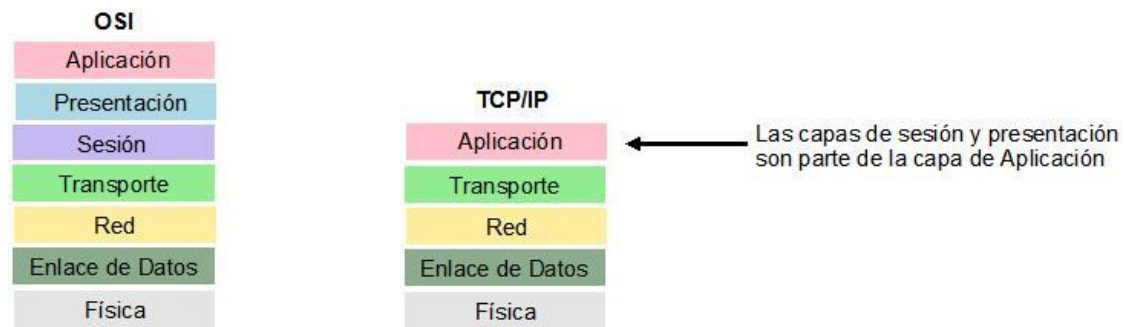
Para reducir la complejidad de diseño, las redes de computadores han sido implementadas en su mayoría utilizando el concepto de capas o niveles. En este esquema, cada capa realiza una determinada función que sirve para ofrecer ciertos servicios a las capas siguientes. Tanto el número como la funcionalidad de los niveles varían de red en red.

Cada capa interactúa sólo con la capa del mismo nivel de la otra máquina. Esta interacción no es física sino que virtual, o sea, los datos se encapsulan hasta llegar a la capa correspondiente del otro host. Los protocolos son el conjunto de reglas y convenciones definidas para la interacción de cada nivel y que, en conjunto con el grupo de capas, definen una arquitectura de red. La figura 1 muestra un esquema genérico de una estructura de red.

Los modelos de red basados en capas, los más importantes son el modelo OSI (Open Systems Interconnection) y el TCP/IP. El primero, aceptado internacionalmente, es un modelo teóricamente completo que sirve para describir otras arquitecturas de red en forma abstracta. Por otro lado, el modelo TCP/IP, a pesar de carecer de ciertas funcionalidades, ha sido el esquema más implementado. Por estas razones, ambos modelos son analizados.

El Departamento de Defensa de Estados Unidos patrocinó a mediados del siglo veinte el proyecto de investigación ARPANET, el cual buscaba conectar diferentes redes de computadores de universidades e instalaciones del gobierno. Producto de la diversidad de tecnologías utilizadas por estas redes, se hizo necesario la creación de un estándar que permitiera la interconexión de ellas. Se buscaba además que la red fuera robusta a ataques que dejaran fuera de servicio a alguno de sus nodos y que fuese flexible a numerosas aplicaciones que en esos años se proyectaban. Es así como en 1974 se define el modelo de referencia TCP/IP. La correspondencia entre el modelo OSI y el modelo TCP/IP se muestra en la figura

1.2 Modelo OSI y modelo TCP/IP



Las funciones de las capas internet, transporte y aplicación del modelo TCP/IP son análogas a las capas correspondientes del modelo OSI. Las capas de presentación y sesión no fueron implementadas en un principio pues muy pocas aplicaciones hacen uso de ellas. Sin embargo, debido al gran aumento de tráfico en tiempo real, se ha necesitado incorporar nuevos protocolos con similares funcionalidades. Bajo la capa de internet el modelo TCP/IP es flexible ya que sólo establece que se debe conectar a la red utilizando un protocolo que permita mandar paquetes, dejando abierta la posibilidad de interconectar tecnologías disímiles.

Para diseñar una red de computadores, se dispone de los siguientes dispositivos básicos: hubs, repetidores, bridges, switches y routers.

1.3 Protocolos de Internet

Los protocolos principales del modelo TCP/IP son el protocolo de Internet (Internet Protocol, IP), el protocolo de control de transmisión (Transmisión Control Protocol, TCP) y el protocolo de datagrama de usuario (User Datagram Protocol, UDP).

1.3.1 Protocolo IP

Perteneciente a la capa de internet, el protocolo IP, permite enrutar información a través de diferentes tecnologías tanto de hardware como de software. Permite direccionar, fragmentar y corregir paquetes.

Dentro de los campos que posee este protocolo se encuentra uno que permite indicar a la subred el tipo de servicio que se desea. Son posibles diversas combinaciones entre confiabilidad y velocidad. En la práctica, este campo ha sido utilizado en redes con calidad de servicio para telefonía IP pertenecientes a redes privadas corporativas. Sin embargo, en redes públicas este campo ha sido ignorado, por lo que IP sólo ofrece un servicio no confiable y no orientado a la conexión.

1.3.2 Protocolo TCP

TCP es un protocolo confiable orientado a la conexión que se implementa en la capa de transporte. Dado que el protocolo IP no es ni confiable ni orientado a conexión, es responsabilidad de TCP ofrecer estos servicios a la capa de aplicación. Debe asegurarse que los datos que fueron transmitidos sean recibidos sin error y en el orden correcto, retransmitiendo información perdida y almacenando y reordenando paquetes en caso necesario.

El protocolo TCP se implementó para adaptarse dinámicamente a las condiciones de red por lo que una de sus principales características es el control de flujo. Esta funcionalidad permite que el protocolo sea capaz de determinar la tasa de transmisión de paquetes, para no sobrecargar a receptores lentos. Además, si aumenta la tasa de pérdida de paquetes (Packet Loss Rate, PLR), el protocolo disminuye exponencialmente su tasa de transmisión, de modo de no saturar la red. Este mecanismo hace que la tasa de transmisión sea conservadora y discontinua.

1.3.3 Protocolo UDP

A diferencia de TCP, el protocolo UDP no es un protocolo confiable ni orientado a conexión. Ofrece a la capa de aplicación un mecanismo para enviar paquetes IP sin establecer conexión. No existe retroalimentación sobre las condiciones de red por lo que no se retransmiten paquetes perdidos ni se realiza control de flujo. El protocolo no disminuirá su tasa de transmisión en presencia de un PLR (Packet Loss Rate) alto, por lo que en la práctica el tráfico UDP resulta ser muy agresivo. El protocolo no garantiza que los paquetes sean recibidos, más aún pueden estos llegar en desorden o duplicados.

2. DIRECCIONAMIENTO IPv4

2.1 Direcciones IP

Las direcciones IP, en sentido general, se utilizan para especificar el destinatario de los datos dentro de una red.

Cada equipo conectado en una red en particular posee una dirección física vinculada al protocolo de acceso, con un formato dependiente de este.

Las direcciones IP, también son conocidas como direcciones de Internet, identifican de forma única y global a cada sistema final y a cada sistema intermedio. Este sistema se vincula a la red a la que pertenece cada equipo y hacen posible el encaminamiento de los paquetes extremo a extremo, a través del complejo entramado que supone el Internet.

Habitualmente se asigna una única dirección IP a un ordenador; aunque cabe la posibilidad de que un ordenador que presente dos conexiones a Internet, cada una de ella a través de una red distinta, posea dos direcciones, se califica entonces al ordenador de “multi-host”. Por otro lado las direcciones deben carecer de ambigüedad, es decir, no debe adjuntarse las mismas direcciones a varios equipos. La unicidad de las direcciones queda garantizada mediante un registro por una autoridad competente en la materia conocida como IANA, que se ocupa de asignar las direcciones de manera que nunca dos se repitan. Finalmente dada la vinculación de las direcciones a las redes, si un ordenador se traslada de una red a otra, debe modificarse su dirección IP.

2.2 Formato de las direcciones IPv4

Las direcciones IP constan de 32 bits, repartidos en tres secciones:

- Un código que indica la clase de red.
- Un identificador de la red.
- Un identificador de la estación dentro de su red

La dirección está codificada para permitir una asignación variable de bits a cada uno de los identificadores. Con ellos se pretende una flexibilidad en el direccionamiento acorde a la variedad de tamaños de las redes presentes en Internet.

Las direcciones IP constan de 32 bits agrupados en 4 octetos, donde cada octeto es un número decimal de (0 – 255). En el Cuadro 2 se muestra el ejemplo de una dirección IP, y el Cuadro 3 las muestras según su rango.

1100 0000	1010 1000	0001 0000	0000 0010
192	168	16	4

Cuadro 2. Ejemplo de una dirección IP

Clase	Rango del 1er Octeto	Octetos de Red (R) y de Host (H)	Máscara de Red	Rango de Red Privada
A	00000001 - 01111111 1 – 127	R.H.H.H	255.0.0.0/8	10.0.0.0
B	10000000 - 10111111 128 - 191	R.R.H.H	255.255.0.0/16	172.16.0.0 172.32.0.0
C	10000000 - 11011111 192 - 223	R.R.R.H	255.255.255.0/24	192.168.0.0
D	11100000 - 11101111 224 - 239	NA-MULTICAST?		
E	11110000-11111111 240-255	NA-Reservada por IETF		

Cuadro 3. Rango de direcciones IP según su tipo de clase

2.2.1 Clase A

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 2^7 , es decir 128 redes; sin embargo, los ordenadores que consta una red pueden alcanzar el número de 2^{24} , la Figura 4 ilustra cómo está organizada una dirección IP de clase A.

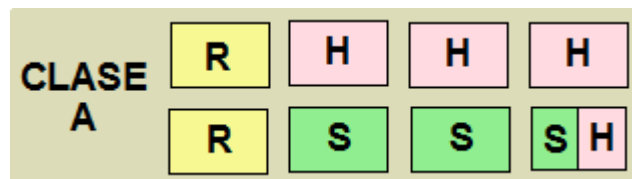


Figura 4. Rango de direcciones IP clase A (Red, Sub Red, y Host)

2.2.2 Clase B

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), la Figura 5 ilustra cómo está organizada una dirección IP de clase B.

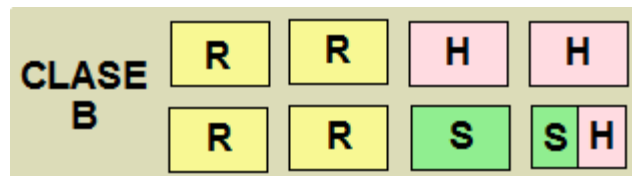


Figura 5. Rango de direcciones IP clase B (Red, Sub Red y Host)

2.2.3 Clase C

Las direcciones de esta clase se destinan a redes con pocos ordenadores, que son lo más frecuentes. Es posible direccionar hasta 2^{21} redes, esto es hasta 2, 097,152. En la Figura 6 se ilustra cómo está organizada una dirección IP de clase B.

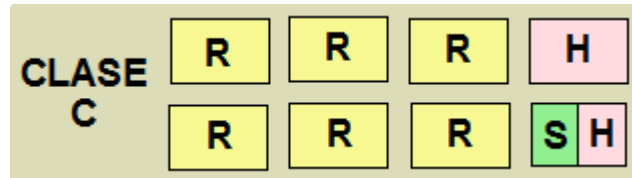


Figura 6. Rango de direcciones IP clase C (Red, Sub Red y Host)

2.2.4 Clase D y clase E

Se trata de direcciones reservadas para usos especiales. Por ejemplo, las direcciones de clase D se utilizan para la multidistribución (multicast).

2.3 Máscara de subred

Las máscaras de subred son valores de 32 bits que permiten a los equipos distinguir qué parte de la dirección IP de los paquetes es id de red (tienen sus bits a 1) y qué otra es id de host (tienen sus bits a 0). El valor resultante se expresa en la misma notación que una dirección IP, según el Cuadro 4:

Clase	Bits para la Máscara de Subred	Máscara de Subred
A	11111111 00000000 00000000 00000000	255.0.0.0
B	11111111 11111111 00000000 00000000	255.255.0.0
C	11111111 11111111 11111111 00000000	255.255.255.0

Cuadro 4. Mascara IP según su tipo de clase

Por ejemplo, cuando la dirección IP es 138.57.7.27 y las máscara de subred es 255.255.0.0, el id de red es 138.57 y el id de host es 7.27.

Es importante que todos los equipos de una misma red lógica utilicen la misma máscara de subred y el mismo id de red, así se evitan problemas de direccionamiento y encaminamiento.

Aunque puede parecer redundante utilizar máscaras de subred, ya que éstas son fácilmente identificables analizando el primer byte de la dirección IP, también se utilizan para dividir aún más un id de red asignado entre varias redes locales. A veces sólo es necesario dividir partes de un byte usando algunos bits para especificar los identificadores de la subred.

Para calcular la máscara de Subred se utiliza la función AND.

Bits		AND
1	1	1
1	0	0
0	1	0
0	0	0

Cuadro 5. Tabla de verdad de la compuerta AND

2.4 Encaminamiento IP (Enrutamiento)

Las redes TCP/IP se conectan mediante encaminadores, gateway o puerta de enlace, dispositivos que pasan los paquetes IP de una red a otra. Un gateway conoce los id de las otras redes, y sólo es necesario cuando se conecte entre sí un conjunto de ellas. Además un gateway tiene unas tablas de encaminamiento, que contienen las direcciones IP de los nodos locales y de las demás subredes.

La búsqueda de la dirección de destino en el encaminador se realiza desde lo específico hasta lo general según el siguiente orden:

- Se busca una coincidencia exacta (ruta de host)
- Se elimina el id de host y se examina la tabla en busca de una coincidencia (ruta de subred)
- Se usa el gateway predeterminado
- Si no se ha especificado ninguno, se descarta el paquete

El proceso de encaminamiento consiste en el envío de los paquetes entre los encaminadores hasta que se encuentra el destino especificado. El programa route añade manualmente rutas a la tabla de encaminamiento para los sistemas o redes más utilizados.

2.5 Subnetting o Subneteo

El subnetting IP hace referencia a cómo están direccionadas las redes IP, y en cómo una red se puede dividir en pequeñas redes.

La función del Subneteo o Subnetting es dividir una red IP física en subredes lógicas (redes más pequeñas) para que cada una de estas trabaje a nivel envío y recepción de paquetes como una red individual, aunque todas pertenezcan a la misma red física y al mismo dominio.

El Subneteo permite una mejor administración, control del tráfico y seguridad al segmentar la red por función. También, mejora la performance de la red al reducir el tráfico de broadcast de nuestra red. Como desventaja, su implementación desperdicia muchas direcciones, sobre todo en los enlaces seriales.

2.5.1 Convertir Bits en Números Decimales

Como sería casi imposible trabajar con direcciones de 32 bits, es necesario convertirlas en números decimales. En el proceso de conversión cada bit de un intervalo (8 bits) de una dirección IP, en caso de ser "1" tiene un valor de "2"

elevado a la posición que ocupa ese bit en el octeto y luego se suman los resultados. Explicado parece medio engorroso pero con la ilustración mostrada en el Cuadro 6, y los ejemplos se va a entender mejor.

Posición y Valor de un Bits								
	2 ⁷	2 ⁶	2 ⁵	2 ⁴	2 ³	2 ²	2 ¹	2 ⁰
Binario	1	0	0	0	0	0	0	0
Decimal	128	0	0	0	0	0	0	0
Binario	0	1	0	0	0	0	0	0
Decimal	0	64	0	0	0	0	0	0
Binario	0	0	1	0	0	0	0	0
Decimal	0	0	32	0	0	0	0	0
Binario	0	0	0	1	0	0	0	0
Decimal	0	0	0	16	0	0	0	0
Binario	0	0	0	0	1	0	0	0
Decimal	0	0	0	0	8	0	0	0
Binario	0	0	0	0	0	1	0	0
Decimal	0	0	0	0	0	4	0	0
Binario	0	0	0	0	0	0	1	0
Decimal	0	0	0	0	0	0	2	0
Binario	0	0	0	0	0	0	0	1
Decimal	0	0	0	0	0	0	0	1
Valor decimal	128	64	32	16	8	4	2	1
Sumatoria	128	192	224	240	248	252	254	255
Octeto N° 1		/2	/3	/4	/5	/6	/7	/8
Octeto N° 2	/9	/10	/11	/12	/13	/14	/15	/16
Octeto N° 3	/17	/18	/19	/20	/21	/22	/23	/24
Octeto N° 4	/25	/26	/27	/28	/29	/30	/31	

Cuadro 6. Tabla de conversión de binario a decimal

Para calcular la cantidad de Subredes es igual a: $2^N - 2$, donde "N" es el número de bits "robados" a la porción de Host y "-2" porque la primer subred y la última subred no son utilizables ya que contienen la dirección de la red y broadcast respectivamente.

Cantidad de Hosts por Subred es igual a: 2^{M-2} , donde "M" es el número de bits disponible en la porción de host y "-2" es debido a que toda subred debe tener su propia dirección de red y su propia dirección de broadcast.

2.5.2 Subneteo Manual de una Red Clase A

Dada la dirección IP Clase A **10.0.0.0/8** para una red, se nos pide que mediante subneteo obtengamos 7 subredes. Este es un ejemplo típico que se nos puede pedir, aunque remotamente nos topemos en la vida real.

Se realizará en dos pasos:

- a. Adaptar la Máscara de Red por Defecto a Nuestras Subred, ya que la máscara por defecto para la red 10.0.0.0 es:

Máscara	Red	Hosts	N° de Barra
255.0.0.0	1111 1111.0000 0000. 0000 0000. 00000000		8

Mediante la fórmula 2^{N-2} , donde N es la cantidad de bits que tenemos que robarle a la porción de host, adaptamos la máscara de red por defecto a la subred.

En este caso particular $2^{N-2} = 7$ (o mayor) ya que nos pidieron que hagamos 7 subredes.

$2^N - 2$	Redes	Máscara Binario	Máscara Decimal
$2^2 - 2$	2	1111 1111.1100 0000. 0000 0000. 00000000	255.192.0.0
$2^3 - 2$	6	1111 1111.1110 0000. 0000 0000. 00000000	255.224.0.0
$2^4 - 2$	14	1111 1111.1111 0000. 0000 0000. 00000000	255.240.0.0
$2^5 - 2$	30	1111 1111.1111 1000. 0000 0000. 00000000	255.248.0.0
$2^6 - 2$	62	1111 1111.1111 1100. 0000 0000. 00000000	255.252.0.0
$2^7 - 2$	126	1111 1111.1111 1110. 0000 0000. 00000000	255.254.0.0

Cuadro 7. Obtención de la máscara de la red clase A.

Una vez hecho el cálculo nos da que debemos robar 4 bits a la porción de host para hacer 7 subredes o más y que el total de subredes útiles va a ser de 14, es decir que van a quedar 7 para uso futuro.

Tomando la máscara Clase A por defecto, a la parte de red le agregamos los 4 bits que le robamos a la porción de host remplazándolos por "1" y así obtenemos 255.240.0.0 que es la máscara de subred que vamos a utilizar para todas nuestras subredes.

Máscara	Red	Hosts	N° de Barra
255.240.0.0	1111 1111.1111 0000. 0000 0000. 00000000		12

- b. Obtener Rango de Subredes y Cantidad de Hosts

Para obtener las subredes se trabaja únicamente con la dirección IP de la red, en este caso 10.0.0.0. Para esto vamos a modificar el mismo octeto de bits (el segundo) que modificamos anteriormente en la máscara de red pero esta vez en la dirección IP.

Para obtener el rango hay varias formas, la que me parece más sencilla a mí es la de restarle a 256 el número de la máscara de subred adaptada. En este caso sería: $256-240=16$, entonces 16 va a ser el rango entre cada subred. En el Cuadro 8 muestra cómo se pueden obtener las cantidades de subredes y de host en una subred clase A.

N° de Subred *	Rango IP **		Host Asignables por Subred
	Desde	Hasta	
1	10.0.0.0	10.15.255.255	1, 048,574
2	10.16.0.0	10.31.255.255	1, 048,574
3	10.32.0.0	10.47.255.255	1, 048,574
4	10.48.0.0	10.63.255.255	1, 048,574
5	10.64.0.0	10.79.255.255	1, 048,574
6	10.80.0.0	10.95.255.255	1, 048,574
7	10.96.0.0	10.111.255.255	1, 048,574
8	10.112.0.0	10.127.255.255	1, 048,574
9	10.128.0.0	10.143.255.255	1, 048,574
10	10.144.0.0	10.159.255.255	1, 048,574
11	10.160.0.0	10.175.255.255	1, 048,574
12	10.176.0.0	10.191.255.255	1, 048,574
13	10.192.0.0	10.207.255.255	1, 048,574
14	10.208.0.0	10.223.255.255	1, 048,574
15	10.224.0.0	10.239.255.255	1, 048,574
16	10.240.0.0	10.255.255.255	1, 048,574

* La primera y la última Subred no se asignan ya que contienen la dirección de Red y Broadcast de la red global

** La primera y la última Dirección IP no se asignan ya que contienen la dirección de Red y Broadcast de cada Subred

Cuadro 8. Obtención de subredes y cantidad de host por cada subred clase A

Si queremos calcular cuántos hosts vamos a obtener por subred debemos aplicar la fórmula 2^M-2 , donde M es el número de bits disponible en la porción de host y -2 es debido a que toda subred debe tener su propia dirección de red y su propia dirección de broadcast. En este caso particular sería:

$$2^{20}-2=1.048.574 \text{ host utilizables por subred}$$

2.5.3 Subneteo Manual de una Red Clase B

Dada la red Clase B **132.18.0.0/16** se nos pide que mediante subneteo obtengamos un mínimo de 50 subredes y 1000 hosts por subred.

- a. Adaptar la Máscara de Red por Defecto a Nuestras Subredes, ya que la máscara por defecto para la red 132.18.0.0 es:

Máscara	Red	Hosts	N° de Barra
255.255.0.0	1111 1111.1111 1111. 0000 0000. 00000000		16

Usando la fórmula 2^N-2 , donde N es la cantidad de bits que tenemos que robarle a la porción de host, adaptamos la máscara de red por defecto a la subred.

En este caso particular $2^N-2=50$ (o mayor) ya que necesitamos hacer 50 subredes.

$2^N - 2$	Redes	Máscara Binario	Máscara Decimal
$2^2 - 2$	2	1111 1111.1111 1111.1100 0000.0000 0000	255.255.192.0
$2^3 - 2$	6	1111 1111.1111 1111.1110 0000.0000 0000	255.255.224.0
$2^4 - 2$	14	1111 1111.1111 1111.1111 0000.0000 0000	255.255.240.0
$2^5 - 2$	30	1111 1111.1111 1111.1111 1000.0000 0000	255.255.248.0
$2^6 - 2$	62	1111 1111.1111 1111.1111 1100.0000 0000	255.255.252.0
$2^7 - 2$	126	1111 1111.1111 1111.1111 1110.0000 0000	255.255.254.0
$2^8 - 2$	254	1111 1111.1111 1111.1111 1111.0000 0000	255.255.255.0
$2^9 - 2$	510	1111 1111.1111 1111.1111 1111.1000 0000	255.255.255.128
$2^{10} - 2$	1022	1111 1111.1111 1111.1111 1111.1100 0000	255.255.255.192

Cuadro 9. Obtención de la máscara de la red clase B.

El cálculo nos da que debemos robar 6 bits a la porción de host para hacer 50 subredes o más y que el total de subredes útiles va a ser de 62, es decir que van a quedar 12 para uso futuro. Entonces a la máscara Clase B por defecto le agregamos los 6 bits robados reemplazándolos por "1" y obtenemos la máscara adaptada 255.255.252.0.

Máscara	Red	Hosts	N° de Barra
255.255.0.0	1111 1111.1111 1111.1111 1100. 00000000		22

- b. Obtener Cantidad de Hosts y Rango de Subredes

Como también nos piden una cantidad específica de hosts, 1000 hosts por subred, deberemos verificar que sea posible obtenerlos con la nueva máscara. Para eso utilizamos la fórmula 2^M-2 , donde M es el número de bits disponible en la porción de host y -2 es debido a que toda subred debe tener su propia dirección de red y su propia dirección de broadcast

$$2^{10} - 2 = 1022 \text{ hosts por subred}$$

Para obtener las subredes se trabaja con la dirección IP de la red, en este caso 132.18.0.0 modificando el mismo octeto de bits (el tercero) que modificamos en la máscara de red pero esta vez en la dirección IP.

Para obtener el rango hay varias formas, la que me parece más sencilla a mí es la de restarle a 256 el número de la máscara de subred adaptada. En este caso sería: $256 - 252 = 4$, entonces 4 va a ser el rango entre cada subred. En el gráfico solo puse las primeras 10 subredes y las últimas 5 porque iba a quedar muy largo, pero la dinámica es la misma. En el Cuadro 10 muestra cómo se pueden obtener las cantidades de subredes y de host en una subred clase B.

N° de Subred *	Rango IP **		Host Asignables por Subred
	Desde	Hasta	
1	132.18.0.0	132.18.3.255	1,022
2	132.18.4.0	132.18.7.255	1,022
3	132.18.8.0	132.18.11.255	1,022
4	132.18.12.0	132.18.15.255	1,022
5	132.18.16.0	132.18.19.255	1,022
60	132.18.236.0	132.18.239.255	1,022
61	132.18.240.0	132.18.243.255	1,022
62	132.18.244.0	132.18.247.255	1,022
63	132.18.248.0	132.18.251.255	1,022
64	132.18.252.0	132.18.255.255	1,022
* La primera y la última Subred no se asignan ya que contienen la dirección de Red y Broadcast de la red global			
** La primera y la última Dirección IP no se asignan ya que contienen la dirección de Red y Broadcast de cada host			

Cuadro 10. Obtención de subredes y cantidad de host por cada subred clase B

Si queremos calcular cuántos hosts vamos a obtener por subred debemos aplicar la fórmula $2^M - 2$, donde M es el número de bits disponible en la porción de host y -2 es debido a que toda subred debe tener su propia dirección de red y su propia dirección de broadcast. En este caso particular sería:

$$2^{10} - 2 = 1,022 \text{ host utilizables por subred}$$

Nota: Para el caso de la Clase C, se tendrían que tomar en cuenta los mismos pasos tomados en dos clases anteriormente mencionadas.

2.5.4 Ejemplos de subneteo

A. Su red utiliza la dirección IP 172.30.0.0/16. Inicialmente existen 25 subredes con un mínimo de 1000 hosts por subred. Se proyecta un crecimiento en los próximos años de un total de 55 subredes. ¿Qué máscara de subred se deberá utilizar?

Solución:

Para encontrar el número de subredes proyectadas al crecimiento de la empresa, se determinará el número de bits que serán usado en la nueva máscara por lo que:

$2^N - 2$, donde N es el número de bits que se estarán tomando para la nueva máscara de red.

$2^6 - 2 = 62$, que serán el número hábiles de subredes que se estarán usando, ya que la primer y la última subred no son hábiles (dirección de Red y Broadcast de la red global).

La máscara de red original es: 255.255.0.0, es decir contendrá 16 bits, a esta se le agregaran los 6 bits que se han tomado para la nueva máscara, por lo que la nueva máscara quedaría:

Decimal	255.255.252.0
Binario	1111 1111. 1111 1111. 1111 1100 .0000 0000

El resto de bits no habilitados son los que se estarán usando para la asignación de los host, donde deben de haber como mínimo 1000 por cada subred.

$2^N - 2$, donde N es el número de bits que se estarán tomando los números de host.

$2^{10} - 2 = 1022$, que serán el número hábiles de host (La primera y la última Dirección IP no se asignan ya que contienen la dirección de Red y Broadcast de cada host).

Por lo tanto la máscara de subred que se deberá utilizar será: 255.255.252.0

B. Una red está dividida en 8 subredes de una clase B. ¿Qué máscara de subred se deberá utilizar si se pretende tener 2500 host por subred?

Solución:

Existen $2^4 - 2 = 14$ subredes validadas, como es de clase B la dirección IP por defecto es:

Decimal	255.255.0.0
Binario	1111 1111. 1111 1111.0000 0000. 0000 0000

A la máscara de red se le agregarán los cuatros bits que serán tomados para obtener las subredes que pide el enunciado; por lo tanto nuestra nueva máscara de red quedaría de la siguiente manera:

Decimal	255.255.240.0
Binario	1111 1111. 1111 1111. 1111 0000. 0000 0000

Se necesita comprobar si el número de bits restante nos podrán alcanzar para asignar los 2500 hosts por subred.

La cantidad de bits restante son 12, por lo tanto $2^N - 2$, donde N será nuestra cantidad de bits que quedaron libres en la nueva máscara.

$2^{12} - 2 = 4094$, que será la cantidad de hosts que habrá por subred.

C. ¿Cuáles de las siguientes subredes no pertenece a la misma red si se ha utilizado la máscara de subred 255.255.224.0?

- a) 172.16.66.24
- b) 172.16.65.33
- c) 172.16.64.42
- d) 172.16.63.51**

Solución:

Para encontrar la cual es la subred que no pertenece a la misma red se hará uso de la compuerta AND entre las cada una de las direcciones IP y la máscara de subred, la cual es de clase B.

- a) 172.16. 66. 24

1010 1100.0000 1000.0100 0010.0001 1000 Dirección de Subred

1111 1111.1111 1111.1110 0000.0000 0000 Dirección de máscara

1010 1100.0000 1000.0100 0000.0000 0000 Dirección de red

Dirección de red: **172.16.64.0**, por lo tanto la dirección de subred pertenece al rango de dirección de red.

- b) 172.16.65.33

1010 1100.0000 1000.0100 0001.0010 0001 Dirección de Subred

1111 1111.1111 1111.1110 0000.0000 0000 Dirección de mascara

1010 1100.0000 1000.0100 0000.0000 0000 Dirección de red

Dirección de red: **172.16.64.0**, por lo tanto la dirección de subred pertenece al rango de dirección de red.

- c) 172.16.64 42

1010 1100.0000 1000.0100 0000.0010 1010 Dirección de Subred

1111 1111.1111 1111.1110 0000.0000 0000 Dirección de máscara

1010 1100.0000 1000.0100 0000.0000 0000 Dirección de red

Dirección de red: **172.16.64.0**, por lo tanto la dirección de subred pertenece al rango de dirección de red.

d) 172.16.63.51

1010 1100.0000 1000.0011 1111.0011 0011 Dirección de Subred

1111 1111.1111 1111.1110 0000.0000 0000 Dirección de máscara

1010 1100.0000 1000.0010 0000.0000 0000 Dirección de red

Dirección de red: **172.16.32.0**, por lo tanto la dirección de subred **NO** pertenece al rango de dirección de red.

D. Se tiene una dirección IP 172.17.111.0 máscara 255.255.254.0, ¿Cuántas subredes y cuantos host validos habrá por subred?

Solución:

La dirección IP es de clase B, por lo tanto la máscara de red en binario será: 1111 1111.1111 1111.**1111 1110**.0000 0000

Hay 7 bits (los que están a 1 en el tercer octeto) para conocer la cantidad de subredes válidas; por lo tanto:

$$2^N - 2 = 2^7 - 2 = 126 \text{ subredes válidas}$$

Hay 9 bits (el cero que resta del tercer octeto y los ochos del cuarto octeto) para conocer la cantidad de direcciones válidas para hosts; por lo tanto:

$$2^N - 2 = 2^9 - 2 = 510 \text{ hosts válidos}$$

En la dirección 172.17.111.0/23 existen 126 subredes con 510 hosts cada subred.

E. A partir de la dirección IP 192.168.85.129 con máscara 255.255.255.192, ¿Cuál es la dirección de subred y de broadcast a la que pertenece el host?

192.168.85. 129

1100 0000.1010 1000.0101 0101.1000 0001 Dirección de Subred

1111 1111.1111 1111.1111 1111.1100 0000 Dirección de máscara

1100 0000.1010 1000.0101 0101.1000 0000 Dirección de red

Dirección de red: 192.168.85.128

Dirección Broadcast: 10111111

Con lo anterior se toma el último octeto de la dirección de red ya que esta dirección es clase C y se deja intacto en el último octeto los 2 primeros bits ya que en la máscara son unos (1), se completa el resto de ceros (0) con unos (1) y de esta forma sabemos que la dirección de Broadcast es **192.168.85.191** para esta subred.

Por lo tanto la dirección de sería: 192.168.85.128, y la dirección broadcast sería: 192.168.85.191

2.5.5 Ejercicios Propuestos

¿Cuál es la dirección de red y difusión que corresponde a la IP 12.252.255.48 con una máscara 255.255.255.248?

- a. 12.252.255.255
- b. 12.252.255.52
- c. 12.252.255.40
- d. 12.252.255.7
- e. 12.255.255.255

1. Un host posee una dirección IP 192.168.100.50 con una máscara de 255.255.255.192 ¿Cuál sería su dirección de difusión a la subred a la que este pertenece?

- a. 192.168.100.63
- b. 192.168.100.255
- c. 192.168.100.62
- d. 192.168.100.100
- e. 192.168.255.255

2. El administrador de redes de una institución, tiene provisto un crecimiento de 5 subredes por año con ID clase C en los próximos 5 años, si actualmente la institución posee 15 subredes ¿Qué máscara de red deberá usar el administrador de la institución para lograr hacer las asignaciones previstas?

- a. 255.255.255.224
- b. 255.255.255.192
- c. 255.255.255.240
- d. 255.255.255.255
- e. 255.255.255.252

3. ¿Cuál es el número máximo de subredes y host por subred asignables que tendrá Juan, si a su empresa tiene la dirección 172.16.0.0 y la máscara de subred es 255.255.240.0?

- a. No podrá asignar, ya que la máscara es inválida para esa dirección
- b. 16 subredes y 4096 host por cada subred
- c. 16 subredes y 4094 host por cada subred
- d. 14 subredes y 4092 host por cada subred
- e. 14 subredes y 4094 host por cada subred

4. La dirección 192.168.15.24/27 está siendo dividida en subredes ¿Cuál es la cantidad de subredes y hosts por subredes son los que se obtendrán?
 - a. No podrá asignar, ya que la máscara es inválida para esa dirección
 - b. 6 subredes y 30 host por cada subred
 - c. 6 subredes y 32 host por cada subred
 - d. 8 subredes y 32 host por cada subred
 - e. 8 subredes y 30 host por cada subred

5. Teniendo una dirección 172.123.42.2 se necesitan crear 120 subredes con 500 hosts por cada subred ¿Qué máscara de subred se deberá utilizar?
 - a. /22
 - b. /24
 - c. /25
 - d. /23
 - e. /20

6. Dada la dirección IP de una subred 195.106.14.0/24, ¿Cuál es el número total de redes y el número total de nodos por red que se obtiene?
 - a. 1 red con 254 nodos
 - b. 2 redes con 128 nodos
 - c. 4 redes con 64 nodos
 - d. 6 redes con 30 nodos
 - e. 8 redes con 32 nodos

7. La red 192.141.27.0/24, se ha subnetado en 16 subredes. Es decir, la nueva máscara de subred para las 16 subredes nuevas es un /28 o 255.255.255.240. Cada una de estas 16 subredes tendrá 14 direcciones útiles (es decir, asignables a interfaces de hosts o routers). Dadas las siguientes direcciones IP de hosts, identifique cuáles direcciones son válidas (elija 3).
 - a. 192.141.27.33
 - b. 192.141.27.112
 - c. 192.141.27.119
 - d. 192.141.27.126
 - e. 192.141.27.175
 - f. 192.141.27.208

8. Se tiene la siguiente dirección IP de un host con su máscara: IP 172.16.10.22 / 255.255.255.240 ¿Cuál de los siguientes es el rango de nodo válido para la dirección de la subred a la que pertenece dicho host?
- a. 172.16.10.20 a 172.16.10.22
 - b. 172.16.10.1 a 172.16.10.255
 - c. 172.16.1.16 a 172.16.10.23
 - d. 172.16.10.17 a 172.16.10.31
 - e. 172.16.10.17 a 172.16.10.30
9. ¿Cuál es la dirección de broadcast de la dirección de subred 192.168.99.20 /255.255.255.252?
- a. 192.168.99.127
 - b. 192.168.99.63
 - c. 192.168.99.23
 - d. 192.168.99.31
 - e. 192.168.99.30
10. Un host tiene la dirección IP 176.32.65.32 / 20. ¿Cuál de las siguientes direcciones de host, todos con la misma máscara /20, NO pertenecen a la misma subred del host antes citado? Pista: Debe primero calcular la dirección de red del host citado inicialmente.
- a. 176.32.66.15
 - b. 176.32.80.22
 - c. 176.32.77.130
 - d. 176.32.65.18
 - e. 176.32.78.24
11. Un host tiene la dirección IP 160.126.129.131 / 22 (Clase B). ¿Cuál es la dirección de red y de difusión a la que pertenece dicho host?
- a. La dirección de red es: 160.126.129.0 Broadcast es: 160.126.129.255
 - b. La dirección de red es: 160.126.129.0 Broadcast es: 160.126.130.255
 - c. La dirección de red es: 160.126.128.0 Broadcast es: 160.126.131.255
 - d. La dirección de red es: 160.126.128.0 Broadcast es: 160.126.135.255
 - e. Ninguna de las anteriores.
12. Dada la siguiente dirección IP y máscara: 176.32.65.32 / 24. ¿Cuál de las siguientes afirmaciones es verdadera?
- a. La dirección de red es 176.32.65.0 y hay 4096 direcciones útiles en dicha subred.
 - b. La dirección de red es 176.32.64.0 y hay 4096 direcciones útiles en dicha subred.

- c. El rango efectivo de la subred en la que se ubica este host va de 176.32.65.1 a 176.32.79.254, lo cual permite 4096 direcciones utilizables para esta subred.
 - d. El rango efectivo de la subred en la que se ubica este host va de 176.32.64.1 a 176.32.79.254, lo cual permite 4096 direcciones utilizables para esta subred.
 - e. El rango efectivo de la subred en la que se ubica este host va de 176.32.65.0 a 176.32.79.254, lo cual permite 4094 direcciones utilizables para esta subred.
13. Si un nodo de una red tiene la dirección 172.16.45.14/30, ¿Cuál es la dirección de la subred y difusión a la cual pertenece ese nodo?
- a. Subred 172.16.45.12 Difusión 172.16.45.15
 - b. Subred 172.16.45.12 Difusión 172.16.45.14
 - c. Subred 172.16.45.0 Difusión 172.16.45.3
 - d. Subred 172.16.45.4 Difusión 172.16.45.7
 - e. Subred 172.16.45.8 Difusión 172.16.45.11
 - f. Subred 172.16.45.18 Difusión 172.16.45.21
 - g. Subred 172.16.0.0 Difusión 172.16.0.3
14. Usted se encuentra trabajando en una empresa a la que le ha sido asignada una dirección clase C y se necesita crear 10 subredes. Se le requiere que disponga de tantas direcciones de nodo en cada subred, como resulte posible. ¿Cuál de las siguientes es la máscara de subred que deberá utilizar?
- a. 255.255.255.192
 - b. 255.255.255.224
 - c. 55.255.255.240
 - d. 255.255.255.248
 - e. 255.255.255.242
 - f. Ninguna de las anteriores.
15. ¿Cuántos hosts efectivos pueden haber en la subred 170.16.200.128 con máscara 255.255.255.128?
- a. 128
 - b. 64
 - c. 127
 - d. 126
 - e. 256
 - f. 254

16. Su empresa tiene asignada la dirección clase B 172.12.0.0 / 16. De acuerdo a las necesidades planteadas, esta red debería ser dividida en subredes que soporten un máximo de 459 nodos por subred, procurando mantener en su máximo el número de subredes disponibles ¿Cuál es la máscara que deberá utilizar?
- a. 255.255.254.0
 - b. 255.255.0.0
 - c. 255.255.128.0
 - d. 255.255.224.0
 - e. 255.255.248.0
 - f. 255.255.192.0
17. ¿Cuál es la dirección de difusión que corresponde a la IP 10.254.255.19 /255.255.255.248?
- a. 10.254.255.23
 - b. 10.254.255.24
 - c. 10.254.255.255
 - d. 10.255.255.255
 - e. 10.254.255.7
18. ¿A cuál de las redes siguientes podría pertenecer un host que tuviera la dirección 147.156.135.191?
- a. 147.156.135.128 / 26
 - b. 147.156.135.128 / 25
 - c. 147.156.135.64 / 25
 - d. 147.156.135.192 / 26
 - e. Ninguna de las anteriores.

3. VLANS

Una VLAN (Virtual LAN) es un método de crear redes lógicamente independientes dentro de una misma red física; es una red de ordenadores que se comportan como si estuviesen conectados al mismo cable, aunque pueden estar en realidad conectados físicamente en diferentes segmentos de una red de área local.

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

3.1 Tipos de VLANs

3.1.1 VLANs de Puerto Central

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

3.1.2 VLANs Estáticas

Los puertos de los Switches están ya pre asignados a las estaciones de trabajo.

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

3.1.3 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 DVLAN 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. En la Figura 7 y Figura 8, se muestra como está organizada normalmente una red sin el uso de VLANs.

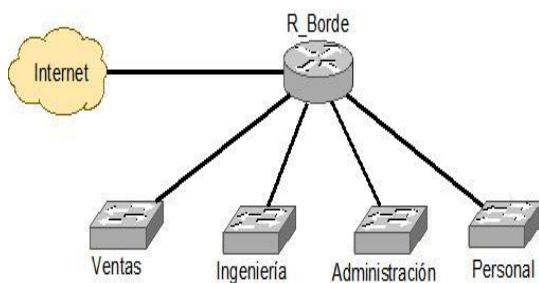


Figura 7. Diseño tradicional de una red

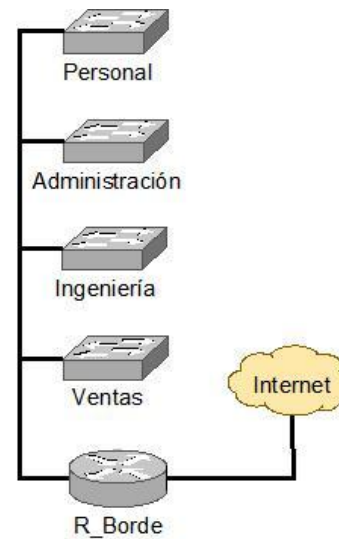


Figura 8. Esquematización de la red en un edificio

Con respecto a los ejemplos de la Figura 7 y Figura 8, no existe problema alguno si cada piso se designa a su departamento correspondiente; sin embargo, la mayoría de las veces se hace necesario por algún motivo especial destinar uno o más hosts a algún piso no correspondiente tal y a como se muestra la Figura 9 y Figura 10.

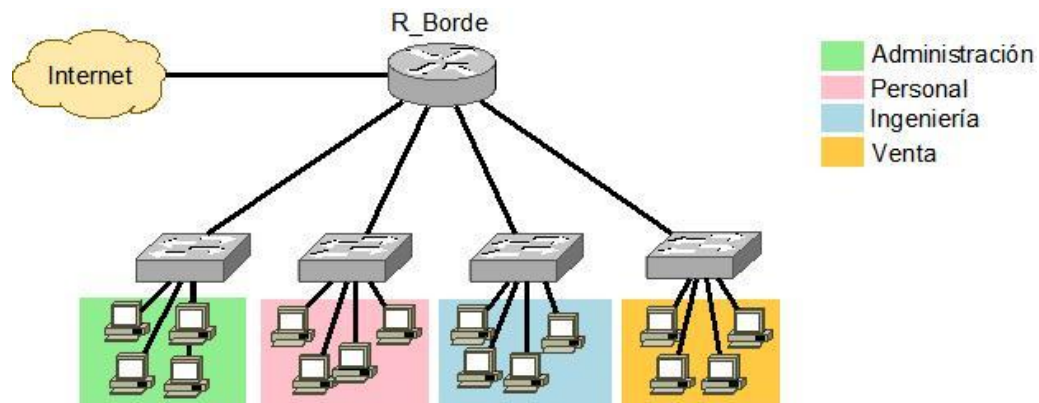


Figura 9. Diseño de una red usando VLANs

Las VLANs resuelven el problema de tener que conectar cables o dispositivos adicionales en un esquema como en el del edificio. Con las VLANs implementadas, se crean cuatro dominios broadcast (Ventas, Ingeniería, Administración y Personal) que son completamente independiente de la tecnología física de la red.

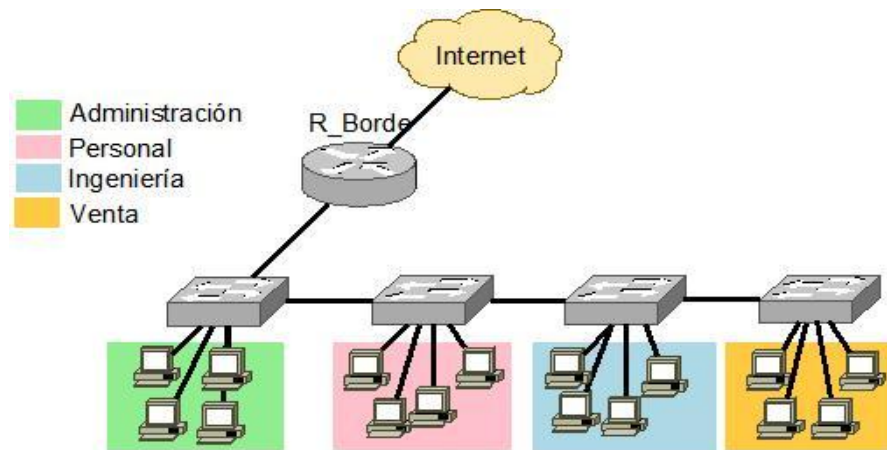


Figura 10. Cambio en el diseño físico de una red usando VLANs

3.2 Tipos de Puertos en la VLANs

3.2.1 Access Port

Solamente una VLAN puede estar asociada a este puerto. Los paquetes enviados a través de este puerto nunca serán etiquetados.

3.2.2 Trunk Port

Todas las VLANs deberán estar asociadas a este tipo de puerto. Los paquetes enviados a estos puertos siempre deberán ser etiquetados.

3.3 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 las VLANs.

- El enlace troncal transporta el tráfico para varias VLAN.
- 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 Switches, 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.

VTP es un protocolo de mensajería de capa 2 que mantiene la coherencia de la configuración VLAN a través de un dominio de administración común, gestionando las adiciones, supresiones y cambios de nombre de las VLAN a través de las redes. Un dominio VTP son varios Switches interconectados que comparten un mismo entorno VTP. Cada Switch se configura para residir en un único dominio VTP, la Figura 11 ilustra ambos tipos de puertos.

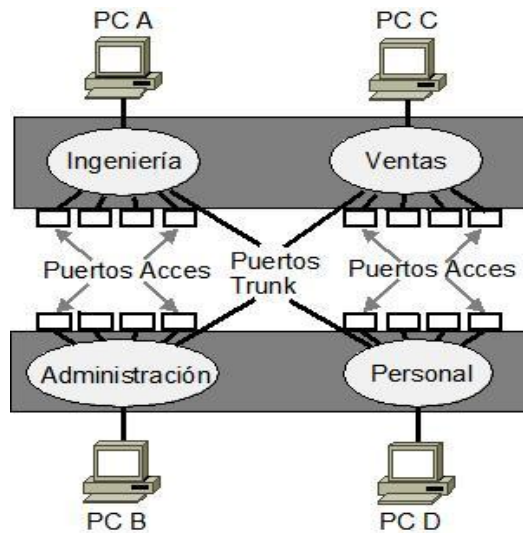


Figura 11. Puertos Trunk y Puertos Acces

3.4 Creación de VLANs

Hay tres pasos prácticos para la creación de VLANs que son:

3.4.1 Crear la VLAN

Para hacerlo, se tomarán en cuenta dos cosas:

- Darle el nombre a la VLAN la cual puede ser hasta de 32 caracteres
- Asignarle un ID, el cual puede ser de 2 a 4095.

3.4.2 Asignarles puertos a cada una de las VLANs

Con esto se le dice al Switch que puertos pertenecerán a cada dominio de broadcast creado por cada VLAN.

3.4.3 Crear el TRUNK

Esto se hace, solo en el caso de que se vayan a conectar más de un Switch que contengan las mismas o diferentes VLANs.

3.5 Configuración de VLANs

Las VLANs funcionan en el nivel 2 (nivel de enlace de datos) del modelo OSI. Sin embargo, los administradores suelen configurar las VLANs como correspondencia directa de una red o subred IP, lo que les da apariencia de funcionar en el nivel 3 (nivel de red).

Los administradores de red configuran las VLANs mediante software en lugar de hardware, lo que las hace extremadamente flexibles. Una de las mayores ventajas de las VLANs surge cuando se traslada físicamente algún ordenador a otra ubicación: puede permanecer en la misma VLAN sin necesidad de ninguna reconfiguración hardware.

En dispositivos Cisco, se usa VTP (VLAN Trunking Protocol) lo que permite definir dominios de VLAN, facilitando las tareas administrativas. VTP también permite «podar», lo que significa dirigir tráfico VLAN específico sólo a los conmutadores que tienen puertos en la VLAN destino. La necesidad de confidencialidad como así el mejor aprovechamiento del ancho de banda disponible dentro de la corporación ha llevado a la creación y crecimiento de las VLANs.

4. SPANNING TREE

4.1 Spanning Tree Tradicional

Un diseño potente de red no solo incluye la transferencia eficiente de paquetes o tramas, sino que también considera la forma de recuperarse rápidamente de fallos en la red. En un entorno de capa 3, los protocolos de enrutamiento en uso hacen un seguimiento de las rutas redundantes a la red de destino para que una ruta secundaria se pueda utilizar con rapidez en caso de que la ruta primaria falle. El enrutamiento de capa 3 permite muchas rutas redundantes activas hacia un destino y permite compartir la carga a través de múltiples caminos.

En un entorno de nivel 2 (conmutación o puenteo), sin embargo, no se utilizan protocolos de enrutamiento y las rutas redundantes activas no están permitidas. En cambio alguna forma de puente proporciona el transporte de datos entre las redes o puertos de conmutador. El protocolo Spanning Tree (STP) proporciona redundancia de enlaces de red, de modo que una red de conmutación de capa 2 puede recuperarse de los fallos sin ninguna intervención de una manera oportuna. El STP se define en el estándar IEEE 802.1D

Recordemos que un conmutador de capa 2 imita la función de un puente transparente. Un puente transparente debe ofrecer segmentación entre dos redes, sin dejar de ser transparente para todos los dispositivos finales conectados a él. Para el propósito de esta discusión considere un conmutador Ethernet de 2 puertos y sus similitudes con un puente transparente de 2 puertos.

Un puente transparente (y un switch Ethernet) opera de las siguientes maneras:

- El puente no tiene conocimiento inicial de la ubicación de cualquier dispositivo final, por lo tanto, el puente debe escuchar las tramas que llegan a cada uno de sus puertos y averiguar en qué red reside un dispositivo. La dirección de origen en una trama entrante es la clave de un dispositivo, el puente asume que el dispositivo fuente se encuentra detrás del puerto que llegó la trama. A medida de que el proceso de escucha continua, el puente se construye una tabla que contiene las direcciones MAC origen y los números de puerto del puente asociados con ellos.
- El puente o switch puede actualizar constantemente su tabla de puenteo al detectar la presencia de una nueva dirección MAC o tras la detección de una dirección MAC que ha cambiado de lugar de un puerto a otro puente. El puente puede entonces enviar tramas observando la dirección de destino, buscando la dirección en la tabla del switch o puente, y el envío de la trama a través del puerto donde se encuentra el dispositivo de destino.
- Si llega una trama con una dirección de destino que no se encuentra en la tabla del puente, el puente no puede determinar que puerto transmitirá la trama. Este tipo de trama se le conoce como unicast (unidifusión) desconocido. En este caso el puente trata la trama como si se tratara de una emisión y lo reenvía a todos los puertos restantes. Después que se escuchó una respuesta a esa trama, el puente aprende la ubicación de la estación desconocida y la añade a la tabla del puente para un uso futuro.
- Tramas reenviadas a través del puente no pueden ser modificadas por el propio puente. Por lo tanto, el proceso de puente es eficazmente transparente.

- Cualquier trama enviada ya sea a un destino conocido o desconocido, se remitirá a cabo el puerto o los puertos apropiados de modo que es probable que se reciba con éxito en el dispositivo final. La Figura 12 muestra un simple interruptor en funcionamiento con 2 puertos como un puente, reenvío de tramas entre 2 dispositivos finales. Sin embargo, este diseño de red no ofrece vínculos o rutas de acceso redundantes adicionales por si el interruptor o una de sus rutas de acceso fallara.

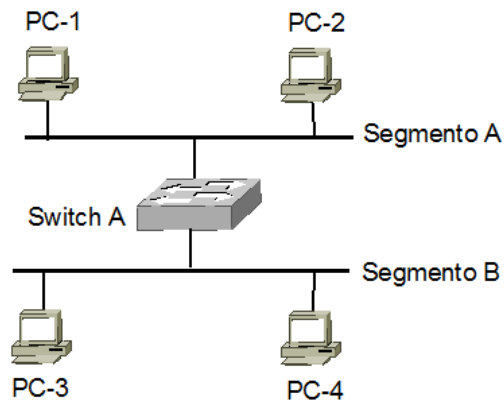


Figura 12. Puente transparente con un Switch

Para agregar algo de redundancia, puede agregar un segundo interruptor entre los dos segmentos de red originales, como se muestra en la Figura 13, ahora 2 Switches ofrecen la función de puente transparente en paralelo.

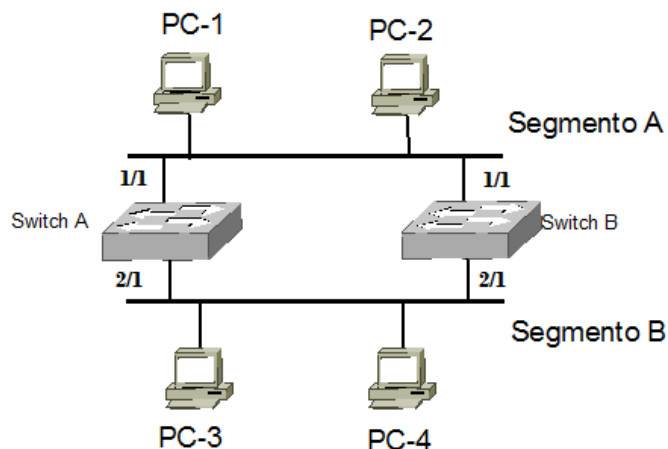


Figura 13. Puente redundante con 2 Switches

Considere lo que sucede cuando PC1 envía una trama a PC4. Por ahora supongamos que PC1 y PC4 son reconocidas por el switch que las conecta a cada una de ellas y sus direcciones MAC están registradas en la tabla de cada uno de los switch. PC1 envía la trama desde el segmento de red A, los interruptores A y B ambos reciben la trama en sus puertos 1/1. Debido a que PC4 ya es reconocido por ambos Switches la trama se reenvía a los puertos 2/1 en cada switch al segmento B. El resultado final es que PC4 recibe 2 copias de la trama de PC1. Esto no es lo ideal en una red

pero tampoco desastroso ya que este capítulo es para ilustrar al estudiante a que tenga una noción de las opciones para que no suceda lo inesperado.

Ahora consideremos el mismo proceso de reenviar una trama de PC1 a PC4. Esta vez sin embargo, ninguno de los Switches sabe nada de PC1 y PC4. PC1 envía la trama a PC4 y la secuencia de eventos es la siguiente:

- Tanto el interruptor A y B reciben la trama de PC1 por sus puertos 1/1. Debido a que la dirección MAC de PC1 aún no se ha visto o gravado, cada switch registra en su tabla de direcciones MAC la dirección de PC1 junto con el número de puerto de recepción 1/1. A partir de esta información ambos switches sabe que el PC1 reside en el segmento A.
- Debido a que la ubicación de PC4 es desconocido, los dos interruptores reenvían la trama a todos sus puertos disponibles (2/1) hacia el segmento B.
- Cada switch coloca una trama en su puerto disponible (2/1) del segmento B. PC4 que se encuentra en el segmento B recibe 2 tramas generadas por los 2 interruptores. Sin embargo el interruptor A escucha la trama remitida por el interruptor B, y el interruptor B escucha la trama remitida por A.
- El interruptor A escucha que la nueva trama es de PC1 para PC4. A partir de la tabla de direcciones del interruptor, se entera que PC1 fue escuchado en el puerto 1/1 por lo tanto pertenece al segmento A. sin embargo, la dirección de la fuente de PC1 acaba de ser escuchada en el puerto 2/1 del segmento B. Por definición el interruptor debe aprender la ubicación de PC1. (el conmutador B sigue el mismo procedimiento basándose en la trama recibida generada por A).
- En estos momentos ni el interruptor A ni B han aprendido la ubicación de PC4 debido a que no se han recibido los marcos con PC4 como dirección origen. Por lo tanto la trama debe ser enviada a todos los puertos disponibles en un intento de encontrar a PC4, cuando lo encuentra la trama de respuesta de PC4 se escucha en cada switch de manera que la dirección MAC de PC4 se grava en la tabla de direcciones de cada switch.
- Ahora la ubicación tanto de PC1 como PC4 ya son reconocidas por los interruptores y todas las tramas generadas hacia el segmento B o viceversa repetirán el mismo paso.

Este proceso de transmisión de una única trama alrededor y alrededor de 2 conmutadores se conoce como un circuito de puente. Ni un interruptor es consistente del otro, por lo que cada uno felizmente reenvía el mismo marco de ida y vuelta entre sus segmentos. Tenga en cuenta también que debido a que dos interruptores estén involucrados en el bucle, la trama original ha sido involucrada y ahora es enviado en dos bucles, lo que deja que la trama pueda ser reenviada para siempre.

Los bucles de puenteo se forman de la misma manera que se mencionó antes. Las tramas de difusión siguen circulando para siempre. Ahora los usuarios situados en ambos segmentos y los interruptores A y B procesan las tramas de

difusión. Este tipo de tormenta de broadcast puede hacer que la red colapse e impedir que los host finales tengan comunicación.

4.2 Prevención de bucles con el protocolo Spanning Tree

Los bucles de puente se forman debido a conmutadores en paralelo ya que ninguno es consciente del otro. STP fue desarrollado para superar la posibilidad de tener bucles en conmutadores en paralelo y las rutas redundantes podrían utilizarse para su beneficio. Básicamente el protocolo permite a los Switches a tomar conciencia de sí para que puedan negociar un camino libre de bucles a través de la red.

Los bucles se descubren antes de que estén disponibles para su uso, los enlaces redundantes se cierran para evitar la formación de bucles. En el caso de enlaces redundantes, los Switches pueden ser conscientes de que un enlace cerrado puede ser llevado rápidamente en caso de una falla del enlace. Esto se estará discutiendo en secciones posteriores en este capítulo.

STP se comunica con todos los interruptores conectados en una red. Cada Switch ejecuta el algoritmo de árbol de expansión **STA** basado en la información acerca de otros Switches vecinos. El algoritmo elige un punto de referencia en la red y calcula todas las rutas redundantes a ese punto de referencia. Cuando se detectan rutas redundantes el algoritmo **STA** toma un camino por el que transmitirá tramas y deshabilita el reenvío de las otras rutas redundantes.

En este momento la red de conmutación se encuentra libre de bucles. Sin embargo, si un puerto de reenvío falla o se desconecta, el algoritmo **STA** vuelve a calcular la topología para que los enlaces bloqueados se vuelvan a reactivar.

4.3 Comunicación Spanning Tree: Bridge Protocol Data Units (BPDU)

STP opera como un interruptor se comunica con otro. Los mensajes de datos se intercambian en unidades de datos de protocolo de puente (**BPDU**). Un interruptor envía una trama BPDU a un puerto, utilizando la dirección MAC única como una dirección de puerto de origen, el interruptor no tiene conocimiento de los otros a su alrededor. Por lo tanto, la trama BPDU tiene una dirección de destino conocida por STP como dirección de multidifusión 01:80:c2:00:00:00 para llegar a todos los interruptores de escucha.

4.3.1 Tipos de BPDU

4.3.1.1 BPDU de configuración

Que se utiliza para el cálculo del STP.

4.3.1.2 Notificación de cambio de topología (TCN) BPDU

Que se utiliza para anunciar los cambios en la topología de red.

El mensaje BPDU de configuración contiene los campos que se muestran en el Cuadro 11, la BPDU TCN se discute en la sección de cambio de topología más adelante en este capítulo.

Descripción	Número de bytes
ID de protocolo (siempre 0)	2
Versión (siempre 0)	1
Tipo de mensaje (Configuración o TCN BPDU)	1
Banderas	1
ID de puente raíz	8
Costo de ruta raíz	4
ID de puente de recepción	8
ID de puerto	2
Tiempo del mensaje	2
Máximo tiempo del mensaje	2
Tiempo de saludo	2
Retardo de reenvío	2

Cuadro 11. Contenido del mensaje de BPDU de configuración

El intercambio de mensajes BPDU trabaja con el objetivo de elegir a los puntos de referencia como base para una topología de árbol de expansión estable. Los bucles también pueden ser identificados y eliminados mediante la colocación de los puertos redundantes en un estado de bloque o espera. Tenga en cuenta que algunos campos de BPDU están relacionados con identificación de puente, coste de las rutas y valores de temporizador. Estos trabajan en conjunto para que la red de conmutadores pueda converger en una topología de árbol de expansión común y seleccionar los mismos puntos de referencia dentro de la red. Estos puntos de referencia definen en las secciones posteriores.

BPDU se envían a todos los puertos de un switch cada 2 segundos para que intercambiar información actual de la topología y los bucles se identifiquen rápidamente.

4.4 Elección del puente raíz

Para estar de acuerdo en una topología libre de bucles, un interruptor común de referencia debe de existir para utilizarlo como guía. Este punto de referencia se le denomina como puente raíz (cuando hablamos de puente nos referimos a Switch).

Un proceso de elección entre todos los interruptores conectados elige el puente raíz. Cada interruptor tiene un ID de puente único que identifica a otros Switches. El ID de puente es un valor de 8 bytes que consta de los siguientes campos:

4.4.1 Prioridad de puente (2 bytes)

La prioridad o peso de un interruptor en relación con todos los demás. El valor de la prioridad puede tener un valor de **0 a 65535** y por defecto **32768**.

4.4.2 Dirección MAC (6 bytes)

Esta dirección es codificada y única, donde el usuario no puede cambiarlo.

El proceso de elección se procede de la siguiente manera: Cada interruptor se inicia por el envío de BPDUs con un ID de puente raíz igual a su propio ID de puente y un ID de puente del remitente de su propio ID de puente. La identificación del puente del remitente simplemente le dice a los demás Switches que es el remitente real del mensaje BPDUs. (Después de que un puente raíz se elige, las BPDUs de configuración solo se envían por el puente raíz. Todos los otros puentes deben de transmitir las BPDUs añadiendo sus propios ID de puentes de remitentes al mensaje).

Los mensajes BPDUs se analizan para ver si se anuncia un mejor puente raíz. Un puente raíz se considera mejor cuando el valor de ID de puente raíz es más bajo que el otro. Sabiendo que este campo BPDUs (ID de puente raíz) se divide en 2 campos diferentes, a como lo es la dirección MAC y la prioridad. En este caso se analiza el valor de prioridad y si son iguales con las de otro switch se elige como raíz al que tenga la dirección MAC más baja.

Cuando el interruptor se entera que tiene un mejor ID de puente, este sustituye el ID de puente raíz anterior por el suyo propio en la BPDUs, quedando como puente raíz y empezando a anunciarse hacia los demás Switches.

Como un ejemplo, considere la pequeña red que se muestra en la Figura 14. Para simplificar, supongamos que cada switch tiene una dirección MAC de todos 0s con el último dígito hexadecimal igual a la etiqueta del interruptor.

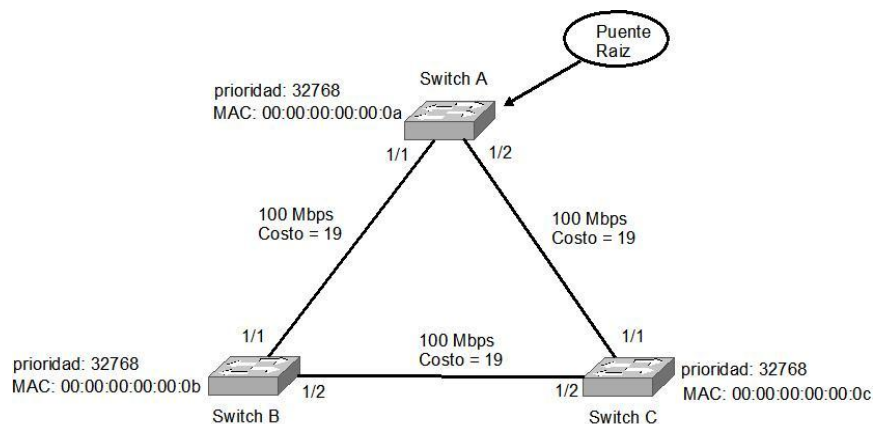


Figura 14. Ejemplo de elección de puente raíz

En esta red cada switch tiene una prioridad de puente por defecto de 32768. Los Switches están conectados con enlaces FastEthernet, tienen un costo de ruta por defecto de 19. Los tres Switches tratan de elegirse como puente raíz, pero todos ellos tienen un valor de prioridad de puente igual. La elección es determinada por la dirección MAC más baja, es este caso el switch A quedaría como raíz.

Una vez que tenemos elegido el puente raíz una opción y la más óptima utilizada en la actualidad es tener un respaldo raíz o puente raíz secundario, de manera que este servirá en algún momento como puente raíz si el principal falla. Las

formas de elegir un puente raíz es dejar por defecto que el STP lo elija, o alterando su prioridad de manera que si la opción es la última se espera que al switch afectado se elija ya sea como primario o secundario por defecto.

4.5 Eligiendo puertos raíz

Ahora que un punto de referencia ha sido nominado y elegido para toda la red conmutada, cada switch no raíz debe averiguar dónde se encuentra en relación con el puente raíz. Esta acción se puede realizar mediante la selección de un solo puerto raíz en cada switch no raíz.

STP utiliza el concepto de costos para determinar muchas cosas. Selección de un puerto de la raíz implica evaluar el coste de la ruta raíz. Este valor es el costo acumulado de todos los enlaces que conducen al puente raíz. Un interruptor de enlace en particular tiene un costo asociado con él, también, llamado el coste de la ruta. Para entender la diferencia entre estos valores, recuerda que sólo el coste de la ruta raíz se realiza dentro de la BPDU. (Ver Tabla 1 otra vez.) Como el coste de la ruta raíz viaja a lo largo, otros Switches pueden modificar su valor a que sea acumulativo. El coste de la ruta, sin embargo, no está contenido en el BPDU. Se sabe sólo para el conmutador local donde reside el puerto.

Costes de la ruta se define como un valor de 1 byte, con los valores por defecto que se muestran en la Tabla 2. En general, cuanto mayor sea el ancho de banda de un enlace, menor será el coste de transporte de datos a través de ella. El estándar original IEEE 802.1D define Coste de ruta como 1000 Mbps dividido por el ancho de banda de enlace Mbps. Estos valores se muestran en la columna central de la tabla. Las redes modernas utilizan comúnmente GigabitEthernet y OC-48 ATM, que son a la vez demasiado cerca de o mayor que el valor máximo de 1,000 Mbps. El IEEE ahora utiliza una escala no lineal de coste de la ruta, como se muestra en la columna de la derecha del Cuadro 12.

Ancho de banda de enlaces	Costo STP viejo	Costo STP nuevo
4 Mbps	250	250
10 Mbps	100	100
16 Mbps	63	62
45 Mbps	22	39
100 Mbps	10	19
155 Mbps	6	14
622 Mbps	2	6
1 Gbps	1	4
10 Gbps	0	2

Cuadro 12. Costo de ruta STP

El valor de coste de la ruta raíz se determina de la siguiente manera:

- El puente raíz envía una BPDU con un trazado de valor Costo Raíz de 0, ya que sus puertos se sientan directamente sobre el puente raíz.

- Cuando el vecino más cercano recibe la BPDU, añade el coste de la ruta de su propio puerto, donde llegó el BPDU. (Esto se hace como se recibe la BPDU.).
- El vecino envía BPDU con este nuevo valor acumulativo como el coste de la ruta raíz.
- Finalmente cada switch que recibe la BPDU tiene que añadir el valor de costo de puerto, de manera que se ira acumulando.

Después de incrementar el coste de la ruta raíz, un interruptor también registra el valor en su memoria. Cuando un BPDU se recibe en otro puerto y el nuevo coste de la ruta de la raíz es menor que el valor registrado anteriormente, este valor más bajo se convierte en el nuevo coste de la ruta raíz. Además, el costo más bajo indica que el interruptor de la ruta de acceso al puente raíz debe ser mejor usando este puerto de lo que era en otros puertos. El interruptor se ha determinado cuáles de sus puertos tiene la mejor ruta a la raíz.

La Figura 15, muestra el proceso de selección de puerto raíz:

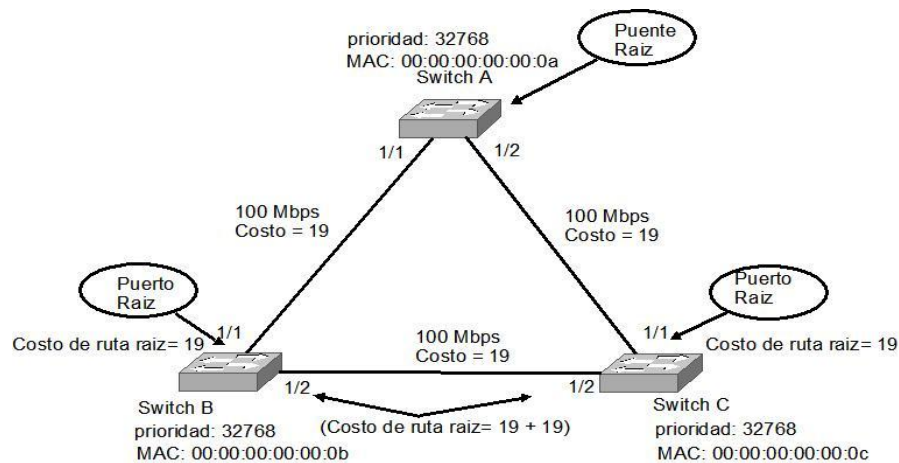


Figura 15. Ejemplo de elección de puerto raíz

El puente raíz, Switch A, ya ha sido elegido. Por lo tanto, cada switch en la red debe elegir un puerto que tiene el mejor camino para el puente raíz. Switch B selecciona el puerto 1/1, con un coste de la ruta raíz de 0 más 19. El puerto 1/2 no se elige porque su coste de la ruta de la raíz es 0 (BPDU de switch A) más 19 (Coste de ruta de enlace AC) más 19 (Coste de ruta de enlace CB), o un total de 38. Switch C toma una decisión similar de puerto 1/1.

4.6 Elección de puertos designado

A estas alturas, usted debe comenzar a ver el proceso de desarrollo: un punto de partida o punto de referencia ha sido identificado, y cada switch "conecta" a sí mismo hacia el punto de referencia con el único enlace que tiene el mejor camino. Una estructura de árbol está comenzando a emerger, pero los vínculos se han identificado sólo en este punto. Todos los enlaces están todavía conectados y pueden ser activos, dejando bucles de puenteo.

Para eliminar la posibilidad de bucles de puenteo, STP hace un cómputo final para identificar un puerto designado para cada segmento de la red. Supongamos que dos o más interruptores tienen puertos conectados a un solo segmento de

red común. Si aparece un marco en ese segmento, todos los puentes intentan para alejar a su destino. Recordemos que este comportamiento era la base de un circuito de puente y debe ser evitado.

En cambio, sólo uno de los enlaces en un segmento debe reenviar el tráfico desde y hacia ese segmento. Esta ubicación es el puerto designado. Los interruptores deben elegir un puerto designado en base al coste de la ruta raíz acumulada más baja hasta el puente raíz. Por ejemplo, un interruptor siempre tiene una idea de su propio coste de la ruta raíz, que anuncia en sus propias BPDUs. Si un conmutador vecino en un segmento LAN compartida envía una BDU anunciando un coste de la ruta raíz inferior, el vecino debe tener el puerto designado. Si un switch aprende solamente de los costos de ruta de raíz más altos de otros BDU recibidos en un puerto, sin embargo, a continuación, asume correctamente que su propio puerto de recepción es el puerto designado para el segmento.

Tenga en cuenta que todo el proceso de determinación de STP sólo ha servido para identificar los puentes y puertos. Todos los puertos siguen activos y bucles de puenteo aún podrían estar al acecho en la red. STP tiene un conjunto de estados progresivos que cada puerto tiene que pasar, independientemente del tipo o identificación. Estos estados prevenir activamente la formación de bucles y se describen en la siguiente sección.

En cada proceso de determinación discutido hasta ahora, dos o más enlaces tienen costos ruta raíz idénticos es posible. Esto se traduce en una condición de lazo, a menos que se consideran otros factores. Todas las decisiones STP se basan en la siguiente secuencia de cuatro condiciones:

- Menor ID de puente raíz.
- Menor coste de la ruta raíz hacia el puente raíz.
- Menor ID de puente del remitente.
- Menor ID de puerto del remitente.

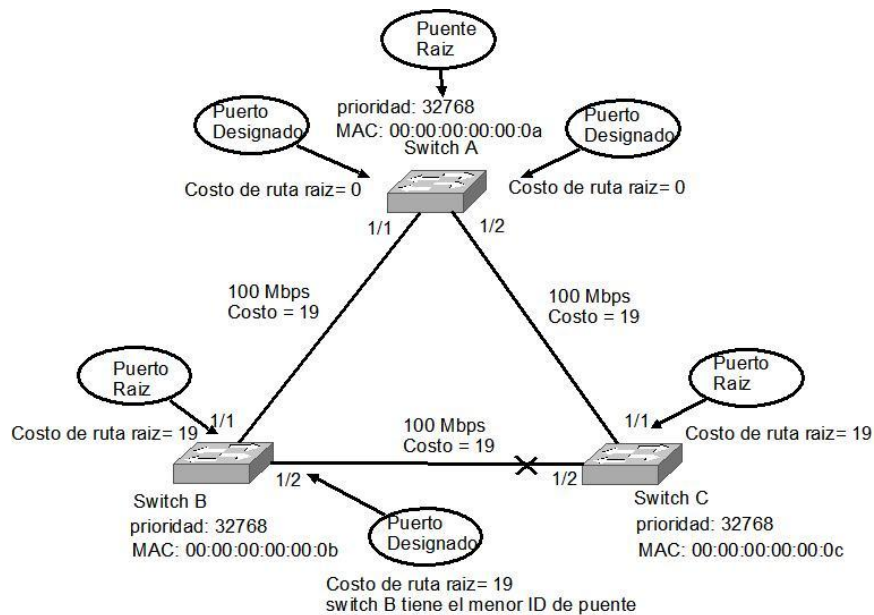


Figura 16. Ejemplo de elección de puerto designado

Los tres interruptores escogieron sus puertos designados (PD) por las siguientes razones:

- **Switch A:** Debido a este cambio es el puente raíz, todos sus puertos activos son Puertos designados por definición. En el puente raíz, el coste de la ruta raíz de cada puerto es 0.
- **Switch B:** Switch A puerto 1/1 es la DP para el segmento AB, ya que tiene la raíz más baja costo de ruta (0). Switch B puerto 1/2 es la DP para el segmento BC. El coste de la ruta raíz para cada extremo de este segmento es 19, determinada a partir de la BPDU entrante en el puerto 1/1. Debido a que el coste de la ruta de la raíz es igual en ambos puertos del segmento, el PD debe ser elegido por los siguientes criterios: el remitente ID de puente más bajo. Cuando el switch B envía una BPDU al switch C, que tiene la dirección MAC más baja en el ID de puente. El switch C también envía una BPDU al switch B, pero su ID de remitente puente es mayor debido a su MAC que es más alta. Por lo tanto, el puerto del switch B se selecciona como PD del segmento.
- **Switch C:** El puerto 1/2 del switch A es la PD para el segmento AC, ya que tiene el coste de la ruta raíz más bajo (0). El puerto 1/2 del switch B es la DP para el segmento B-C. Por lo tanto, el puerto 1/2 del switch C no será ni un puerto raíz ni un puerto designado. Como se discute en la siguiente sección, cualquier puerto que no es elegido a cualquiera de las posiciones entra en el estado de bloqueo, de manera que la red queda libre de bucles.

Recuerde que los papeles de un puerto pueden ser:

- Raíz.
- Designado
- No designado o bloqueado

4.7 Estados STP

Para participar en STP, cada puerto de un conmutador debe progresar a través de varios estados. Un puerto comienza su vida en un estado de movilidad reducida, moviéndose a través de varios estados pasivos y, por último, en un estado activo si se le permite reenviar el tráfico. Los estados STP de un puerto son los siguientes:

- **Deshabilitado:** Los puertos que están administrativamente cerrado por el administrador de la red o por el sistema debido a una condición de fallo, se encuentran en el estado desactivado. Este estado es especial y no es parte de la progresión normal de un puerto STP.
- **Bloqueo:** Después de que un puerto se inicia, comienza en el estado de bloqueo para que no se produzcan bucles de puenteo. En el estado de bloqueo, un puerto no puede recibir o transmitir datos y no se puede agregar las direcciones MAC a su tabla de direcciones. En cambio, un puerto se le permite recibir solamente BPDU para que el switch pueda saber de otros switches vecinos. Además, los puertos que se ponen en modo de espera para eliminar un bucle puente entran en el estado de bloqueo.
- **Escucha:** El puerto se traslada de bloqueo para escuchar solo si el interruptor piensa que el puerto se puede seleccionar como un puerto raíz o Puerto designado. En otras palabras, el puerto está en camino para comenzar el reenvío de tráfico. En el estado de escucha, el puerto todavía no puede enviar o recibir secuencias de datos. Sin embargo, el puerto se le permite recibir y enviar BPDU para que pueda participar activamente en el proceso de topología Spanning Tree. Este caso se da cuando el puerto raíz o designado falla.
- **Aprende:** Después de un período de tiempo denominado retardo de reenvío (Forward Delay) en el estado de escucha, se permite que el puerto para entrar en el estado de aprendizaje. El puerto aún envía y recibe BPDU como antes. Además, el cambio ahora puede aprender nuevas direcciones MAC para añadir a su tabla de direcciones. Esto le da al puerto de un período adicional de la participación pasiva y permite que el conmutador grave al menos alguna información en la tabla de direcciones.
- **Reenvío:** Después de otro periodo Forward Delay de tiempo en el estado de aprendizaje, se permite que el puerto pase al estado de reenvío. El puerto ya puede enviar y recibir tramas de datos, recopilar direcciones MAC en su tabla de direcciones, y enviar y recibir BPDU. El puerto es ahora un puerto del switch en pleno funcionamiento dentro de la topología del árbol de expansión.

Con lo que respecta a los puertos se le pueden configurar algunos parámetros con el objetivo de influir en las decisiones del STP.

Estos parámetros son el costo de puerto o ya sea la prioridad del puerto. Esto en el caso de que queramos que una ruta o un puerto siempre sean elegidos.

4.8 Temporizadores STP

STP opera con interruptores que envían BPDU el uno al otro en un esfuerzo para formar una topología libre de bucles. Las BPDU tienen una cantidad limitada de tiempo para viajar desde el interruptor para cambiar. Además, las noticias de un cambio de topología (tal como un enlace o el fracaso puente raíz) puede sufrir retardos de propagación como el anuncio se desplaza de un lado de una red a la otra.

STP utiliza tres contadores de tiempo para asegurarse de que una red converge correctamente antes de que pueda formar un bucle de puente. Los temporizadores y sus valores predeterminados son los siguientes:

- **Hello Time:** El intervalo de tiempo entre las BPDUs de configuración enviados por el puente raíz. El valor de Hello Time configurado en el conmutador de puente raíz determina el Hello Time para todos los switches que no son raíz, ya que sólo transmiten la BPDUs de configuración a medida que se reciben desde la raíz. Sin embargo, todos los interruptores tienen un tiempo de Hello configurado localmente, que se utiliza para tiempo TCN BPDUs cuando se retransmiten. El estándar IEEE 802.1D especifica un valor de tiempo Hello predeterminado de 2 segundos.
- **Forward Delay:** El intervalo de tiempo que un puerto de switch pasa tanto en los estados de escucha y aprendizaje. El valor predeterminado es de 15 segundos.
- **Max (máximo) Age:** El intervalo de tiempo en que un switch almacena una BPDUs antes de descartarlo. Durante la ejecución de la STP, cada puerto del switch mantiene una copia de la "mejor" BPDUs que ha escuchado. Si la fuente de la BPDUs pierde el contacto con el puerto del conmutador, el conmutador se da cuenta de que se produjo un cambio en la topología, después de que transcurra el tiempo y la edad máxima la BPDUs caduca. El valor Age Max predeterminado es de 20 segundos.

Los temporizadores de STP se pueden configurar o ajustar desde la línea de comandos switch. Sin embargo, los valores del temporizador no deben ser cambiados de los valores por defecto a la ligera. A continuación, los valores se deben cambiar sólo el interruptor puente raíz. Recordemos que los valores de los temporizadores se anuncian en los campos de la BPDUs. El puente raíz se asegura de que los valores de los temporizadores se propagan a todos los demás interruptores.

4.9 Cambio en la topología

Para anunciar un cambio en la topología de red activa, los interruptores envían una BPDUs TCN. El Cuadro 13 muestra el formato de estos mensajes.

Descripción	Número de bytes
ID de protocolo siempre (0)	2
Versión	1
Tipo de mensaje	1

Cuadro 13. Contenido de mensaje TCN BPDUs

Un cambio de topología se produce cuando un interruptor o bien se mueve de un puerto en el estado Forwarding o mueve un puerto de reenvío o de aprendizaje en el estado de bloqueo. El conmutador envía una BPDUs TCN cabo su puerto raíz por lo que, en definitiva, el puente raíz recibe la noticia del cambio de topología. Observe que el BPDUs TCN no lleva los datos sobre el cambio, pero informa a los beneficiarios sólo que se ha producido un cambio. Observe también que el interruptor se No envíe BPDUs TCN si el puerto se ha configurado con **PortFast** habilitado.

El cambio continúa enviando BPDUs TCN cada intervalo de tiempo de saludo hasta que se obtiene una reconocimiento por parte de un vecino. En cuanto a los vecinos recibe el BPDUs TCN, que se propagan en dirección al puente raíz. Cuando el puente raíz recibe la BPDUs, el puente raíz también envía un acuse de recibo.

4.10 Tipos de STP

Hasta ahora, este capítulo ha discutido STP en términos de su funcionamiento para evitar bucles y para recuperarse de los cambios de topología en forma oportuna. STP fue desarrollado originalmente para operar en un entorno de

puente, básicamente de apoyo una sola LAN (o una VLAN). Implementación de STP en un entorno conmutado ha requerido el examen y la modificación adicional para soportar múltiples VLANs. Debido a esto, el IEEE y Cisco se han acercado a STP diferente. En esta sección se examinan los tres tipos tradicionales de STP que se encuentran en las redes de conmutación y cómo se relacionan entre sí. No hay comandos de configuración específicos están asociados con los diversos tipos de STP. Más bien, se necesita un conocimiento básico de cómo interoperar en una red.

4.10.1 Spanning Tree común (CST)

El estándar IEEE 802.1Q especifica cómo se van a propagar las VLANs entre switches. También especifica una única instancia de STP para todas las VLAN. Esta instancia se conoce como el árbol de extensión común (CST). Todas las BPDU CST se transmiten a través de la VLAN nativa como las tramas sin etiquetar.

Tener un solo STP para muchas VLAN simplifica la configuración del switch y reduce la carga de la CPU del interruptor durante los cálculos de STP. Sin embargo, tener sólo una instancia STP puede causar limitaciones, también. Enlaces redundantes entre switches serán bloqueados sin capacidad para equilibrar la carga. Las condiciones también pueden ocurrir que haría que la transmisión en un enlace que no sea compatible con todas las VLAN, mientras que otros enlaces serían bloqueados.

4.10.2 Per-VLAN Spanning Tree (PVST)

Cisco tiene una versión propia de STP que ofrece más flexibilidad que la versión CST. Per-VLAN Spanning Tree (PVST) opera una instancia independiente de STP para cada VLAN individual. Esto permite que el STP en cada VLAN que se configura de forma independiente, ofreciendo un mejor rendimiento y el ajuste para condiciones específicas. Spanning Tree múltiples también hacen posible el equilibrio de carga sobre enlaces redundantes cuando los enlaces están asignados a diferentes VLAN.

Debido a su naturaleza propia, PVST requiere el uso de Cisco Inter-Switch Link (ISL) en la encapsulación trunking entre switches. En las redes donde TSVP y CST coexisten, se producen problemas de interoperabilidad. Cada uno requiere un método de concentración de enlaces diferentes, por lo que las BPDU no serán intercambiadas entre tipos de STP.

4.10.3 Per-VLAN Spanning Tree Plus (PVST+)

Cisco tiene una segunda versión propietaria de STP que permite a los dispositivos interactúan tanto con PVST y CST. Per-VLAN Spanning Tree Plus (PVST +) apoya efectivamente tres grupos de STP que opera en la misma red del campus.

Para ello, PVST+ actúa como un traductor entre grupos de interruptores y grupos de interruptores PVST CST. PVST+ se puede comunicar directamente con PVST utilizando troncos de ISL.

Se debe tener en cuenta que con el paso del tiempo han venido saliendo versiones nuevas de STP, cada una con diferentes características y mejoras a las respuestas al momento de converger. Estas nuevas versiones se muestran a continuación:

➤ **Propiedades de Cisco**

- PVST
- PVST+
- PVST+ rápido

➤ **Estándar IEEE**

- RSTP
- MSTP

4.11 Convergencia de enlaces redundantes

Algunos métodos adicionales que existen para permitir una convergencia más rápida STP en el caso de un fallo de enlace incluyen los siguientes:

- **PortFast:** Permite una conectividad rápida a establecerse en los puertos de switch de capa de acceso a puestos de trabajo que están iniciando.
- **UplinkFast:** Permite rápida recuperación de fallos de enlace ascendente en un switch de capa de acceso cuando dos uplinks están conectados en la capa de distribución.
- **BackboneFast:** Permite una rápida convergencia de la red troncal (core) después que se produce un cambio en la topología del árbol de expansión.

En lugar de modificar los valores de temporizador, estos métodos funcionan mediante el control de la convergencia en los puertos ubicados específicamente dentro de la jerarquía de la red.

4.12 Protección ante BPDU inesperadas en puertos de puente

Para mantener una topología eficiente, la colocación del puente raíz debe ser predecible. Con suerte, se ha configurado un conmutador para convertirse en el puente raíz y otro para ser la raíz secundaria. ¿Qué ocurre cuando un interruptor del "exterior" o invasor está conectado a la red, y que el interruptor de repente es capaz de convertirse en el puente raíz? Cisco añade dos características STP que ayudan a prevenir lo inesperado: **guardia raíz y BPDU**.

4.12.1 Root Guard

Supongamos que otro interruptor se introduce en la red con una Prioridad de puente que es más deseable (menor) que el puente raíz actual. El nuevo switch se convertiría entonces en el puente raíz y la topología STP podría re converger a una nueva forma. Esto es totalmente permitida por la STP, el interruptor con el ID de puente más bajo siempre gana las elecciones Root. Sin embargo, esto no siempre es deseable que, el administrador de la red, debido a que la nueva topología STP podría ser algo totalmente inaceptable. Esta tecnología se puede implementar en los puertos de switch donde posiblemente se puede conectar un nuevo puente.

4.12.2 BPDU Guard

Recordemos que la tradicional STP ofrece la función PortFast, donde se permite a los puertos del switch para entrar de inmediato al estado Forwarding tan pronto como la conexión funcione. Normalmente, PortFast proporciona acceso

a la red rápida de dispositivos de usuario final, en los que nunca se espera que los bucles de puente para formar. Incluso mientras PortFast está habilitado en un puerto, STP sigue funcionando y puede detectar un circuito puente. Sin embargo, un bucle puede ser detectado sólo en una cantidad finita de tiempo de la longitud de tiempo requerida para mover el puerto a través de los estados STP normal. Utilice BPDU guardia en todos los puertos del switch donde PortFast STP está habilitado. Esto impide cualquier posibilidad de un interruptor que se añade al puerto, ya sea intencionalmente o por error. Una aplicación obvia para BPDU está en los puertos de switch de capa de acceso donde los usuarios y dispositivos finales se conectan. BPDU normalmente no se espera que allí se detecta si un switch o hub estaba conectado inadvertidamente.

4.13 Protección ante la pérdida de BPDU

STP BPDU se utilizan como sondas para aprender acerca de una topología de red. Cuando los interruptores que participan en STP convergen en una topología libre de bucles común y coherente, BPDU todavía debe ser enviada por el puente raíz y retransmitida por cualquier otro cambio en el dominio STP. La integridad de la topología STP depende entonces de un flujo continuo y regular de BPDU desde la raíz.

¿Qué sucede si un interruptor no recibe BPDU en forma oportuna, o cuando no recibe ninguna en absoluto? El interruptor se puede ver que la condición aceptable, la topología debe haber cambiado, los puertos bloqueados por lo se puede desbloquear de nuevo. Si la ausencia de BPDU es en realidad un error, bucles de puenteo se pueden formar fácilmente. Cisco ha añadido tres características STP que ayudan a detectar o prevenir lo inesperado:

4.13.1 Detección de inclinación BPDU

Recordemos que el puente raíz es la fuente central de BPDU en un dominio STP y que las propaga cada 2 segundos (Hello Time). Después de interruptores aprenden el intervalo de tiempo hola de la raíz, que esperan recibir las BPDU en ese intervalo. Bajo algunas condiciones, un interruptor podría recibir una BPDU pero no puede transmitir por un período de tiempo. Por ejemplo, la CPU interruptor podría ser ocupado realizando otras tareas antes de la BPDU se puede procesar y propagado, o una cola de entrada en un interruptor podría ser tan llena que una BPDU entrante simplemente se pierde. Detección de inclinación BPDU mide la cantidad de tiempo que transcurre desde el momento que una BPDU se espera que cuando llegue en realidad. Esta diferencia de tiempo se denomina tiempo de inclinación. Detección de inclinación BPDU sólo detecta la condición y lo informa a través de mensajes Syslog. El interruptor no puede hacer nada más sobre la enfermedad.

4.13.2 Loop Guard

Supongamos que un puerto del conmutador está recibiendo BPDU, y el puerto del conmutador está en el estado de bloqueo. El puerto constituye una ruta redundante, sino que es el bloqueo, ya que no es ni un puerto raíz ni un puerto designado. Si, por alguna razón, ya no sale BPDU, conocido el último BPDU se mantiene hasta que se agote el tiempo máximo de edad. Entonces, BPDU que se vacía, y el interruptor piensa que ya no hay una necesidad de bloquear el puerto. El puerto mueve a través de los estados STP hasta que comienza a reenviar el tráfico y formar un bucle de puente. En su estado final, el puerto se convierte en un puerto designado. Para evitar esta situación, puede utilizar la

función de loop guard de STP. Cuando está habilitado, el lazo protector hace un seguimiento de la actividad en los puertos BPDU no designadas. Mientras que se reciben las BPDU, el puerto se le permite comportarse normalmente. Cuando BPDU se pierden, loop guardia mueve el puerto en el estado de bucle inconsistente. El puerto está bloqueando efectivamente en este punto para evitar la formación de un bucle y para mantenerlo en el papel no designadas. Después BPDU se reciben en el puerto nuevo, loop guardia permite que el puerto se mueva a través de los estados STP normal. De esta manera, el lazo guardia regula automáticamente los puertos sin necesidad de intervención manual.

4.13.3 UDLD

En una red, los interruptores se conectan entre sí por enlace bidireccional (full dúplex). Claramente, si un enlace tiene un problema de la capa física, los dos interruptores que se conectan detectan un problema y el enlace se muestra como no conectado. ¿Qué pasaría si un solo lado (recibir o transmitir) del enlace tuvo un fallo extraño? En algunos casos, los dos interruptores todavía verían un enlace funcional. Sin embargo, el tráfico podría ser entregado en una sola dirección, y ninguno de los interruptores se diera cuenta, de manera que si en un extremo no se recibe BPDU de un lado el switch hará cuenta de que ese puerto que está bloqueado no necesite permanecer en ese estado, de manera que se convierte en un puerto de reenvío y forme bucles en la red. Para evitar esta situación, puede utilizar el enlace de detección de características STP unidireccional (UDLD). Cuando se activa, UDLD supervisa interactivamente un puerto para ver si el enlace es verdaderamente bidireccional. El conmutador envía tramas de capa 2 UDLD especiales que identifican el puerto de conmutación a intervalos regulares. UDLD espera que el interruptor de extremo lejano eco de esos marcos de regreso a través del mismo enlace, con la identificación del puerto del conmutador de extremo lejano añadió. Si se recibe una trama UDLD a cambio, y los dos puertos vecinos están identificados en el marco, el enlace debe ser bidireccional. Sin embargo, si los marcos de eco no se ven, el enlace es unidireccional. En este caso el enlace pasa a un estado STP de error. Esta tecnología trabaja en 2 modos que son: normal y agresivo.

La Figura 17, muestra estas tecnologías de protección STP en uso y muestra cuales se pueden mezclar y cuáles no.

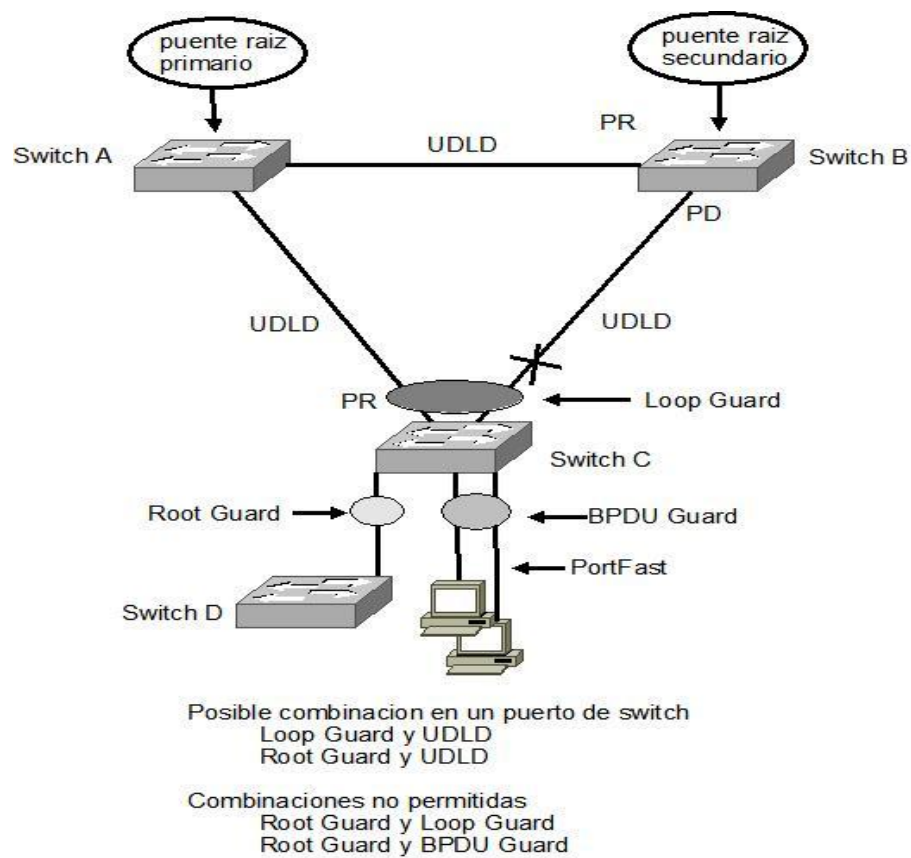


Figura 17. Aplicando protección STP

5. ETHERNET CHANNEL

5.1 EtherChannel

EtherChannel es una tecnología que fue inventada por Kalpana en los años 90, donde fue adquirido en 1994 por Cisco System. Cisco ahora ofrece este método de ancho de banda mediante la agregación o agrupación de enlaces en paralelo desde 2 hasta 8 eslabones de cualquier tipo ya sea FastEthernet (FEC), GigabitEthernet (GEC) respectivamente. Este paquete proporciona un ancho de banda full-dúplex de hasta 1600Mbps (8 enlaces de FastEthernet) o 16 Gbps (8 enlaces de GigabitEthernet).

Esto proporciona un medio fácil para crecer o ampliar la capacidad de un enlace entre 2 switches sin tener que comprar Hardware para la siguiente magnitud de producción. Por lo general si se tiene múltiples enlaces en paralelos entre 2 switches, crea la posibilidad de tener bucles de puente de una manera no deseable. EtherChannel evita esta situación mediante la agregación de los enlaces paralelos en un solo enlace lógico único que puede actuar ya sea como un enlace de acceso o un troncal. Los interruptores o dispositivos en cada extremo EtherChannel deben comprender y utilizar la tecnología para su correcto funcionamiento.

Aunque un enlace EtherChannel es visto como un único enlace lógico, el vínculo no tiene un ancho de banda total inherente igual a la suma de sus vínculos físicos que lo componen. Por ejemplo, supongamos que un enlace de FEC se compone de cuatro enlaces FastEthernet full-dúplex, 100-Mbps. Aunque es posible para el enlace de FEC para llevar a un rendimiento de 800 Mbps, el único enlace FEC resultante no funciona a esta velocidad. En cambio, el tráfico se equilibra entre los eslabones individuales de la EtherChannel. Cada uno de estos enlaces funciona a su velocidad inherente (200 Mbps full-duplex para FE), pero sólo transmite tramas que se le plantean por el hardware EtherChannel. El proceso de balanceo de carga se explica con más detalle en la siguiente sección.

EtherChannel también proporciona redundancia con varios enlaces físicos agrupados. Si uno de los enlaces en el paquete falla, el tráfico enviado a través de ese enlace se traslada a un enlace de al lado. Esta conmutación por error se produce en menos de unos pocos milisegundos y es transparente para el usuario final. Del mismo modo, como enlaces se restauran, se redistribuye la carga entre los enlaces activos.

5.1.1 Funcionamiento EtherChannel

Hay mucha gente que no conoce el concepto de EtherChannel y por culpa de este desconocimiento se pierde una gran ventaja que ofrece esta tecnología. Supongamos que tenemos la topología de la Figura 18.

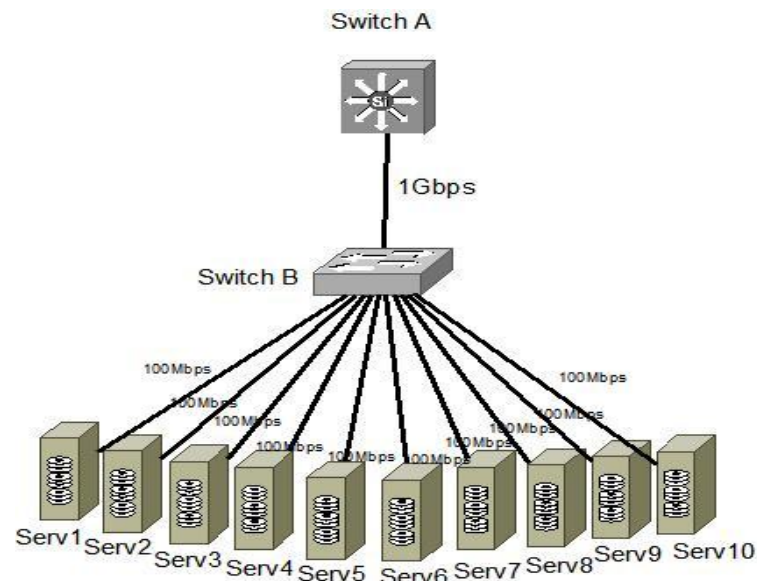


Figura 18. Funcionalidad sin implementación EtherChannel

Cuando tenemos una serie de servidores que salen por único enlace troncal, puede ser que el tráfico generado llegue a colapsar el enlace. Una de las soluciones más prácticas que se suele implementar en estos casos es el uso de EtherChannel. A como se mencionó antes, cuando generamos un EtherChannel lo que se hace es sumar la velocidad de los puertos que agregamos al enlace lógico obteniendo los siguientes resultados.

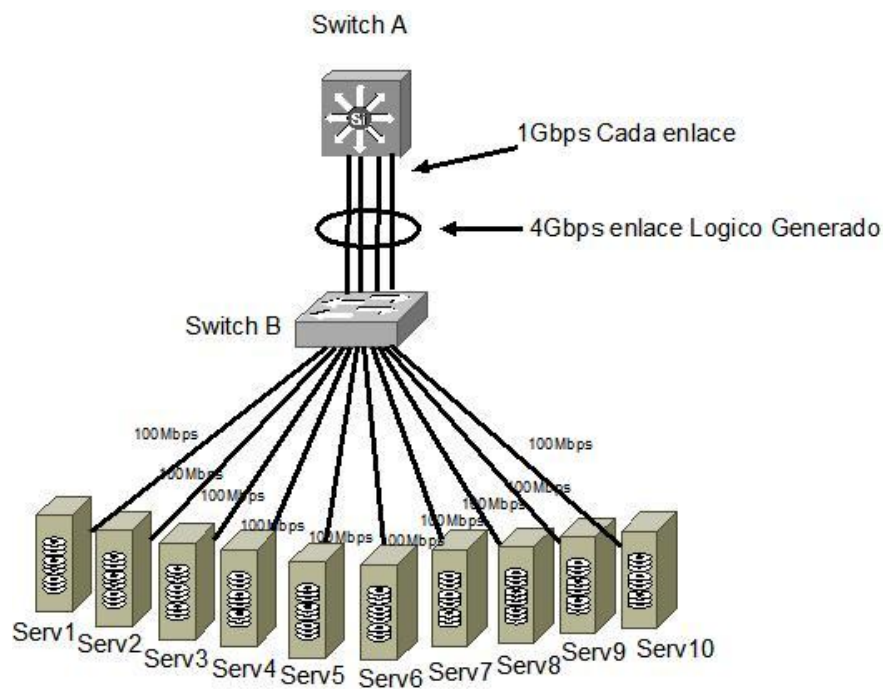


Figura 19. Funcionalidad con implementación EtherChannel

5.1.2 Requisitos de configuración

Algunas restricciones de configuración existen para asegurar que los enlaces agrupados se encuentran configurados de manera similar. Alguno de estos requisitos mencionados a continuación:

- Una de estas restricciones es que todos los puertos deben permanecer en la misma VLANs.
- Si se utiliza como un puerto troncal, todos los puertos agrupados deben de estar en modo troncal, de manera que tengan tanto la VLANs nativa como el resto de VLANs existente.
- Cada uno de los puertos debe tener la misma velocidad y ser full-duplex antes de ser empaquetados.

5.1.3 Distribución de tráfico en EtherChannel

El tráfico en un EtherChannel se distribuye a través de los enlaces que componen el paquete de una manera determinista. Sin embargo, la carga no es necesariamente igualmente equilibrada en todos los enlaces. En su lugar, las tramas se transmitirán en un enlace específico, como resultado de un algoritmo de hash. El algoritmo puede utilizar la dirección IP de origen, dirección IP de destino, o una combinación de las direcciones de origen y destino IP, direcciones MAC de origen y de destino, o los números de puerto TCP/UDP. El algoritmo de control calcula un patrón binario que selecciona un número de enlace en el paquete para cada trama.

Si es ordenado sólo una dirección o un número de puerto, el switch envía cada trama mediante el uso de uno o más bits de orden bajo del valor hash como un índice en los enlaces agrupados. Si son ordenadas dos direcciones o números de puerto, un interruptor realiza una operación OR exclusiva (XOR) en uno o más bits de orden inferior de las direcciones o números de puerto TCP / UDP como un índice en los enlaces agrupados.

Por ejemplo, un EtherChannel que consta de dos enlaces agrupados requiere un índice de un bit para la representación (0 y 1). O se utiliza la dirección de orden de bits más baja o el XOR de la última parte de la dirección en el cuadro. Un paquete de cuatro brazos utiliza un hash de los dos últimos bits. Del mismo modo, un paquete de ocho enlaces utiliza un hash de los últimos tres bits. El resultado de la operación de hash selecciona enlaces externos del EtherChannel. El Cuadro 14 muestra los resultados de un XOR en un paquete de dos enlaces, el uso de direcciones de origen y destino.

La operación XOR se lleva a cabo de forma independiente en cada posición de bit en el valor de la dirección. Si los dos valores de dirección tienen el mismo valor de bit, el resultado XOR es 0. Si los dos bits de dirección son diferentes, el resultado XOR es 1. De esta manera, las tramas pueden ser distribuidas estadísticamente entre los enlaces con la suposición de que MAC o direcciones IP se distribuyen estadísticamente a lo largo de la red. En un enlace de cuatro EtherChannel, el XOR se lleva a cabo en los dos bits más bajos de los valores de dirección que resulta en un valor XOR 2 bits (cada bit se calcula por separado) o un número de enlace de 0 a 3.

Dirección en Binario	XOR EtherChannel con 2 enlaces y el número del enlace
Addr1: ...xxxxxxx0 Addr2: ...xxxxxxx0	xxxxxxx0: Se usa el enlace 0
Addr1: ...xxxxxxx0 Addr2: ...xxxxxxx1	xxxxxxx1: Se usa el enlace 1
Addr1: ...xxxxxxx1 Addr2: ...xxxxxxx0	xxxxxxx1: Se usa el enlace 1
Addr1: ...xxxxxxx1 Addr2: ...xxxxxxx1	xxxxxxx0: Se usa el enlace 0

Cuadro 14. Distribución de tráfico en un EtherChannel de 2 enlaces

Como ejemplo, considere un paquete que se envía desde la dirección IP **192.168.1.1** a **172.31.67.46**. Debido que EtherChannel puede construirse de dos a ocho eslabones individuales, sólo en el extremo derecho (menos significativo) se necesitan tres bits como un índice de referencia. Estos bits son 001 (1) y 110 (6), respectivamente. Para un EtherChannel de dos enlaces, un XOR de un bit se lleva a cabo en la dirección del bit más a la derecha: **1 XOR 0 = 1**, causando que el enlace 1 sea utilizado. Un EtherChannel de cuatro enlaces produce un XOR de dos bits: **01 XOR 10 = 11**, causando que el enlace 3 sea utilizado. Por último, un enlace de ocho EtherChannel requiere un XOR de tres bits: **001 XOR 110 = 111**, donde se selecciona el enlace 7 en el paquete para ser utilizado.

Una conversación entre dos dispositivos se envía siempre a través del mismo enlace EtherChannel porque las dos direcciones de extremos permanecen igual. Sin embargo, cuando un dispositivo se comunica con otros dispositivos, lo más probable es que las direcciones de destino se distribuyen por igual a 0 y 1 en el último bit. Tenga en cuenta que una conversación entre dos dispositivos finales para crear un desequilibrio de la carga se puede utilizar uno de los enlaces en un paquete porque todo el tráfico entre un par de estaciones usará el mismo enlace.

5.1.4 Configuración del método EtherChannel para el balanceo de carga

La operación del hash se puede realizar en cualquiera de las direcciones MAC e IP, y puede estar basada únicamente en las direcciones de origen o de destino, o ambos.

Valor del método	Entrada Hash	Operación Hash	Modelo de Switch
src-IP	Dirección IP origen	Bits	6500/4500
dst- IP	Dirección IP destino	Bits	6500/4500
src-dst- IP	Dirección IP origen y destino	XOR	6500/4500/3550
src- MAC	Dirección MAC origen	Bits	6500/4500/3550
dst- MAC	Dirección MAC destino	Bits	6500/4500/3550

src-dst- MAC	Dirección MAC origen y destino	XOR	6500/4500
src-port	Número de puerto origen	Bits	6500/4500
dst-port	Número de puerto destino	Bits	6500/4500
src-dst-port	Puerto origen y destino	XOR	6500/4500

Cuadro 15. Tipos de métodos de balanceo de carga EtherChannel

Debe de tener en cuenta que no todos estos métodos ofrecen un balanceo de tráfico equitativo en los enlaces que conformen el grupo, de manera que se deberá elegir el que mejor funcionalidad preste a la red.

5.1.5 Protocolos de negociación de EtherChannel

Un EtherChannel puede ser negociado entre dos switch de dos formas que son Dinámica y manualmente. A su vez cabe destacar que la negociación se hace de manera dinámica existen dos protocolos, las cuales son: PAgP que es un protocolo propietario de cisco y el LACP que es un protocolo propio del estándar IEEE, cada uno de ellos se abordan en las siguientes secciones.

5.1.5.1 Protocolo de agregación de puerto (PAgP)

Cuando se configura PAgP el switch negocia con el otro extremo que puertos deben ponerse como activos, en versiones anteriores de este protocolo los puertos que no eran compatibles se dejaban desactivados, no obstante en la versión reciente si un puerto o es compatible se deja fuera del EtherChannel quedando activo y trabajando de manera independiente. PAgP se encarga de agrupar puertos de características similares de forma automática. PAgP es capaz de agrupar puertos de las mismas velocidades, modo dúplex, troncales o de asignación a una misma VLANs.

PAgP puede opera de 2 maneras que son:

- **Auto:** Establece puertos en una negociación pasiva el puerto responderá a paquetes PAgP cuando los reciba, pero nunca iniciara la negociación.
- **Desirable:** Establece el puerto en modo de negociación activa, este puerto negociara el estado cuando reciba paquetes PAgP y también podrá iniciar una negociación contra otros puertos.

5.1.5.2 Protocolo de agregación de enlaces de control (LACP)

EL protocolo LACP es definido en el estándar IEEE 802.3ad y puede ser implementado en switches Cisco. LACP y PAgP funcionan de manera similar ya que LACP puede agrupar puertos por su velocidad, modo duplex, troncales, VLANs.

LACP también tiene 2 modos de configuración:

- **Activo:** Un puerto en este estado es capaz de iniciar negociaciones con otros puertos para establecer el grupo.

- **Pasivo:** Un puerto en este estado es un puerto que no iniciara ningún tipo de negociación, pero si responderá a las negociaciones generadas por otros puertos.

Tener en cuenta que al igual que en PAgP, dos puertos pasivos nunca podrán formar un grupo.

5.1.6 Modo de configuración de EtherChannel

A como se mencionó antes que se puede configurar EtherChannel tanto dinámicamente como de forma manual, esta última forma de configuración se describe a continuación:

- **Configuración modo ON (manual):** El modo ON es un modo de configuración en el cual se establece toda la configuración del puerto de forma manual, no existe ningún tipo de negociación entre los puertos para establecer un grupo. En este tipo de configuración es necesariamente que ambos lados estén en modo ON.
 - **Limitaciones:** Una limitación de EtherChannel es que todos los puertos físicos en el grupo de agregación deben residir en el mismo conmutador. El protocolo SMLT avaya elimina esta limitación al permitir que los puertos físicos sean divididos entre 2 switches en una configuración de triangulo o 4 o más switches en una configuración de malla.
 - **Ventajas:** Usar EtherChannel tiene muchas ventajas, una de las mejores es el aprovechamiento del ancho de banda. Utilizando un máximo de 8 puertos, con un ancho de banda total de 800Mbps, de 8Gbps o de 80Gbps, dependiendo de la velocidad del cable utilizado. Esto asume que hay una mezcla de tráfico, pues esas velocidades no se aplican a un solo uso.

Puede ser utilizado con Ethernet, que funciona en fibra sin blindaje del cableado del par trenzado (UTP), unimodo y con varios modos de funcionamiento. EtherChannel se aprovecha del cableado existente haciéndolo escalable. Puede ser utilizado en todos los niveles de la red para crear acoplamientos más altos en la anchura de banda mientras que las necesidades del tráfico en la red aumentan.

- **Configuración EtherChannel:** Para cada EtherChannel en un interruptor, se debe elegir el protocolo de negociación EtherChannel y asignar los puertos de conmutación individuales al EtherChannel. Tanto la configuración de PAgP y LACP se describen en las siguientes secciones. También puede configurar un EtherChannel utilizar el modo ON, que incondicionalmente haces los enlaces. En este caso no serán enviados o recibidos paquetes PAgP ni LACP. En los puertos se configuran para ser miembros de un EtherChannel, el parámetro genera automáticamente una interfaz de canal de puerto lógico. Esta interfaz representa el canal en su conjunto.

Tarea	Sintaxis de Comando
Seleccionar un método de balanceo de carga para el Switch	port-channel load-balance método
Usar el protocolo PAgP	channel-protocol pagp channel-group número mode {on auto desirable}
Usar el protocolo LACP	channel-protocol lacp channel-group número mode {on pasive active}

Cuadro 16. Configuración de EtherChannel PAgP, LACP, ON

6. PPP

6.1 Definición de PPP

El protocolo PPP, definido en el RFC 1661 y en otras posteriores, efectúa la detección de errores, reconoce múltiples protocolos, permite la negociación de direcciones IP en el momento de la conexión y da soporte a la verificación de la autenticidad. Resumiendo, tiene todas las prestaciones, y algunas más, de las que adolece el protocolo SLIP.

El protocolo PPP (Point to Point Protocol) o Protocolo Punto a Punto, proporciona básicamente tres funciones: Es un método de conformado de tramas que determina sin ambigüedades sus límites y que gestiona la detección de errores. Es un protocolo de control de enlace, LCP (Link Control Protocol), para activar Líneas, testearlas, negociar opciones o parámetros y desactivarlas ordenadamente cuando ya no son necesarias. Y por último, es un mecanismo encargado de negociar opciones de la capa de red con independencia del protocolo de red empleado; el método para implementar este mecanismo consiste en tener disponible un NCP (Network Control Protocol), Protocolo de Control de Red, para cada capa de red reconocida.

PPP se puede utilizar en varias interfaces físicas diferentes, incluyendo las siguientes:

- Serial asíncrono
- Interfaces serial de alta velocidad (HSSI)
- ISDN serie síncrona

6.2 PPP en el Protocolo Stack

PPP se puede utilizar a través de conexiones tanto asíncrono y síncronos en la capa física del modelo de referencia OSI.

6.3 Protocolo de Control de Enlace (LCP).

Se utiliza en la capa de enlace de datos para establecer, configurar y probar la conexión.

6.4 Protocolos de Control de Red (NCP)

Permitir el uso simultáneo de múltiples protocolos de capa de red y se requieren para cada protocolo que utiliza PPP.

6.5 Estructura de PPP

El protocolo PPP se basa en, y similar a, el protocolo HDLC con algunas diferencias importantes, en que utiliza PPP LCP, NCP y soportes: autenticación, cifrado y compresión.

Estructura de PPP	
Capa de Red	IP, IPX, AppleTalk. IPCP, IPXCP, AppleTalkCP.
PPP en la capa de enlace de datos	NCP (Network Control Protocol): Hacia arriba el servicio a la capa 3 (capa de red) de múltiples capas a través de soporte para el protocolo NCP. PPP puede cambiarla capa 3 direccionamiento.

	LCP (Link Control Protocol): Maneja la configuración de enlace, la terminación, la autenticación, encriptación, compresión y detección de errores
PPP en la capa física	Serie síncrona o asíncrona Interfaz serial de alta velocidad Líneas telefónicas Líneas Alquiladas

Cuadro 17. Estructura de PPP

6.6 Protocolo de control de enlace

Maneja la configuración de enlace y terminación, así como incluyendo la mayoría de las opciones de configuración: autenticación, compresión de detección de errores, y multilink (balanceo de carga).

- **Autenticación:** PPP puede ser configurado para utilizar PAP (Protocolo de autenticación de contraseña) y CHAP (Challenge Handshake Authentication Protocol). Ambos protocolos autenticar conexiones en serie con un nombre de usuario y contraseña, pero CHAP es más seguro, ya que utiliza un protocolo de enlace de tres vías durante la negociación LCP enlace inicial, a continuación, utiliza tres desafíos handshake manera a veces al azar después, asegurando que el enlace está todavía confiaba.
- **Compresión:** El plan de estudios CCNA cubre dos opciones de configuración de compresión PPP: Agrupador y predictor. PPP puede lograr un mejor rendimiento mediante la compresión de los datos que viajan a través de los enlaces.
- **Detección de errores:** PPP puede ser configurado para comprobar y comunicar la calidad del enlace en serie, también puede detectar las interfaces en un estado de bucle de retorno.
- **Multilink:** PPP se puede configurar para equilibrar la carga entre múltiples interfaces en serie para mayor ancho de banda.

6.7 Formato de la trama

PPP se basa en el control de enlace de datos de alto nivel (HDLC)

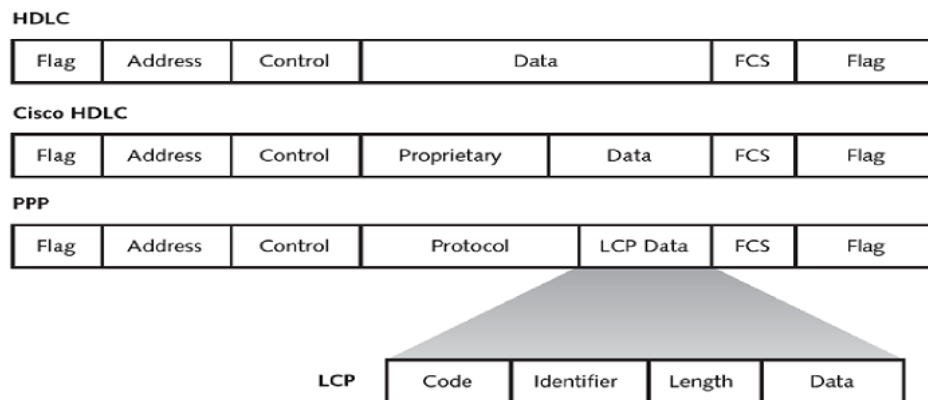


Figura 20. Estructura del paquete HDLC y PPP

LCP campo del paquete PPP puede contener muchos diferentes piezas de información, incluyendo las siguientes:

- Mapa carácter asíncrono
- Máximo tamaño de la unidad recibe
- Compresión
- Autenticación
- Número mágico
- Enlace Monitoreo de la Calidad (LQM)
- Multilink

6.8 LCP enlace proceso de configuración

Modifica y mejora las características predeterminadas de una conexión PPP. Incluye las siguientes acciones:

- Establecimiento del Enlace.
- Autenticación (opcional).
- Calidad del Enlace determinación (opcional).
- Red de capa de protocolo de negociación de configuración.
- Terminación del Enlace.

6.8.1 Clases de tramas LCP existe

- **Establecimiento de enlace:** Las tramas se utilizan para establecer y configurar un enlace.
- **Terminación de enlace:** Las tramas se utilizan para terminar una relación.
- **Tramas de enlace de mantenimiento:** Se utiliza para administrar y depurar un enlace.

Opciones de configuración LCP			
Opción	función	protocolo	comando
Autenticación	Requiere una contraseña>>	PAP	PPP authentication pap
	Realiza un apretón de manos >>	CHAP	PPP authentication CHAP
Compresión	Comprime los datos en la fuente>>	Stacker	PPP compression Stacker
	Reproduce los datos en el destino>>	Predictor	PPP compression predictor
Detección de errores	Monitoriza los datos colocados en el enlace	quality, número mágico	PPP quality <número1-100>
Multilink	Realiza el equilibrio de carga a través de múltiples enlaces	MP	PPP multilink

Cuadro 18. Opciones de configuración LCP

6.9 Establecimiento de comunicaciones PPP

Involucra las siguientes acciones:

- Enlace establecimiento.
- Autentificación opcional.
- Red de capa de protocolo de negociación de configuración.

La fase de establecimiento del enlace implica la configuración y prueba del enlace de datos.

El proceso de autenticación puede utilizar dos tipos de autenticación con conexiones PPP: PAP y CHAP. PPP es un tipo de encapsulación para las comunicaciones de la interfaz serie.

Para configurar una conexión PPP, debe acceder al modo de configuración de interfaz para la interfaz específica que desea configurar.

- Después de LCP ha terminado de negociar los parámetros de configuración
- Los protocolos de capa de red se pueden configurar individualmente por el NCP apropiado.

6.10 Configurar la autenticación PPP

- Uso de la autenticación con conexiones PPP es opcional.
- Usted específicamente debe configurar la autenticación PPP en cada host PPP para que el host pueda utilizar.
- Usted puede optar por habilitar CHAP, PAP, o ambos en su conexión PPP, en cualquier orden.

Una vez establecido el tipo de autenticación

- No obstante, debe configurar un nombre de usuario y la contraseña para la autenticación.
- Debe salir del modo de configuración de interfaz y entrar en el modo de configuración global.
- Tipo de nombre de usuario seguido del nombre de host del router remoto.
- A continuación, escriba la contraseña seguida de la contraseña de la conexión.
- Confirmación de Comunicaciones PPP
- Con el comando show interface.

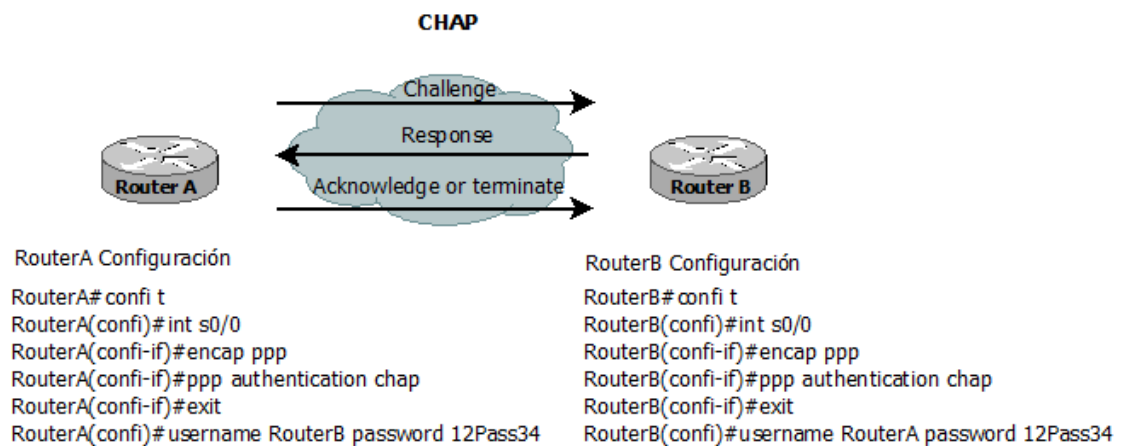


Figura 21. Configuración de PPP Chap

7. HSRP

7.1 Definición de HSRP

El Hot Standby Router Protocol es un protocolo propiedad de CISCO que permite el despliegue de Routers redundantes tolerantes a fallos en una red. Este protocolo evita la existencia de puntos de fallo únicos en la red mediante técnicas de redundancia y comprobación del estado de los Routers.

El funcionamiento del protocolo HSRP es el siguiente: Se crea un grupo (también conocido por el término inglés Clúster) de Routers en el que uno de ellos actúa como maestro, enrutando el tráfico, y los demás actúan como respaldo a la espera de que se produzca un fallo en el maestro. HSRP es un protocolo que actúa en la capa 3 del modelo OSI administrando las direcciones virtuales que identifican al Router que actúa como maestro en un momento dado.

Supongamos que disponemos de una red que cuenta con dos Routers redundantes, Router A y Router B. Dichos Routers pueden estar en dos posibles estados diferentes: maestro (Router A) y respaldo (Router B). Ambos Routers intercambian mensajes, concretamente del tipo HSRP hello, que le permiten a cada uno conocer el estado del otro. Estos mensajes utilizan la dirección multicast 224.0.0.2 y el puerto UDP 1985.

Si el Router maestro no envía mensajes de tipo hello al Router de respaldo dentro de un determinado periodo de tiempo, el Router respaldo asume que el maestro está fuera de servicio (ya sea por razones administrativas o imprevistas, tales como un fallo en dicho Router) y se convierte en el Router maestro. La conversión a Router activo consiste en que uno de los Router que actuaba como respaldo obtiene la dirección virtual que identifica al grupo de Routers.

Para determinar cuál es el Router maestro se establece una prioridad en cada Router. La prioridad por defecto es 100. El Router de mayor prioridad es el que se establecerá como activo. Hay que tener presente que HSRP no se limita a 2 Routers, sino que soporta grupos de Routers que trabajen en conjunto de modo que se dispondría de múltiples Routers actuando como respaldo en situación de espera.

El Router en espera toma el lugar del Router maestro, una vez que el temporizador holdtime expira (un equivalente a tres paquetes hello que no vienen desde el Router activo, timer hello por defecto definido a 3 y holdtime por defecto definido a 10).

Los tiempos de convergencia dependerán de la configuración de los temporizadores para el grupo y del tiempo de convergencia del protocolo de enrutamiento empleado.

Por otra parte, si el estado del Router maestro pasa a down, el Router decrementa su prioridad. Así, el Router respaldo lee ese decremento en forma de un valor presente en el campo de prioridad del paquete hello, y se convertirá en el Router maestro si ese valor decrementado es inferior a su propia prioridad. Este proceso decremental puede ser configurado de antemano estableciendo un valor por defecto del decremento (normalmente, de 10 en 10).

7.2 HSRP-Ejemplo práctico

- Entre los Routers del grupo HSRP se intercambian mensajes hello para conocer el estado en el que se encuentran.
- Estos mensajes utilizan la dirección multicast 224.0.0.2 y el puerto UDP 1985.

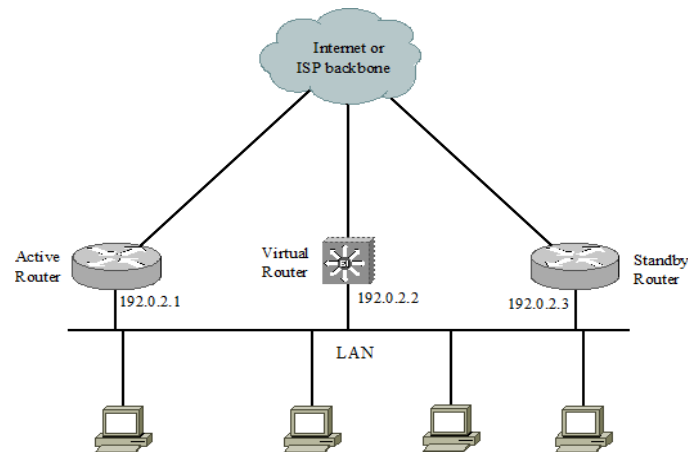


Figura 22. Ejemplo de HSRP

- Si el Router maestro no envía mensajes tipo hello a los Routers respaldo dentro de un periodo de tiempo, otro Router del grupo se coloca como Router maestro.
- Esto consiste en que el nuevo Router obtiene la dirección virtual que identifica al grupo.

7.3 HSRP - Formato de paquetes

Cabecera MAC				Paquete HSRP				Cabecera IP				Paquete UDP			
00	01	02	03	04	05	06	07	08	09	10	11	12	13	14	15
Version								codigo operacion							
Tiempo de espera								Prioridad							
Datos de autenticación															
Dirección IP virtual															
16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31
Estado								Tiempo del Hello							
Grupo								Reservados							
Datos de autenticación															
Dirección IP virtual															

Figura 23. Formato de paquetes HSRP

7.4 HSRP-Descripción campos

- Versión (8 bits): Número de versión HSRP.

- Código operación (8 bits): Operación que se llevará a cabo.

Opcod	Función	Descripción
0	Hello	El Router está funcionando y es capaz de convertirse en el Router activo o de reserva.
1	Coup	El Router desea convertirse en el Router activo.
2	Resign	El Router ya no desea ser el Router activo.

- Estado (8 bits): Este campo describe el estado actual del Router al enviar el mensaje.

Estado	Función	Descripción
0	Inicio	Esta es la situación de partida y se indica que HSRP no se está ejecutando este estado es introducido a través de un cambio de configuración o cuando una primera interfaz aparece.
1	Aprender	El Router no ha determinado la dirección IP virtual, y todavía no ha visto autenticado el mensaje hello desde el Router activo. En este estado, el Router sigue a la espera de escuchar desde el Router activo.
2	Escuchar	El Router conoce la dirección IP virtual, pero no es ni el Router activo ni el Router de espera. Está a la espera de comunicaciones de otros Routers.
4		El Router envía periódicamente mensajes Hello y participa activamente en la elección de los activos y/o Routers de espera. un Router no puede entrar en estado speak a menos que tenga la dirección IP virtual.
8	Espera	El Router es un Candidato para convertirse en el Router activo y envía mensajes periódicos Hello. Excluyendo condiciones transitorias, debe haber al menos un Router en el grupo en estado de espera.
16	Activo	El Router actualmente se encuentra reenviando paquetes a la dirección del grupo Mac Virtual. el Router envía periódicamente mensajes de saludo Excluyendo condiciones transitorias, debe haber al menos un Router en estado activo en el grupo.

Cuadro 19. Estado de un Router al enviar un mensaje

- **Hellotime (8 bits) Defecto = 3 segundos.** Este campo sólo tiene sentido en mensajes Hello. Contiene el período aproximado entre los mensajes de saludo que el Router envía. El tiempo se da en segundos. Si el Hellotime no está configurado en un Router, entonces puede ser adquirido en el mensaje Hello del Router activo. Un Router que envía un mensaje Hello debe insertar el Hellotime en el paquete.
- **Holdtime (8 bits) Defecto = 10 segundos.** Contiene la cantidad de tiempo que el actual Hello debe ser considerado válido. El tiempo se da en segundos. Si un Router envía un mensaje Hello, a continuación, los

receptores deberían considerar la validez en el tiempo Holdtime. El Holdtime debe ser de al menos tres veces el valor del Hellotime.

- **Prioridad (8 bits)** Este campo se utiliza para elegir a los Routers activos y reservas. Al comparar las prioridades de los dos Routers, gana el Router con la prioridad más alta. En el caso de empate gana el de la dirección IP más grande.
- **Grupo (8 bits)** Este campo identifica el grupo de espera. Para Token Ring, los valores entre 0 y 2 inclusive son válidos. Para otros valores medios de comunicación entre 0 y 255 inclusive son válidos.
- **Reservados (8 bits)**
- **Datos de autenticación (64 bits)** Este campo contiene una contraseña de 8 caracteres reutilizados. Si no hay datos de autenticación se configura, el valor predeterminado recomendado es de 0x63 0x69 0x73 0x63 0x6F 0x00 0x00 0x00.
- **Dirección IP virtual (32 bits)** La dirección IP virtual utilizada por este grupo. Si la dirección IP virtual no está configurada en un Router, entonces puede ser adquirida en el mensaje Hello desde el Router activo. La dirección sólo se debe aprender si no se ha configurado la dirección y el mensaje Hello esta autenticado.

7.5 Configuración CISCO

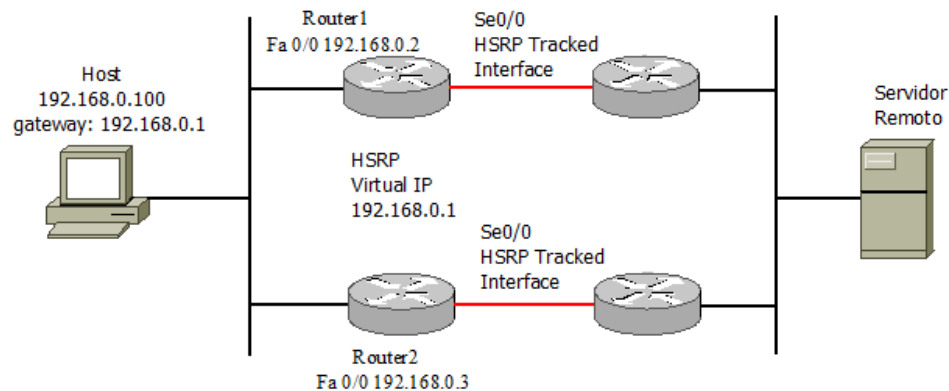


Figura 24. Configuración Cisco de HSRP

Router 1: Añadir IP a la interfaz.

```
hsrp-router1#conf t
hsrp-router1(config)# int fa0/0
hsrp-router1(config-if)# ip address 192.168.0.2 255.255.255.0
Set the Virtual IP Address
```

Poner la IP virtual donde "1" es el grupo HSRP y "192.168.0.1" es la IP virtual para el grupo HSRP

```
hsrp-router1(config-if)# standby 1 ip 192.168.0.1
Enable Preempt
```

Esto sirve para que el Router pase de pasivo a activo cuando se da cuenta de que el Router activo ha caído o él mismo tiene la prioridad más alta.

```
hsrp-router1(config-if)# standby 1 preempt
Set Router Priority
```

La prioridad por defecto es 100, ponemos 110 para que el Router sea el activo en el grupo.

```
hsrp-router1(config-if)# standby 1 priority 110
```

Set Authentication String: Es una cadena de 8 caracteres opcionales que pueden ser usados en el paquetes "hello" multicast para autenticar el grupo HSRP.

```
hsrp-router1(config-if)# standby 1 authentication Local LAN
```

SetTimers: Sets the time period between the "hello" packets and the hold time before assuming an active router is down. Default is 3seconds and 10 seconds respectively.

Poner el período de tiempo entre los paquetes "hello" y el "holdtime" antes de asumir que un Router activo ha caído. Por defecto es 3 10 respectivamente.

```
hsrp-router1(config-if)# standby 1 timers 5 15
```

Track Interface: Poniendo este comando en la interfaz y después el número que quieres restarle a la prioridad en caso de caída de este interface, es decir, si tengo un serial 0/0 y se me cae, tengo el Track en el hsrp con el serial 0/0 y un decremento de 11 este se quedaría como standby ya que $110-11=99$ que es menos que 100 de prioridad del Router que antes estaba como standby. Una vez Si no ponen ningún número después del interface WAN en el comando Track por defecto le resta 10 a la prioridad por lo que no serviría.

```
hsrp-router1(config-if)# standby 1 track se0/0
```

Mismo método para el Router 2:

```
hsrp-router2#conf terminal
hsrp-router2(config)# int fa0/0
hsrp-router2(config-if)# ip address 192.168.0.3 255.255.255.0
hsrp-router2(config-if)# standby 1 ip 192.168.0.1
hsrp-router2(config-if)# standby 1 preempt
hsrp-router2(config-if)# standby 1 priority 100
hsrp-router2(config-if)# standby 1 authentication LocalLAN
hsrp-router2(config-if)# standby 1 timers 5 15
hsrp-router2(config-if)# standby 1 track se0/0
```

8. VRRP

Existen diferentes métodos con los que un sistema final puede determinar cuál es el router de primer salto (**Gateway**) hacia un destino IP en particular, como por ejemplo:

- Protocolos de encaminamiento dinámicos como RIP u OSPF.
- Utilizar un cliente de descubrimiento de router mediante ICMP.
- Proxy ARP.
- Emplear una ruta por defecto configurada estáticamente.

No siempre resulta factible la utilización de un protocolo de enrutamiento dinámico en los equipos, ya que supone un mayor esfuerzo de configuración, requiere una mayor capacidad de procesamiento, pueden aparecer problemas de seguridad, o incluso puede aparecer que en una plataforma determinada no exista la implementación de algún protocolo específico.

Los protocolos de descubrimiento de router requieren la participación activa de todos los equipos de la red, lo que da lugar a la configuración de valores elevados en los temporizadores para reducir la sobrecarga introducida cuando el número de los equipos es muy grande. Esto provoca que la detección de pérdida de un router pueda tener lugar con un retardo significativamente grande.

El uso de una ruta por defecto configurada estáticamente se encuentra ampliamente extendido, minimiza la sobrecarga de configuración en los dispositivos finales y se encuentra soportado prácticamente por cualquier implementación IP. Sin embargo, este modo de funcionamiento crea un punto único de fallo. La pérdida de una ruta por defecto en una LAN resulta catastrófica dado que los equipos quedan aislados en su red, siendo incapaces de encontrar un camino alternativo que incluso puede encontrarse disponible.

Virtual Router Redundancy Protocol (VRRP) es un estándar alternativo basado en el protocolo HSRP, definido en el estándar IETF. VRRP es similar a HSRP, antes de entrar a detalles de configuración y funcionamiento se debe tener conocimientos acerca como opera el protocolo HSRP para la previa comprensión del protocolo VRRP.

El protocolo VRRP fue diseñado para aumentar la disponibilidad de una puerta de enlace por defecto, dando servicios a maquinas en la misma subred. El aumento de fiabilidad se consigue mediante el anuncio de un router virtual como una puerta de enlace por defecto en lugar de un router físico. Dos o más router físicos se configuran representando al router virtual, con solo uno de ellos realizando realmente el enrutamiento. Si el router actual que está haciendo el encaminamiento de tráfico falla, el otro router negocia para sustituirlo. Se denomina router maestro al dispositivo físico que en ese momento este realizando el enrutamiento de tráfico y router de respaldo al router o routers que en ese momento están en espera a que el maestro falle.

Esto permite que cualquiera de las direcciones IP asociadas al router virtual pueda ser utilizada como ruta de primer salto (Gateway) para los dispositivos de esa LAN.

8.1 Descripción del protocolo VRRP

El protocolo VRRP proporciona la funcionalidad de router virtual descrita anteriormente.

8.2 Definiciones

A continuación se exponen un conjunto de definiciones:

8.2.1 Router VRRP

Router que ejecuta el protocolo VRRP. Un router VRRP puede participar en uno o varios routers virtuales.

8.2.2 Router Virtual

Elemento o figura de abstracción creado por un grupo de routers que ejecutan VRRP y que actúa como router por defecto (Gateway) de los equipos de una LAN. Los routers virtuales consisten en un identificador de router virtual y una dirección IP. Un router VRRP puede servir de backup de varios routers virtuales simultáneamente.

8.2.3 Propietario de dirección IP

Router que tiene la **dirección IP virtual** (la asociada al router virtual) como dirección real en algunas de sus interfaces.

8.2.4 Dirección IP principal (multidifusión)

Dirección IP seleccionada dentro del conjunto de direcciones de interfaz reales. Los mensajes de anuncio del protocolo VRRP siempre se envían utilizando la dirección principal como dirección IP origen del paquete.

8.2.5 Master del router virtual

Router VRRP que se encarga de procesar los paquetes encaminados a través de la dirección IP asociada al router virtual y responder a las peticiones ARP de dicha IP virtual.

8.2.6 Virtual router ID

También llamado VRID, esta es una identificación numérica en un enrutador virtual en particular. Este número de identificación debe ser exclusivo en un segmento de red determinado.

8.2.7 Dirección IP virtual

Una dirección IP asociada a un VRID que otras máquinas puedan utilizar para obtener un servicio de red.

8.2.8 Dirección MAC virtual

Para los medios de comunicación que utilizan direccionamiento MAC (Ethernet), las instancias VRRP utilizan una dirección MAC predefinida para todas las acciones VRRP en lugar de las direcciones MAC del adaptador real. Esto aísla el funcionamiento del enrutador virtual desde el router real, proporcionando la función de encaminamiento. El VMAC (Virtual MAC) se deriva de la VRID (ID router virtual).

8.2.9 Prioridad

Diferentes instancias VRRP se les asigna un valor de prioridad, como una forma de determinar que router asumirá el papel de master si el maestro actual falla. La prioridad es un número de 1 a 254 (0 y 255 están reservados). Cuanto mayor sea el número mayor es la prioridad.

8.2.10 Propietario

Si la dirección IP virtual es la misma que que las direcciones ip configuradas en la interfaz de un router, que router es el propietario de la dirección ip virtual. La prioridad de la instancia VRRP cuando es el propietario es de 255, es decir el valor más alto.

8.3 Funcionamiento del protocolo VRRP

El funcionamiento del protocolo VRRP se basa en la simulación de un router virtual entre varios routers VRRP. Al router virtual se le asocia una IP virtual y una dirección MAC virtual. Estas direcciones virtuales se mantienen inmutables con independencia del router real que se encargue del encaminamiento de los paquetes asociados al router virtual.

El protocolo VRRP utiliza los mensajes de anuncio (Advertisements) para indicar que el router que hace de master se encuentra activo. Estos mensajes se envían a la dirección ip multicast 224.0.0.18 asignada por la Internet Assigned Numbers Authority (IANA). El número de protocolo ip establecido por la IANA para el VRRP es el 112 (en decimal). Los Advertisements contienen información del router virtual, su prioridad, etc.

Si durante un periodo de tiempo determinado (Master_Down_Interval) los routers VRRP que están haciendo el papel de backup dejan de recibir mensajes del master, entonces el router de backup con mayor prioridad pasa a convertirse en el nuevo master del router virtual.

Por defecto, un equipo de backup que posea mayor prioridad que el master actual puede expropiar en sus funciones al mismo y convertirse en el nuevo master. Este comportamiento asegura que siempre se encuentre como master el router con mayor prioridad. Sin embargo, si por alguna razón es necesario, se puede deshabilitar administrativamente la expropiación del router virtual.

VRRP combina un grupo de routers (incluyendo al master y a los múltiples backups) en una LAN con un router virtual llamándolo grupo VRRP. Un grupo VRRP tiene las siguientes características:

- Un router virtual tienen una IP virtual. Un host en la LAN solo necesita saber la dirección IP de router virtual y usar esa dirección como el siguiente salto (gateway) de la ruta por defecto.
- Todos los host de la LAN se comunican con las redes externas (internet) a través del router virtual.
- Los routers en el grupo VRRP de acuerdo a sus prioridades eligen un master que actúa como el gateway. Los otros routers actúan como copias de seguridad (Backups). Cuando el maestro falla, para asegurarse que los host en el segmento de red se puedan comunicar sin interrupción con las redes externas, las copias de seguridad en el grupo VRRP eligen una nueva puerta de enlace y asume la responsabilidad del maestro fallado.

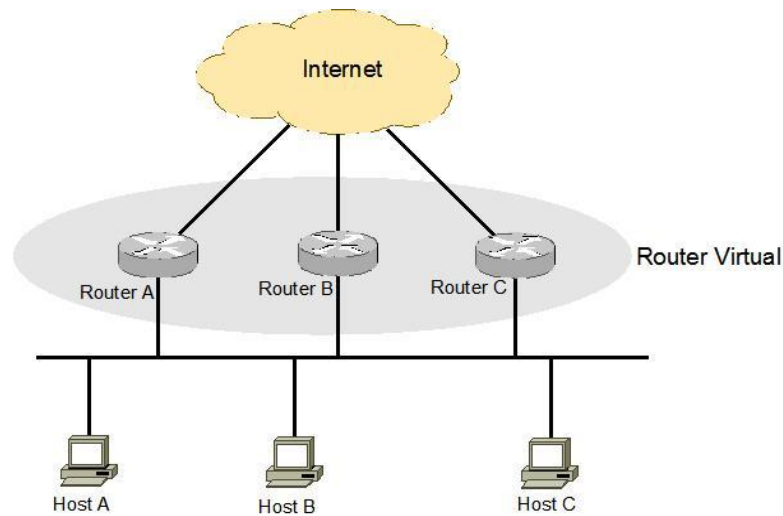


Figura 25. Diagrama de red implementando VRRP con un router virtual

Como se muestra en la Figura 25, el router A, router B y router C forman un router virtual, que tiene su propia dirección IP. Las estaciones de trabajo perteneciente a esa LAN utilizan el router virtual como la puerta de enlace predeterminada. El router con la prioridad más alta entre los tres routers es elegido como el maestro de actuar como puerta de entrada., mientras los B y C son copias de seguridad.

NOTA: La dirección IP del router virtual puede ser una dirección IP no utilizada en el segmento en el que reside el grupo VRRP o la dirección IP de una interfaz de un router que pertenezca al grupo VRRP. En este último caso, el router se llama el propietario de la dirección IP. Solo un propietario de la dirección IP se configura para un grupo VRRP.

8.4 Prioridad VRRP

VRRP determina el papel (Maestro o Backup) de cada router en un grupo VRRP por prioridad. Un router con mayor prioridad tiene más probabilidades de convertirse en el router maestro.

La prioridad VRRP está entre el intervalo de 0 a 255. Cuanto mayor sea el número, mayor es la prioridad. Prioridades 1-254 son configurables. La prioridad 0 se reserva para usos especiales y prioritarias 255 para el propietario de la dirección IP, su prioridad de ejecución es siempre 255. Es decir, el propietario de la dirección IP en un grupo VRRP actúa como maestro mientras funciona correctamente.

8.5 Modo de trabajo

Un router en un grupo VRRP funciona en cualquiera de los siguientes modos:

- **Modo no preferente (non-preemptive):** Cuando un router en el grupo VRRP se convierte en el maestro, se mantiene como el maestro siempre y cuando funcione normalmente, incluso si a una copia de seguridad se le asigna una prioridad más alta más tarde.

- **Modo preferente (preemptive):** Cuando una copia de seguridad encuentra su prioridad más alta que la del maestro, la copia de seguridad envía anuncios VRRP para iniciar una nueva elección principal en el grupo VRRP y se convierte en el maestro. En consecuencia, el maestro original se convierte en una copia de seguridad.

8.6 Modo de autenticación

Para evitar los ataques de los usuarios no autorizados, VRRP añade claves de autenticación en paquetes para su autenticación. VRRP proporciona los siguientes modos de autenticación:

- **Autenticación de texto simple:** Un router envía un paquete añadiendo una clave de autenticación, y el router que recibe el paquete compara su clave de autenticación local con la del paquete recibido. Si las dos claves de autenticación son las mismas, el paquete VRRP recibido se considera legítimo. De lo contrario el paquete recibido se considera no válido.
- **Autenticación MD5:** Un router calcula el producto de digestión de un paquete a enviar mediante el uso de la clave de autenticación y el algoritmo MD5 y guarda el resultado en el encabezado de autenticación. El router que recibe el paquete realiza la misma operación mediante el uso de la clave de autenticación y el algoritmo MD5, y compara el resultado con el contenido de la cabecera de autenticación. Si el resultado es el mismo, el router que recibe los paquetes considera los paquetes como auténticos y válidos.

8.7 Timers VRRP

VRRP timers incluyen advertencias VRRP en intervalos de tiempo y tiempo de retardo VRRP preferente.

- **Intervalos de tiempo de advertencias VRRP:** El router maestro en un grupo VRRP periódicamente envía advertencias VRRP para informar a los otros routers en el grupo VRRP que él está operando perfectamente. Tu puedes ajustar el intervalo de envío de advertencias VRRP configurando el intervalo de tiempo de advertencias VRRP, esta configuración se realizara depende el modelo del equipo.
- **Tiempo de retardo VRRP preferible:** Para evitar los cambios de estado frecuentes entre los miembros de un grupo VRRP y proporcionar las copias de seguridad de tiempo suficiente para recopilar información (tal como información de encaminamiento), cada copia de seguridad espera un período de tiempo (el tiempo de retardo de preferencia) después de que se recibe un anuncio con la prioridad más baja que la prioridad local, a continuación, envía anuncios VRRP para iniciar una nueva elección principal en el grupo VRRP y el router con prioridad más baja se convierte en el maestro.

8.8 Formato de paquete VRRP

El router maestro envía de manera multicast paquetes VRRP para avisar de su existencia. Los paquetes VRRP son también usados para verificar los parámetros del router virtual y elección del maestro.

Los paquetes VRRP son encapsulados en paquetes IP, con el número del protocolo al inicio. La Figura 26, muestra el formato de un paquete VRRPV2.

0	3	7	15	23	31
Version	Tipo	ID de Router Virtual	Prioridad	cuenta de direcciones IP	
Tipo de Autenticacion		Intervalo de Advertencia	Checksum		
Direccion IP 1					
.					
Direccion IP n					
Dato de autenticacion 1					
Dato de autenticacion 2					

Figura 26. Formato del paquete VRRPv2

Un paquete VRRP se compone de los siguientes campos:

- **Versión:** Número de versión del protocolo, 2 para el VRRPv2 y 3 para el VRRPv3.
- **Tipo:** Tipo de paquetes VRRPv2 y VRRPv3, solo un tipo de paquete VRRP es presente, que es, advertencia VRRP, quien se representa con 1.
- **ID de router virtual:** El ID del router virtual, que es, el ID del grupo VRRP, en el rango de 1-255.
- **Prioridad:** Es la prioridad del router en el grupo VRRP., en el rango de 0-255. Un valor alto representa una prioridad alta.
- **Cuenta de direcciones IP:** Número de las direcciones IPv4 virtuales del grupo VRRP. Un grupo VRRP puede tener múltiples direcciones virtuales.
- **Tipo de autenticación:** Tipo de autenticación. 0 cuando no hay autenticación, 1 cuando es una autenticación de texto simple y 2 cuando es una autenticación MD5. VRRP versión 3 no soporta autenticación MD5.
- **Intervalo de advertencia:** Es el intervalo para enviar paquetes de advertencia. Para el VRRPv2 el intervalo es en segundo y por defecto es 1. Para VRRPv3 el intervalo es en centisegundos y por defecto es 100.
- **Checksum:** son 16 bits de checksum para validar el dato en el paquete VRRP.
- **Dirección IP:** Es la dirección IP virtual del grupo VRRP. La cuenta de direcciones IP define el número de direcciones virtuales pertenecientes al grupo VRRP.
- **Dato de autenticación:** Llave de autenticación. Este campo se usa para autenticación simple y es 0 para cualquier otro modo de autenticación.

8.9 Tipos de paquetes VRRP

Modo de protocolo estándar VRRP sólo define anuncio VRRP. Sólo el maestro de un grupo VRRP envía periódicamente anuncios VRRP, y las copias de seguridad no envían anuncios VRRP.

El modo de equilibrio de carga de VRRP define los siguientes tipos de paquetes:

- **Advertisement:** VRRP anuncia el estado y la información acerca de la VF que está en el estado activo grupo VRRP. Tanto el maestro y las copias de seguridad envían periódicamente anuncios VRRP.
- **Request:** Si una copia de seguridad no es el propietario VF, envía una petición para pedir al maestro para asignar una dirección MAC virtual.
- **Reply:** Cuando se recibe una petición, el maestro envía una respuesta al enrutador de respaldo para asignar una dirección MAC virtual. Después de recibir la respuesta, el router de copia de seguridad crea una VF que se corresponde con la dirección MAC virtual, y luego se convierte en el dueño de esta VF.
- **Release:** Cuando un propietario de VF falla, el router que se hace cargo de su responsabilidad envía un comunicado después de un período determinado de tiempo para notificar a los demás routers del grupo VRRP para eliminar la VF del dueño fallido VF.

8.10 Aplicación de VRRP ejemplos

- **Master/Backup:** En el modo master/backup, solo el master reenvía paquetes. Cuando el master falla, un nuevo master es elegido de los backups, este modo requiere solo un grupo VRRP, en el que cada router tendrá una prioridad diferente, y el único con la prioridad más alta será el master, a como se muestra en la Figura 27

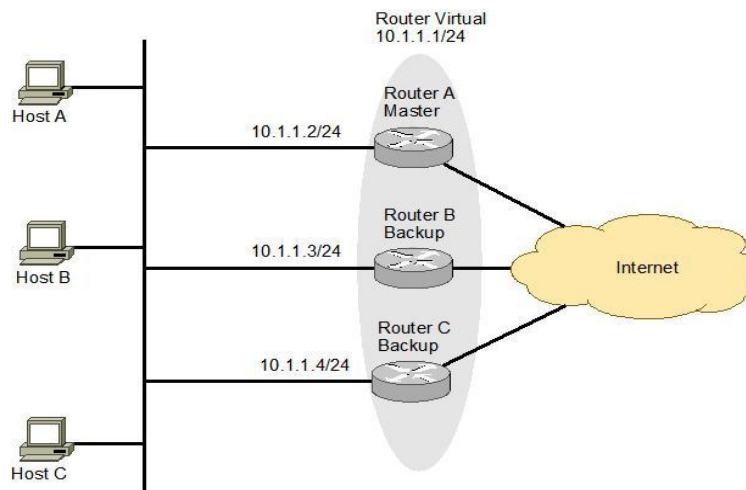


Figura 27. VRRP en modo master/backup

Asumiendo que el router A es el maestro, y también puede enviar paquetes a las redes externas, donde el Router B y el Router C son los respaldos y son quienes están en estado de escucha. Si el Router A falla, Router B y Router C elijen un nuevo master para el reenvío de paquetes hacia los demás host en la LAN.

8.11 Reparto de carga

Más de un grupo VRRP se puede crear en una interfaz de un router para permitir que el router sea el maestro de un grupo VRRP, pero una copia de seguridad de otro al mismo tiempo.

En el modo de reparto de carga, múltiples routers proporcionan servicios al mismo tiempo. Este modo requiere dos o más grupos VRRP, cada uno de los cuales comprende un maestro y uno o más copias de seguridad. Los patrones de los grupos VRRP son asumidos por los diferentes routers, como se muestra en la Figura 28:

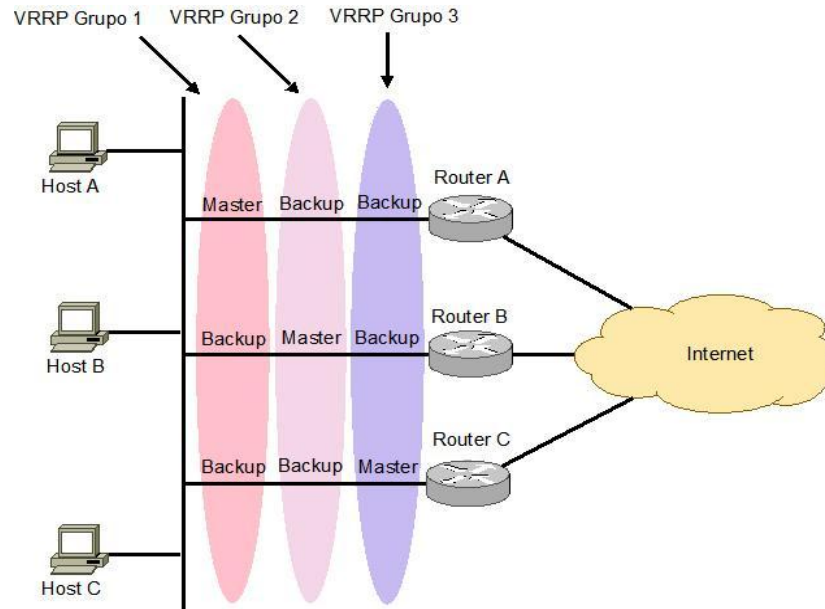


Figura 28. VRRP en modo de reparto de carga

Un router puede estar en varios grupos VRRP y mantenga una prioridad diferente en un grupo diferente.

Como se muestra en la Figura 28, los siguientes grupos VRRP están presentes:

- **VRRP grupo 1:** Router A es el maestro; Router B y Router C son los respaldos.
- **VRRP grupo 2:** Router B es el maestro; Router A y Router C son los respaldos.
- **VRRP grupo 3:** Router C es el maestro; Router A y Router B son los respaldos.

Para compartir la carga entre el Router A, Router B y Router C, las máquinas de la LAN se deben configurar para utilizar VRRP grupo 1, 2 y 3 como las puertas de enlace predeterminadas. Al configurar las prioridades de VRRP, asegúrese de que cada router tiene una prioridad tal en cada grupo VRRP que tomará el papel que se espera en el grupo.

9. FRAME RELAY

9.1 Definición de Frame Relay

Frame Relay es un protocolo WAN de alto rendimiento que funciona en las capas físicas y de enlace de datos del modelo de referencia OSI.

Frame Relay se ha convertido en la tecnología WAN más utilizada del mundo. Grandes empresas, gobiernos, ISP y pequeñas empresas usan Frame Relay, principalmente a causa de su precio y flexibilidad. A medida que las organizaciones crecen y dependen cada vez más de un transporte de datos fiable, las soluciones de líneas arrendadas tradicionales se vuelven imposibles de costear. El ritmo de los cambios tecnológicos y las fusiones y adquisiciones en la industria de networking demandan y exigen más flexibilidad.

Frame Relay reduce los costos de redes a través del uso de menos equipo, menos complejidad y una implementación más fácil. Aún más, Frame Relay proporciona un mayor ancho de banda, mejor fiabilidad y resistencia a fallas que las líneas privadas o arrendadas. Debido a una mayor globalización y al crecimiento de excesivas topologías de sucursales, Frame Relay ofrece una arquitectura de red más simple y un menor costo de propiedad.

9.2 Rentabilidad de Frame Relay

Frame Relay es una opción más rentable por dos motivos. En primer lugar, con líneas dedicadas, los clientes pagan por una conexión de extremo a extremo. Esto incluye el bucle local y el enlace de red. Con Frame Relay, los clientes sólo pagan por el bucle local y por el ancho de banda que compran al proveedor de red. La distancia entre los nodos no es importante. Mientras están en un modelo de líneas dedicadas, los clientes usan líneas dedicadas proporcionadas en incrementos de 64 kbps, los clientes Frame Relay pueden definir sus necesidades de circuitos virtuales con más granularidad, con frecuencia en incrementos pequeños como 4 kbps.

El segundo motivo de la rentabilidad de Frame Relay es que comparte el ancho de banda en una base más amplia de clientes. Comúnmente, un proveedor de red puede brindar servicio a 40 clientes o más de 56 kbps, en un circuito T1. El uso de líneas dedicadas requeriría más DSU/CSU (uno para cada línea), así como también enrutamiento y conmutación más complicados. Los proveedores de red ahorran dado que hay menos equipos para comprar y mantener.

9.3 Flexibilidad de Frame Relay

Un circuito virtual proporciona considerable flexibilidad en el diseño de red. Si observa la figura, puede ver que todas las oficinas de Span se conectan a la nube Frame Relay a través de sus respectivos bucles locales. Lo que sucede en la nube no debe preocuparle en este momento. Lo único que importa es que cuando una oficina de Span desea comunicarse con cualquier otra oficina de Span, no necesita más que conectarse a un circuito virtual que lleve a la otra oficina.

En Frame Relay, el extremo de cada conexión tiene un número de identificación denominado identificador de conexión de enlace de datos (DLCI, Data Link Connection Identifier). Cualquier estación puede conectarse con otra simplemente si escribe la dirección de esa estación y el número de DLCI de la línea que necesita usar. En una sección posterior,

aprenderá que cuando se configura Frame Relay, todos los datos de todos los DLCI configurados fluyen a través del mismo puerto del router. Intente reflejar la misma flexibilidad usando líneas dedicadas. No sólo es complicado, sino que también requiere un número considerablemente mayor de equipos.

9.4 Requisitos de Frame Relay vs líneas dedicadas

Las figuras muestra una comparación del costo representativa para conexiones ISDN y Frame Relay comparables. Si bien los costos iniciales en el caso de Frame Relay son más altos que para ISDN, el costo mensual es considerablemente más bajo. Frame Relay es más fácil de administrar y configurar que ISDN. Además, los clientes pueden incrementar su ancho de banda a medida que sus necesidades crezcan en el futuro.

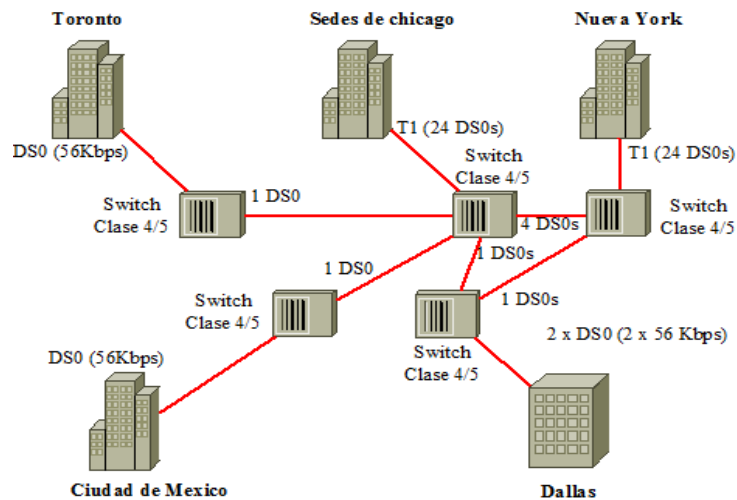


Figura 29. Requisitos líneas dedicadas

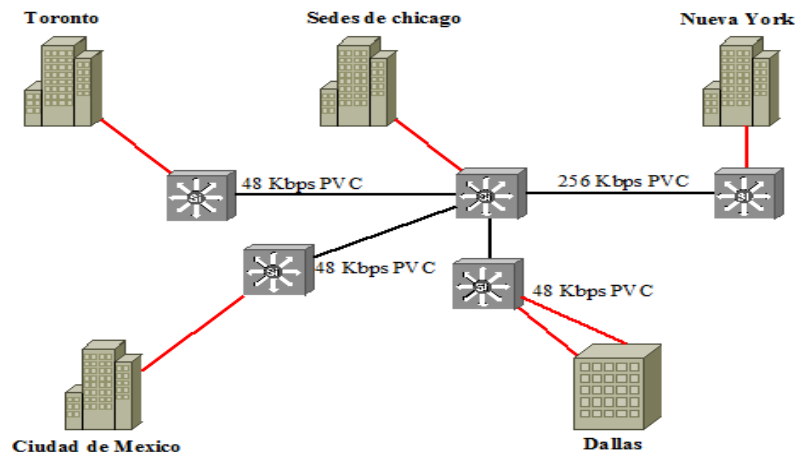


Figura 30. Requisitos Frame Relay

	ISDN de 64Kbps	Frame Relay de 54 Kbps
Cargo mensual por bucle Local	\$185	\$85
Configuración del ISP	\$380	\$750
Equipo	\$700	\$1600
Cargo mensual del ISP	\$195	\$195
Cargo por única vez	\$1080	\$2660
Cargo mensuales	\$380	\$280
<i>Equipos : Router ISDN \$700</i>		
<i>Router Cisco \$1600</i>		

Cuadro 20. Costos de Frame Relay

9.5 Funcionamiento de Frame Relay

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.

Cuando las empresas de comunicaciones usan Frame Relay para interconectar las LAN, un router de cada LAN es el DTE. Una conexión serial, como una línea arrendada T1/E1, conecta el router al switch Frame Relay de la empresa de comunicaciones en el punto de presencia (POP) más cercano para la empresa. 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. Se encuentra en las instalaciones del cliente y se conecta con el puerto del switch en la red del proveedor de servicio. A su vez, el proveedor de servicio interconecta los switches Frame Relay.

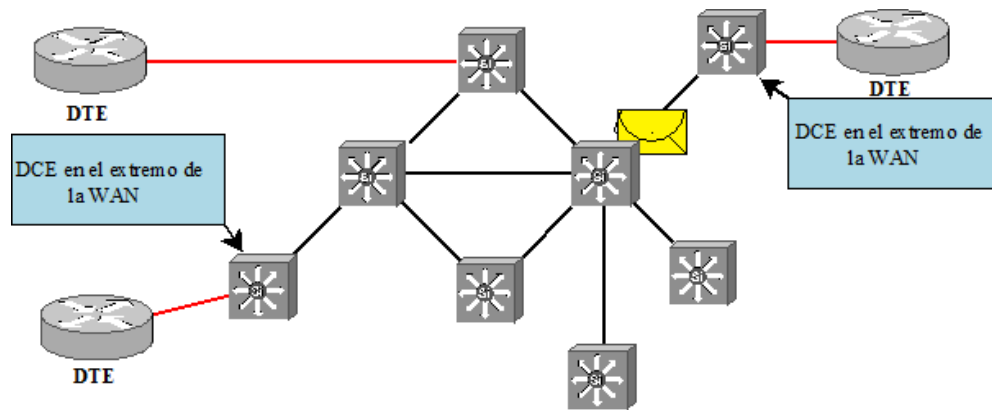


Figura 31. Operación Frame Relay

- DTE envía tramas al switch de DCE en el extremo de la WAN.
- Las tramas se mueven de Switchs a Switchs 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.

9.6 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.

Hay dos formas de establecer VC:

- 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.

9.7 DLCI

Data Link Connection Identifier (DLCI) 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.

- Conexiones Frame Relay identifican circuitos virtuales por identificador de conexión de datos Link (DLCI) números.
- Un número DLCI asocia una dirección IP con un circuito virtual específico.
- Números DLCI sólo tienen importancia local.
- Números DLCI se asignan generalmente por el proveedor de Frame Relay.

Muy probablemente no el mismo en cualquier lado del Switch Frame Relay

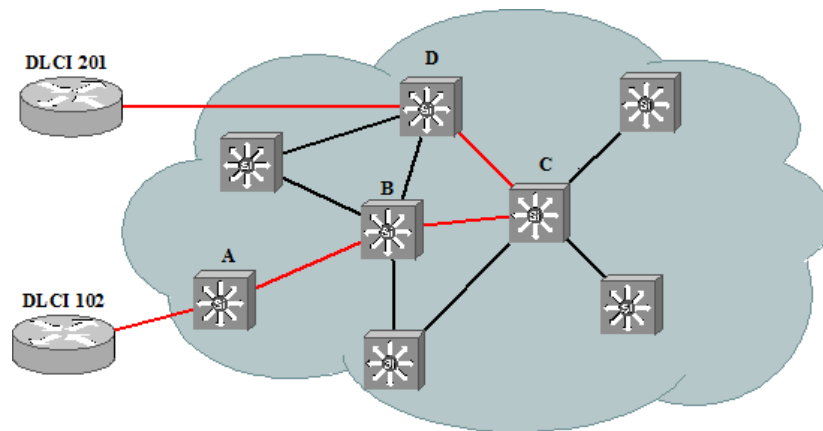


Figura 32. Circuitos Virtuales

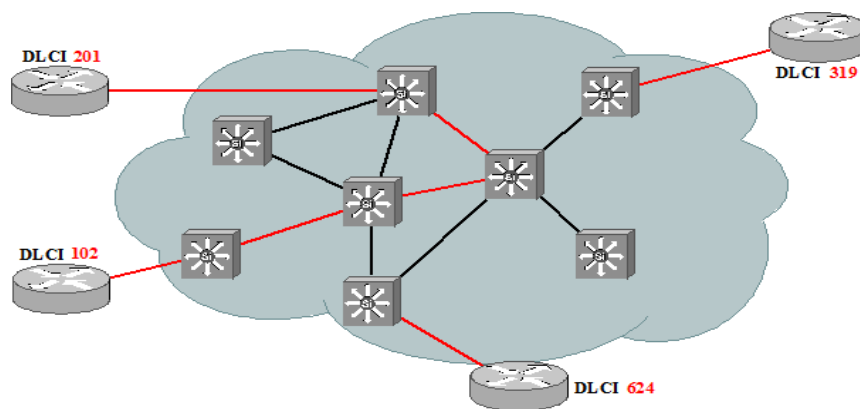


Figura 33. Importancia DLCI

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.

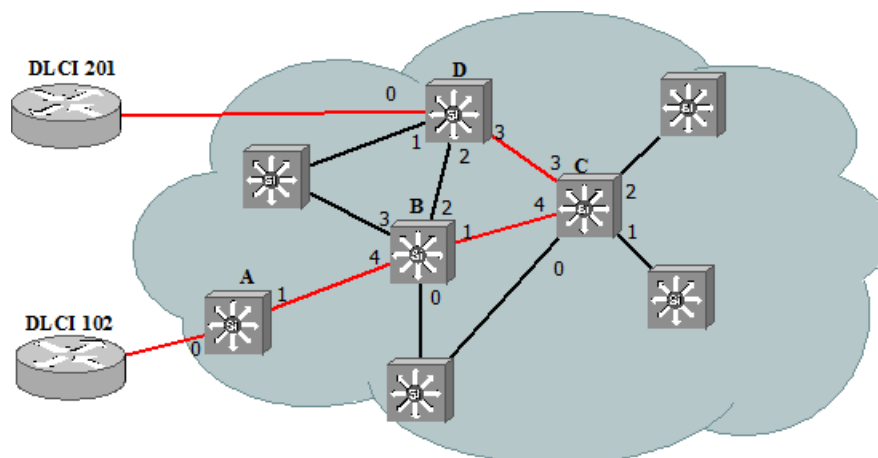


Figura 34. Identificación de VC

Tramo	VC	Puerto	VC	Puerto
A	102	0	432	1
B	432	4	119	1
C	119	4	579	3
D	579	3	201	0

Cuadro 21. Identificación de VC por cada equipo

9.8 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 construido de forma automática o estática según la topología de Frame Relay

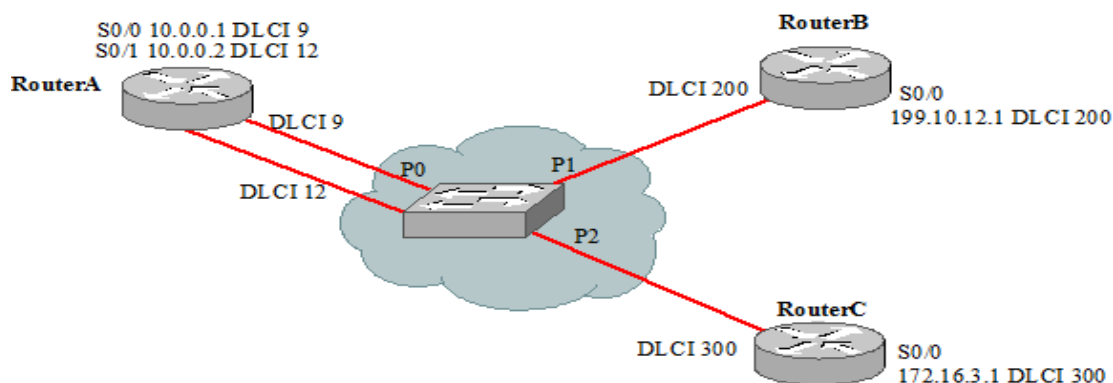


Figura 35. Configuración de Frame Relay

Frame Relay Switching table (Frame Relay Switch)			
IN_Port	IN_DLCI	OUT_Port	OU_DLCI
P0	9	P1	200
P0	12	P1	300

Cuadro 22. Configuración de Frame Relay por cada equipo

Frame Relay Map Router	
199.10.12.1	DLCI 9
172.16.3.1	DLCI 12

Cuadro 23. Map Frame Relay

9.8.1 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, subinterfaces son tratados como interfaces físicas. Con subinterfaces, el costo de implementar múltiples circuitos virtuales de Frame Relay se reduce, debido a que sólo un puerto es necesario en el router.

9.9 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. Para comprender el funcionamiento, resulta útil entender cómo se relaciona con los niveles más bajos del modelo OSI.

En principio, 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. Los señaladores están representados como el número hexadecimal 7E o como el número binario 01111110. Después de haber encapsulado el paquete, Frame Relay pasa la trama a la capa física para su transporte.

El router CPE encapsula cada paquete de Capa 3 dentro de un encabezado Frame Relay y tráiler antes de enviarlo a través del VC. El encabezado y el tráiler están definidos por la especificación de Servicios de portadora del Procedimiento de acceso al enlace para Frame Relay (LAPF), ITU Q.922-A. Específicamente, el encabezado Frame Relay (campo de dirección) incluye lo siguiente:

- **DLCI:** el DLCI de 10 bits es la esencia del encabezado Frame Relay. Este valor representa la conexión virtual entre el dispositivo DTE y el switch. Cada conexión virtual multiplexada en el canal físico está representada por un único DLCI. Los valores de DLCI tienen importancia local solamente, lo que significa que son únicos sólo para

el canal físico en el que residen. Por lo tanto, los dispositivos situados en los extremos opuestos de una conexión pueden usar valores DLCI distintos para referirse a la misma conexión virtual.

- **Dirección extendida (EA):** si el valor del campo EA es 1, el byte actual está determinado como el último octeto DLCI. Si bien las implementaciones actuales de Frame Relay usan un DLCI de dos octetos, esta capacidad no permite DLCI más largos en el futuro. El octavo bit de cada byte del campo Dirección indica la EA.
- **C/R:** el bit que sigue al byte de DLCI más significativo en el campo Dirección. El bit C/R no está definido en este momento.
- **Control de congestión:** incluye 3 bits que controlan los mecanismos de notificación de congestión de Frame Relay. Los bits FECN, BECN y DE son los últimos tres bits en el campo Dirección. El control de congestión se explica en un tema posterior.

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. La FCS determina si hubo errores en el campo Dirección de Capa 2 durante la transmisión. La FCS se calcula antes de la transmisión a través del nodo emisor, y el resultado se inserta en el campo FCS. En el otro extremo, un segundo valor de FCS se calcula y compara con la FCS de la trama. Si los resultados son iguales, se procesa la trama. Si existe una diferencia, la trama se descarta. Frame Relay no notifica el origen cuando se descarta una trama. El control de errores tiene lugar en las capas superiores del modelo OSI.

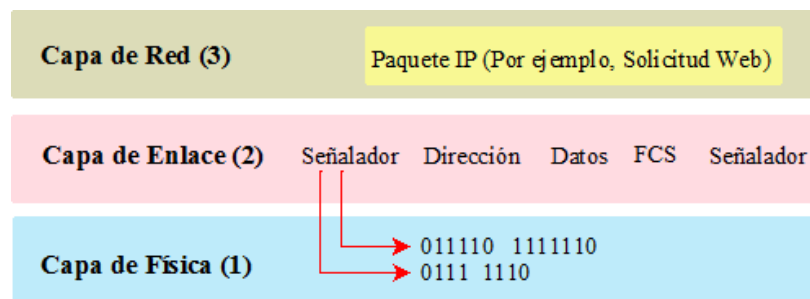


Figura 36. Encapsulación Frame Relay

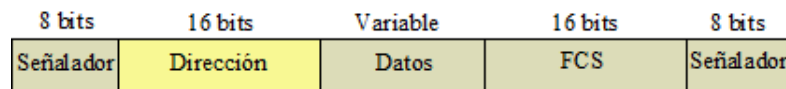


Figura 37. Trama Frame Relay

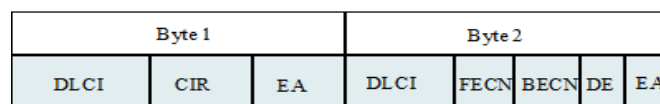


Figura 38. Formato de trama Frame Relay

9.10 LMI

LMI básicamente extiende la funcionalidad de Frame Relay a través de:

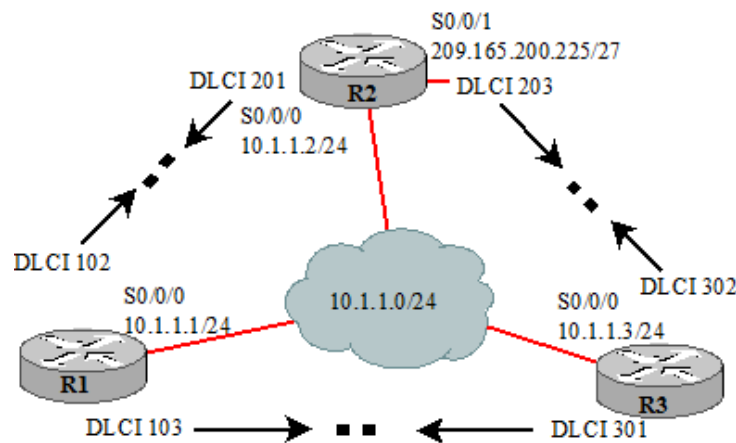
- Hacer que los DLCI importancia mundial en lugar de un importante impacto local.
- Creación de un mecanismo de señalización entre el Router y el Switch Frame Relay, que podría informar sobre el estado del enlace.
- Apoyo a la multidifusión.
- Proporcionar números DLCI que son de importancia mundial hace que la configuración automática de la hoja de Frame Relay posible.

LMI utiliza paquetes de controles para verificar el enlace Frame Relay y para asegurar el flujo de datos.

Cada circuito virtual, representado por su número DLCI, puede tener uno de tres estados de conexión:

- Activo
- Inactivo
- Suprimido

El switch Frame Relay reporta la información de estado para el mapa Frame Relay en el router local.



```
R1 # show frame - raley Lmi
LMI Statistics for interface serial 0/0/0 (frame relay DTE) LMI
TYPE = ANSI
Invalid unnumbered info 0 Invalid Prot Disc 0
Invalid dummy Call Ref 0 Invalid Msg Type 0
Invalid Status Message 0 Invalid Lock shift 0
Invalid Information ID 0 Invalid Report IE Len 0
Invalid Report Request 0 Invalid Keep IE Len 0
Nun Status Eng. Sent 9 Nun Status msgs Rcvd 0
Nun Update Status Rcvd 0 Nun Status Timeouts 9
```

Figura 39. Estadísticas LMI

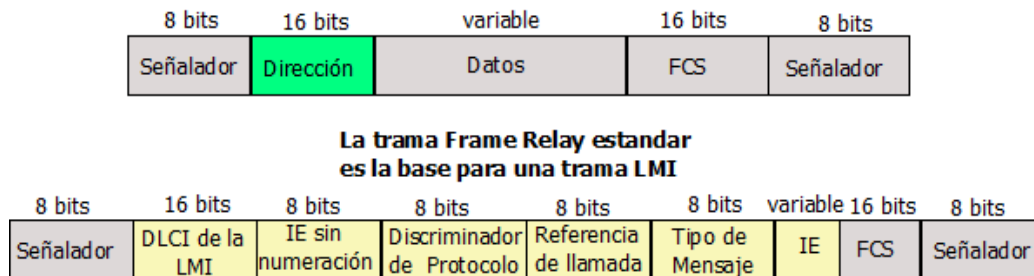
Identificadores de VC	Tipos de VC
0	LMI (ANSI, UIT)
1...15	Se reserva para uso futuro
992...1007	0LLM
2008...1022	Se reserva para uso futuro (ANSI, UIT)
1019...1020	Multicasting (Cisco)
1023	LMI Cisco

Cuadro 24. Identificadores LMI

9.11 Formato de trama LMI

Los mensajes LMI se envían a través de una variante de las tramas LAPF. El campo de dirección lleva uno de los DLCI reservados. Seguido al campo DLCI se encuentran los campos de control, de discriminación de protocolos y de referencia de llamada, los cuales no cambian. El cuarto campo indica el tipo de mensaje LMI.

Los mensajes de estado ayudan a verificar la integridad de los enlaces físicos y lógicos. Esta información resulta fundamental en un entorno de enrutamiento, ya que los protocolos de enrutamiento toman decisiones según la integridad del enlace.

**Figura 40. Formato de Trama LMI**

9.12 ARP inverso

En configuraciones multipunto donde los enrutadores utilizan el protocolo ARP inverso para enviar una consulta utilizando el número de DLCI para encontrar una dirección IP remota.

Como otros routers responden a las consultas ARP Inverso, el router local puede construir su mapa de Frame Relay automáticamente.

Para mantener el mapa Frame Relay, routers intercambiar mensajes Inverse ARP cada 60 segundos por defecto.

10. LISTAS DE CONTROL DE ACCESO

10.1 Definición de Listas de Control de Acceso (ACL)

Las Listas de Control de Acceso (ACL) son una colección secuencial de sentencias de permiso o rechazo que se aplican a direcciones o protocolos de capa superior. Los routers proporcionan capacidades de filtrado de tráfico a través de las listas de control de acceso (ACL). En esta práctica, conocerá las ACL estándar y extendidas como medio de controlar el tráfico de red y de qué manera se usan las ACL como parte de una solución de seguridad (cortafuegos).

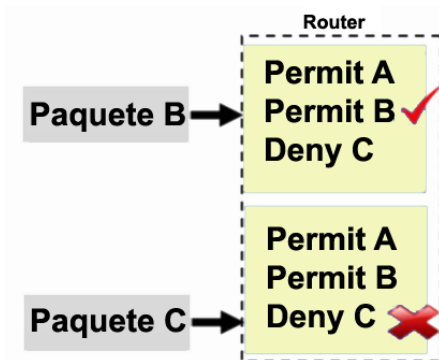


Figura 41. Ejemplo del modo de operación de las ACL

Las Listas de Control de Acceso (ACL) son listas de instrucciones que se aplican a una interfaz del router, a como se muestra en la Figura 41. Estas Listas indican al router qué tipos de paquetes se deben aceptar y qué tipos de paquetes se deben denegar. La aceptación y rechazo se pueden basar en ciertas especificaciones, como dirección origen, dirección destino y número de puerto. Cualquier tráfico que pasa por la interfaz debe cumplir ciertas condiciones que forman parte de la ACL. Las ACL se pueden crear para todos los protocolos enrutados de red, como IP e IPX, para filtrar los paquetes a medida que pasan por un router. Es necesario definir una ACL para cada protocolo habilitado en una interfaz si desea controlar el flujo de tráfico para esa interfaz. Por ejemplo, si su interfaz de router estuviera configurada para IP, AppleTalk e IPX, sería necesario definir por lo menos tres ACL. Cada ACLs sobre cada interfaz, actúa en un sentido, distinguiendo tanto sentido de entrada como de salida. Se puede definir diferentes ACLs y luego instalarlas sobre los interfaces del router según convenga al administrador de la red.

10.2 Razones para el uso de ACLs

Hay muchas razones para crear ACLs. Por ejemplo, las ACL se pueden usar para:

- Limitar el tráfico de red y mejorar el rendimiento de la red. Por ejemplo, las ACL pueden designar ciertos paquetes para que un router los procese antes de procesar otro tipo de tráfico, según el protocolo. Esto se denomina colocación en cola, que asegura que los router no procesarán paquetes que no son necesarios. Como resultado, la colocación en cola limita el tráfico de red y reduce la congestión.
- Brindar control de flujo de tráfico. Por ejemplo, las ACL pueden restringir o reducir el contenido de las actualizaciones de enrutamiento. Estas restricciones se usan para limitar la propagación de la información acerca de redes específicas por toda la red.

- Proporcionar un nivel básico de seguridad para el acceso a la red. Por ejemplo, las ACL pueden permitir que un host acceda a una parte de la red y evitar que otro acceda a la misma área. Al Host A se le permite el acceso a la red de Recursos Humanos, y al Host B se le deniega el acceso a dicha red. Si no se configuran ACL en su router, todos los paquetes que pasan a través del router supuestamente tendrían acceso permitido a todas las partes de la red.
- Se debe decidir qué tipos de tráfico se envían o bloquean en las interfaces del router. Por ejemplo, se puede permitir que se enrute el tráfico de correo electrónico, pero bloquear al mismo tiempo todo el tráfico de telnet.

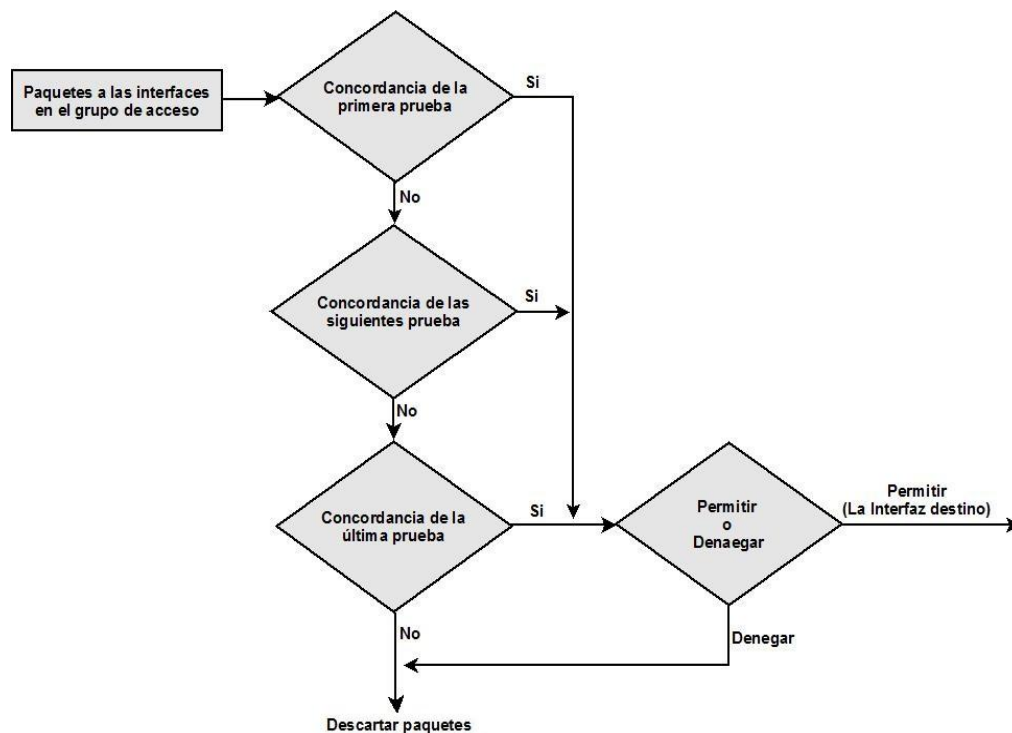


Figura 42. Ciclo de las ACL

10.3 Reglas de Funcionamiento de ACL

- Se puede configurar una sola lista en cada interfaz por sentido del tráfico (entrante/saliente), por protocolo de enrutamiento (IP, IPX, etc.).
- Cada lista de acceso es identificada por un ID único (un número o alfanumérico). En el caso de las listas numeradas, este ID identifica el tipo de lista de acceso y las diferencia de otras semejantes.
- Cada paquete que ingresa o sale de la interfaz es comparado con cada línea de la lista secuencialmente, en el mismo orden que fueron ingresadas, comenzando por la primera ingresada.
- La comparación se sigue realizando hasta tanto se encuentre una coincidencia. Una vez que el paquete cumple con la condición de una línea, se ejecuta la acción indicada y no se sigue comparando.

- Hay un deny all implícito al final de cada lista de acceso, que no es visible. Por lo tanto, si el paquete no coincide con ninguna de las premisas declaradas en la lista será descartado.
- Los filtros de tráfico saliente no afectan el tráfico originado en el mismo router.
- Las listas de acceso IP al descartar un paquete envían al origen un mensaje ICMP de “host de destino inalcanzable”.
- Tener en cuenta que al activar listas de acceso el router automáticamente conmuta de fast switching a process switching.
 - Fast Switching – Feature de los routers ciscos que utiliza un cache del router para conmutar rápidamente los paquetes hacia el puerto de salida sin necesidad de seleccionar la ruta para cada paquete que tiene una misma dirección de destino.
 - Process Switching – Operación que realiza una evaluación completa de la ruta por paquete. Implica la transmisión completa del paquete al CPU del router donde será re-encapsulado para ser entregado a través de la interfaz de destino. El router realiza la selección de la ruta para cada paquete. Es la operación que requiere una utilización más intensiva de los recursos del router.

10.4 Tipos de listas de acceso IP e IPX

- Listas de acceso estándar
 - Listas IP: Permiten filtrar únicamente direcciones IP de origen.
 - Listas IPX: Filtran direcciones IPX de origen como destino.
- Listas de acceso extendida
 - Listas IP: Verifican direcciones de origen y destino, protocolo de capa 3 y puerto de capa 4.
 - Listas IPX: Permiten filtrar además de las direcciones de IPX de origen y destino, a través del valor del campo del protocolo del encabezado de capa 3 y el número de socket del encabezado de capa 4.
- Listas de acceso nombradas: Listas de acceso IP tanto estándar como extendidas que verifican direcciones de origen y destino, protocolos de capa 3 y puertos de capa 4, identificados con una cadena de caracteres alfanuméricos. A diferencia de las listas de acceso numeradas, se configuran en un submodo propio y son editables.
- Filtros IPX SAP: Se utilizan para controlar el tráfico de paquetes SAP tanto a nivel LAN como WAN. Son un mecanismo útil para controlar el acceso a los dispositivos IPX.
- Lista de acceso entrante: controlan el tráfico que ingresa al router a través del puerto en el que esta aplicada, y antes de que sea conmutado a la interfaz de salida.
- Listas de acceso saliente: controlan el tráfico saliente del router a través del puerto en que esta aplicada, una vez que ya ha sido conmutado.

10.5 Número de lista de acceso

El tipo de lista de acceso y protocolo de capa 3 que filtra se especifica a partir de un número de lista de acceso. El listado completo de números de lista de acceso se muestra en el Cuadro 25.

Número	Tipo según su numeración
1-99	IP estándar
100-199	IP extendida
200-299	Protocolo Type Code
300-399	DECnet
400-499	XNS estándar
500-599	XNS extendida
600-699	Apple Talk
700-799	48 bit MAC address estándar
800-899	IPX estándar
900-999	IPX extendida
1000-1099	IPX SAP
1100-1199	48 bit MAC address extendida
1200-1299	IPX Summary address

Cuadro 25. Número de listas de acceso según su tipo

10.6 Número de Puerto

En las listas de acceso IP extendidas, al especificar protocolo TCP o UDP se puede también especificar la aplicación que se desea filtrar a través del número de puerto, a como se muestra en el Cuadro 26.

Nombre del Protocolo	Número del Protocolo	Protocolo de Transporte
FTP Data	20	TCP
FTP	21	TCP
TELNET	23	TCP
SMTP	25	TCP
Time	37	UDP
TACACS	49	UDP
DNS	53	UDP
DHCP Server	67	UDP
DHCP Client	68	UDP
TFTP	69	UDP
Gopher	70	UDP
Finger	79	UDP

HTTP	80	TCP
POP3	110	TCP
RPC	111	UDP
NetBIOS name	137	UDP
NetBIOS Datag	138	UDP
NetBIOS session	139	UDP
SNMP	161	UDP
IRC	194	UDP

Cuadro 26. Número de puertos y protocolos**10.7 Mascara de Wildcard**

Las mascara de Wildcard son direcciones de 32 bits divididas en 4 octetos de 8 bits utilizadas para generar filtros de direcciones IP. Se utilizan en combinación de una dirección IP.

En las mascara de Wildcard el dígito en 0 (cero) indica una posición que debe ser comprobada, mientras que el dígito 1 (uno) indica una posición que carece de importancia.

La máscara se aplica a una dirección IP específica contenida en la declaración de la ACL y a la dirección de los paquetes a evaluar. Si ambos resultados son iguales, entonces se aplica al paquete el criterio de permisos o denegación enunciado en la línea correspondiente.

172.16.14.33 0.0.0.0

Indica que se debe seleccionar únicamente la dirección IP declarada

172.16.14.33 0.0.0.255

Selecciona todas las direcciones IP comprendidas entre 172.16.14.0 y 172.16.14.255 (no discrimina respecto del cuarto octeto).

172.16.14.33 0.0.255.255

Selecciona todas las direcciones IP de la red 172.16.0.0 comprendidas entre 172.16.0.0 y 172.16.255.255 (no discrimina respecto del número de host – los dos últimos bits).

10.7.1 Reglas de cálculos de una máscara Wildcard

- a. Cuando se desea filtrar una red completa, la máscara de Wildcard es el complemento de la máscara de red por defecto:

Dirección de red	172.16.0.0
Mascara de subred por defecto	255.255.0.0
Mascara de Wildcard	0.0.255.255

- b. Cuando se desea filtrar una subred completa, la máscara de Wildcard es el complemento de la máscara de subred que se esté aplicando:

Dirección de red	172.16.30.0/24
Máscara de subred por defecto	255.255.255.0
Máscara de Wildcard	0.0.0.255
Dirección de red	172.16.32.0/20
Máscara de subred por defecto	255.255.240.0
Máscara de Wildcard	0.0.15.255

10.8 Configuración de las Listas de Control de Acceso (ACL)

En el proceso de configuración de las listas de acceso deben distinguirse 2 etapas:

- Creación de la lista de acceso en modo de configuración global
- Asignación de esa lista a un puerto de modo entrante o saliente

10.8.1 Listas de acceso IP estándar

```
Router(config)#access-list [permit/deny] [IP origen]
Router(config)#interface serial 1
Router(config)#ip access-group [1-99] [in/out]
```

10.8.2 Listas de acceso IP extendidas

```
Router(config)#access-list [100-199][permit/deny][protocolo][IP origen] [IP
destino] [tipo de servicio]
Router(config)#interface serial 1
Router(config)#ip access-group [100-199] [in/out]
```

10.8.3 Listas de acceso IP nombradas

```
Router(config)#access-list [satandar/extended] [IP origen]
```

10.8.4 Lista de acceso IP nombrada estándar

```
Router(config)#[permit/deny] [IP origen]
```

10.8.5 Lista de acceso IP nombrada extendida

```
Router(config)#[permit/deny] [protocolo] [IP origen] [ip destino] [tipo de
servicio]
```

10.8.6 Lista de acceso IPX estándar

```
Router(config)#acces-list [800-899][permit/deny][IPX origen] [IPX
destino]
```

Nota: el valor -1 en el campo correspondiente a las direcciones de origen o destino indica cualquier origen o destino.

```
Router(config)#interface serial 1
Router(config)#ip access-group [800-899] [in/out]
```

10.8.7 Listas de acceso IPX extendida

```
Router(config)#access-list[900-999][permit/deny][protocolo][IPX  origen][IPX
destino][socket]
Router(config)#interface serial 1
Router (config)#ip access-group [900-999] [in/out]
```

10.8.8 Filtros IPX SAP

```
Router(config)#access-list[900-999][permit/deny][protocolo][IPX
origen][servicio]
```

Nota: el valor 0 en el campo servicio indica “todos los servicios”

```
Router(config)#interface serial 1
Router(config)#ip access-group [input/output]-sap-filter [1000-1099]
```

10.9 Filtros aplicados a terminales virtuales

Se pueden aplicar a las terminales virtuales (acceso por telnet) listas de acceso estándar a fin de limitar las direcciones IP a partir de las cuales se podrá conectar al dispositivo vía telnet.

```
Router(config)#access-list permit host 172.16.10.3
Router(config)#line vty 0 4
Router(config)#access-class 10 in
```

10.10 Comandos especiales

En algunos casos especiales, un comando puede reemplazar una máscara Wildcard, con el mismo efecto:

xxx.xxx.xxx.xxx	0.0.0.0	host	xxx.xxx.xxx.xxx
0.0.0.0	255.255.255.255	any	
-1		any IPX network	
remark		Utilizado en lugar de la opción permit/deny, permite insertar comentarios en una lista de acceso.	

Si no especifica una máscara por defecto, el sistema operativo asume la máscara por defecto: 0.0.0.0 en consecuencia, las siguientes formas son equivalentes:

```
Router(config)#access-class 10 permit 172.16.10.3 0.0.0.0
Router(config)#access-class 10 permit host 172.16.10.3
```

```
Router(config)#access-class 10 permit 172.16.10.3
```

10.11 Monitoreo de las listas

Muestra las listas y el contenido de todas las ACL o una en particular. No permite verificar a que interfaz esta aplicada.

```
Router# show ip Access-list [#]
```

Muestras solamente la configuración de ACL IP

```
Router#show ip Access-list
```

Muestras solamente la configuración de ACL IPX

```
Router#show ipx Access-list [#]
```

Muestra los puertos que tienen aplicadas ACL IP

```
Router#show ip interface
```

Muestra los puertos que tienen aplicadas ACL IPX

```
Router#show ipx interface
```

Muestra tanto las listas de acceso configuradas, como la que se encuentran aplicadas a cada interfaz

```
Router#show running-config
```

10.12 Tips de aplicación

- Antes de comenzar a trabajar sobre una lista de acceso existente, desvínculela de todas las interfaces a las que se encuentre asociada.
- No se puede remover una única línea de una lista de acceso numerada (no son editables).
- Cada vez que se agrega una línea a la lista de acceso, esta se ubicara a continuación de las líneas existente, al final.
- Organice su lista de acceso de modo que los criterios más específicos estén al comienzo de la misma, y luego las premisas más generales.
- Coloque primero los permisos y luego las denegaciones.
- Toda lista de acceso debe incluir al menos un comando *permit*
- Las listas de acceso no filtran el trafico originado en el router
- Una misma lista de acceso puede ser asignada a varias interfaces en el mismo dispositivo, tanto en modo entrante como en modo saliente.
- Las *listas de acceso estándar* deben colocarse lo más cerca posible del destino del trafico
- Las *listas de acceso extendidas* deben colocarse lo más cerca posible del origen del tráfico que será denegado.

10.13 Borrar las ACLs existentes

Comprobar si existe alguna ACL definida

```
show access-list
```

Borrar las ACL

```
no access-list <n°>
```

Comprueba también si algún interfaz tiene aplicada ninguna ACL

```
show ip interface
```

Desactivar las ACL, desde el modo de configuración del interfaz.

```
no ip access-group <n°> in/out
```

10.14 Ejemplo de usos de Listas de Control de Accesos (ALC)

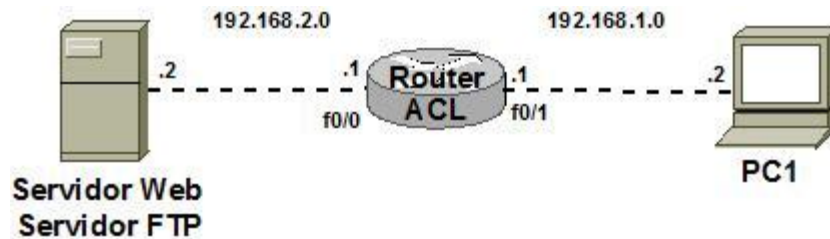


Figura 43. Diagrama de muestra de las ACL

Escenario 1:

Suponiendo que se desea denegar cualquier tipo de tráfico de la PC1 hacia el Servidor Web y FTP, ¿Cuál sería el tipo y las lista de acceso a implementar?

Solución:

La lista de acceso que se implementaría sería de tipo estándar, debido a que no será necesario especificar la dirección origen y destino, tomando en cuenta de igual manera que no se hará especificación de ningún protocolo.

A la hora de crear la lista de acceso quedaría de la siguiente manera

```
Router ACL(config)#access-list 10 deny 192.168.2.2 0.0.0.0
Router ACL(config)#access-list 10 permit any
```

A la hora de implementarlo en Router ACL se hará de la siguiente manera:

```
Router ACL(config)#interface fastEthernet 0/1  
Router ACL(config-if)#ip access-group 10 in
```

Escenario 2:

Suponiendo que se desea denegar el tráfico de la PC1 al Servidor Web y la transferencia de archivos mediante el mismo servidor que también funciona como servidor FTP.

Solución:

Las listas de control de accesos que se implementarían serían de tipo nombrada-extendida, ya que se necesitaría conocer las direcciones origen-destino y además se necesitara saber el protocolo que será denegado, las listas quedarían de la siguiente manera:

```
Router ACL(config)#ip access-list extended denegar  
Router ACL(config-ext-nacl)#deny tcp host 192.168.2.2 host 192.168.1.2 eq 80  
Router ACL(config-ext-nacl)#deny tcp host 192.168.2.2 host 192.168.1.2 eq ftp  
Router ACL(config-ext-nacl)#permit tcp any any
```

A la hora de implementarlas en el Router ACL, se debe aplicar en la interfaz f0/1, recordando que una lista de acceso nombrada-extendida se debe emplear lo más cerca posible de la dirección origen, en este escenario la lo más cerca posible de la PC1, la implementación se haría de la siguiente manera:

```
Router ACL(config)#interface fastEthernet 0/1  
Router ACL(config-if)#ip access-group puerto80 in
```

De esta manera el acceso de la PC1 al servicio web del servidor no es posible llevarse a cabo, porque ha sido denegado.

11. TELEFONÍA SOBRE VoIP

11.1 Protocolo de voz sobre IP

EL protocolo de VoIP se encarga de dividir en paquetes los flujos de audio para transportarlos sobre redes basadas en IP.

11.2 Principales componentes de la VOIP

Para la transmisión de voz sobre una red IP, el estándar define tres elementos fundamentales en su estructura:

- **Terminales:** Son los puntos finales de la comunicación y pueden ser implementados como:
 - **Hardware:** Un teléfono IP es un terminal que tiene soporte VoIP nativo y puede conectarse directamente a una red IP.
 - **Software:** Un softphone es una aplicación audio ejecutable desde PC que se comunica con las PABX a través de la LAN. Para interactuar con el usuario se basa en la utilización de un micrófono y altavoz o mediante un teléfono USB.
 - **Servidor:** Provee el manejo y funciones administrativas para soportar el enrutamiento de llamadas a través de la red. Este servidor puede adoptar diferentes nombres dependiendo del protocolo de señalización utilizado. Así en un sistema basado en el protocolo H.323, el servidor es conocido como Gatekeeper; en un sistema SIP, servidor SIP; y en un sistema basado en MGCP o MEGACO, Call Agent (Agente de llamadas). El servidor es un elemento opcional, normalmente implementado en software, y en caso de existir, todas las comunicaciones pasarían por él.
- **Gateways:** Enlace de la red VoIP con la red telefónica analógica o RDSI. Se encarga de adaptar las señales de estas redes a VoIP y viceversa, actuando de forma totalmente transparente para el usuario. El gateway posee, además de puertos LAN, interfaces de conexión a estas redes: FXO, FXS, E&M, BRI, PRI, G703/G.704
- **Red IP:** Provee conectividad entre todos los terminales. La red IP puede ser una red IP privada, una Intranet o Internet.

11.3 Funcionamiento de VoIP

VoIP funciona digitalizando la voz en paquetes de datos, enviándola a través de la red y reconvirtiéndola a voz en el destino. Básicamente el proceso comienza con la señal analógica del teléfono que es digitalizada en señales PCM por medio del codificador/decodificador de voz (*codec*). Las muestras PCM son pasadas al algoritmo de compresión, el cual comprime la voz y la fracciona en paquetes (*Encapsulamiento*) que pueden ser transmitidos para este caso a través de una red privada WAN. En el otro extremo de la nube se realizan exactamente las mismas funciones en un orden inverso. El flujo de un circuito de voz comprimido es el mostrado en la Figura 44.

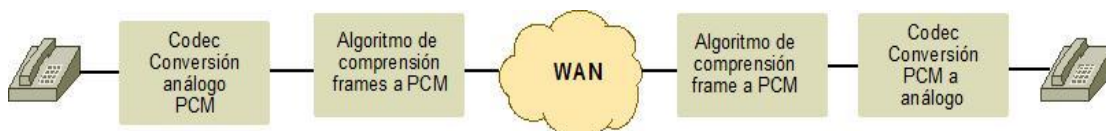


Figura 44. Funcionamiento de VoIP

Es importante tener en cuenta también que todas las redes deben tener de alguna forma las características de direccionamiento, enrutamiento y señalización. El direccionamiento es requerido para identificar el origen y destino de

las llamadas, también es usado para asociar las clases de servicio a cada una de las llamadas dependiendo de la prioridad. El enrutamiento por su parte encuentra el mejor camino a seguir por el paquete desde la fuente hasta el destino y transporta la información a través de la red de la manera más eficiente, la cual ha sido determinada por el diseñador. La señalización alerta a las estaciones terminales y a los elementos de la red su estado y la responsabilidad inmediata que tienen al establecer una conexión.

El Cuadro 27, muestra en detalle la distribución de los protocolos VoIP dentro del modelo OSI:

N° de Capa	Nombre de la Capa	Servicios
7	Aplicación	Asterisk
6	Presentación	G.729/G.711/GDM/Speex
5	Sesión	SIP/H.323/MGCP
4	Transporte	UDP/RTP/SRTP
3	Red	IP/CBWFQ/WRED
2	Enlace	Frame Relay/ATM/PPP
1	Física	Ethernet/xDSL

Cuadro 27. Distribución de los protocolos VoIP dentro del modelo OSI

Como se puede ver en la Figura 44, la voz sobre IP está compuesta por diversos protocolos envolviendo varias capas del modelo OSI. De cualquier forma, VoIP es una aplicación que funciona sobre las redes IP actuales, tratando principalmente las capas de transporte, sesión, presentación y aplicación.

En la capa de transporte, la mayor parte de estos protocolos usa el RTP/RTCP, siendo el primero un protocolo de medio y el segundo un protocolo de control. La excepción es IAX, que implementa un transporte de medio propio. Todos ellos usan UDP para transportar la voz.

En la capa de sesión entran los protocolos de voz sobre IP propiamente dichos, tales como: H323, SIP, MGCP, IAX e SCCP.

En la capa de sesión los CODECs definen el formato de presentación de voz con sus diferentes variaciones de compresión.

11.4 Transporte de la Voz en redes IP

Los efectos del transporte de voz sobre redes IP son la variación del retardo y la pérdida de secuencia en la entrega de paquetes. Para contrarrestarlos podemos usar buffers y números de secuencia.

La IETF ofrece dos protocolos de transporte para proporcionar números de secuencia, marcas de tiempo e identificación del tipo de carga útil, que no influyen en la calidad de la red de transporte IP. Estos dos protocolos trabajan sobre UDP (User Datagram Protocol) y son RTP y RTCP.

11.4.1 RTP (Real Time Protocol, RFC 1889)

El protocolo RTP se usa tanto en SIP como en H.323 y soporta la transferencia de tráfico de audio y voz en redes de conmutación de paquetes. Como protocolo de transporte permite al receptor detectar cualquier pérdida de paquetes y ofrece una marca de tiempo para que pueda compensar el retardo producido por la transmisión.

En la cabecera RTP encontramos información útil de cómo cada codec fragmenta los paquetes para reconstruir correctamente el flujo de datos (stream) enviado.

La seguridad no recae en RTP en forma de autenticación, pero podemos cifrar el contenido mediante algoritmos de bloque.

11.4.1.1 Funciones de RTP

- **Números de secuencia:** Ayudan a la identificación de paquetes perdidos.
- **Identificador de sobrecarga:** Informa del tipo de codec utilizado. En ocasiones es necesario cambiar el codec dinámicamente para ajustar el tráfico al ancho de banda disponible.
- **Indicador de Trama:** Separa dos de las tramas en las que se divide el stream de audio o video.
- **Identificador de origen:** Provee una manera de distinguir diferentes orígenes en sesiones multicast.
- **Sincronización entre medios:** Se utiliza para compensar los retardos de paquetes de diferentes streams. Con marcas de tiempo se ajustan los buffers de salida.

11.4.1.2 Formato de trama de RTP

V=2	P	X	CC	PayLoad Type	Número de Secuencia
Time Stamp					
Synchronization Source Identifier (SSRC)					
Contributing Source Identifier (CSRC)					
Dependiente del Perfil				Tamaño	
Datos					
0				32	

P=Padding, **X**=Extensiones tras CSRC, **CC**=Cuenta CSRC, **M**=Marcador

Payload type=Tipo de sobrecarga, **Número de secuencia**=Comienza con un N° aleatorio

Timestamp=Cuenta de *ticks* desde la emisión del primer paquete (1 *tick* = 1/8000)

SSRC=Origen del envío, a un mismo origen, mismo tiempo y N° de secuencia

11.4.2 RTCP (Real Time Control Protocol, RFC 1890)

El protocolo RTCP trabaja conjuntamente con el protocolo RTP en una sesión RTP para implementar nuevas funciones como:

- **QoS Feedback:** RTCP envía periódicamente información sobre la calidad de servicio.

- **Control de sesión:** Con el uso del paquete BYE se anuncia a los demás participantes que se finaliza la sesión.
- **Identificación:** En los paquetes se incluyen el nombre, dirección de correo y teléfono de todos los participantes de la sesión.
- **Sincronización entre medios:** RTCP provee la información necesaria para que se reproduzcan a la vez los streams transmitidos por separado de audio y vídeo.

El tráfico de paquetes RTCP aumentan con el número de usuarios, incrementando el ancho de banda ocupado. Para prevenir esta masificación se acota el ancho de banda para este protocolo (lo normal es acotarlo a un 5%).

11.4.2.1 Tipos de paquete RTCP

- **SR (Sender Report):** Información sobre transmisión y recepción.
- **RR (Receiver Report):** Información de recepción para los receptores.
- **SDES (Source Description):** Parámetros del origen (CNAME).
- **BYE:** Enviado por un participante cuando abandona la conferencia.
- **APP:** Funciones propietarias específicas de aplicación.

Como ejemplo, vemos el formato de la trama de dos tipos típicos como el SR y el SDES:

a. Sender Report

V=2	P	CC	PT=SR=200	Longitud
Sender (SSRC del autor del Report)				
NTP TimeStamp (byte más significativo)				
NTP TimeStamp (byte menos significativo)				
RTP TimeStamp				
Sender's Packet Count				
Sender's Octet Count				
SSRC 1 (SSRC de la primera fuente)				
Datos RR Adicionales				
SSRC N				
Datos RR Adicionales				

RC: Report Count, **PT:** Carga útil = 200 para SR, **Longitud del Report,**

SSRC: qué lo origina, **NTP timestamp:** Segundos desde el 1/1/1900

RTP timestamp: el mismo instante en *ticks* de RTP (equivalencia)

Paquetes y octetos enviados desde el inicio de la sesión (por SSRC)

Conjunto de RR: uno por cada fuente escuchada

b. Paquetes SDES

Como el anterior pero con las siguientes diferencias:

- **SC:** *Source count*
- **PT:** Carga útil = 202 para SDES

Cada origen tiene un campo SSCR (o CSRC) y una lista con información:

- **CNAME:** Identificador único del formato usuario@ host dirección IP.
- **NAME:** Nombre común del origen.
- E-MAIL, PHONE, LOCATION, entre otros.

11.5 Clasificación de los protocolos sobre VoIP

Si quisiéramos definir en forma teórica, independizándonos de los protocolos ya existentes, un modelo del procedimiento para establecer una comunicación de voz entre dos terminales sobre una red IP, lo primero que deberíamos hacer es definir los distintos tipos de negociaciones que deberían intercambiar las terminales para lograr la comunicación.

La primer idea que surge es la de informar al terminal llamado que deseo establecer la comunicación de voz. Luego el terminal llamado responderá de alguna forma, aceptando o rechazando dicha comunicación. A este tipo de intercambio de información se la suele llamar señalización de llamada (*call signalling*).

Por tratarse de una comunicación de voz sobre una red IP, la voz se transmite codificada en paquetes. Existen una gran variedad de codificadores y hoy en día los más utilizados son G.729, G.711, GSM, entre otros. Además en la mayoría de los casos la voz se transporta sobre segmentos UDP, lo que hace necesario la negociación de los puertos UDP donde el receptor espera recibir el audio. Debido a esto, es necesario intercambiar mensajes donde se negocien estas cuestiones y otras más específicas de cada protocolo. Para el intercambio de este tipo de información se definen los protocolos de control de señalización de llamada (*Call control signaling*).

Una vez establecida la comunicación, se debe enviar el audio codificado en paquetes IP. Las redes IP suelen tener variaciones de retardo altos respecto a las redes de telefonía tradicionales ya que no fueron diseñadas para el transporte de voz. Y además, por ser una red de datagramas, los paquetes de voz podrían llegar desordenados. Debido a estas características de la red IP, se necesita empaquetar la información de voz sobre algún protocolo que minimice o controle estos efectos. A estos protocolos se los denomina protocolos de transporte de "media" (*media transport protocols*), cuya función es la de enviar entre los terminales intervinientes en la comunicación estadísticas sobre jitter, paquetes enviados, paquetes recibidos, paquetes perdidos, etc. La RFC3550 define el protocolo RTP y RTCP que son hoy en día los más utilizados para el transporte y control de la "media".

A esta altura parecería que tenemos todos los elementos necesarios para poder establecer y controlar una comunicación de voz entre dos terminales. Esto es cierto y de hecho en algunas topologías de redes pequeñas con esto es suficiente. Cuando la red empieza a crecer y ya no son solo terminales los que se quieren comunicar sino que también *gateways* para interconectarse con la red de telefonía pública tradicional, se hace necesario centralizar cierto

tipo de información para que la red sea escalable. Para lograr esto se coloca un dispositivo de control que posee la inteligencia de la red, es decir, capacidades de ruteo, transcoding de señalización y localización de dispositivos entre otras funciones. A éste dispositivo se lo suele denominar *softswitch*. Como consecuencia se hace necesaria la comunicación entre gateways o terminales y el dispositivo de control, el *softswitch*. A este tipo de comunicación le llamaremos protocolos de registración y control. Cabe destacar que esta clasificación es un poco ambigua ya que a veces la definición de los protocolos de registración y control está embebida como parámetros dentro de los protocolos de señalización de llamadas.

Los protocolos existentes hoy en día clasificados según sus funciones son los siguientes:

11.5.1 Protocolos de señalización de llamada

Para simplificar la explicación vamos a utilizar un ejemplo de una llamada directa entre dos terminales (teléfonos IP o softphones).

Supongamos que un usuario quiere establecer una llamada, y para ello, digita una dirección IP (no suele ser lo más común pero ésta sería la forma más simple si no hay algún tipo de dispositivo de control que traduzca nombres o números en direcciones IP) que indica el teléfono destino al que se quiere llamar. El teléfono IP llamante envía un paquete al teléfono IP llamado indicándole que quiere establecer una comunicación. El teléfono llamado responde indicando que recibió la llamada y la está procesando. Cuando el llamado avisa por medio del timbre local (*ringing*) que ha recibido una llamada entrante, envía un mensaje al llamante para avisar. Una vez que se levanta el teléfono, se avisa al llamante con otro mensaje y a partir de ese momento la comunicación queda establecida. Todo este envío de mensajes se realiza a través de los *protocolos de señalización de llamada*. Estos son SIP y H.323 (H.225/Q.931) A manera de ejemplo se muestra el flujo de mensajes en la Figura 45.



Figura 45. Flujo de mensajes VoIP

Los protocolos de señalización en VoIP cumplen funciones similares a sus homólogos en la telefonía tradicional, es decir tareas de establecimiento de sesión, control del progreso de la llamada, entre otras. Se encuentran en la capa de sesión del modelo OSI.

Existen algunos protocolos de señalización, que han sido desarrollados por diferentes fabricantes u organismos como la ITU o el IETF, y que se encuentran soportados por Asterisk. Algunos son:

- H.323
- SIP
- MGCP
- IAX
- SCCP



Figura 46. Protocolos de señalización

11.5.1.1 Protocolo H.323

Es un estándar de la ITU-T (International Telecommunications Union). Esta recomendación fue diseñada originalmente para el tráfico de contenido multimedia en redes LAN, pero se extendió para cubrir la voz sobre IP. Por este motivo no se tiene en cuenta la calidad de servicio (QoS). La estandarización de H.323 permite que productos de diferentes fabricantes cooperen sin problemas.

H.323 define la señalización necesaria para llamadas y conferencias con los codecs más utilizados y utiliza para el transporte los protocolos RTP/RTCP aunque no se utilizan todas las capacidades de RTCP como los paquetes SDP.

11.5.1.1.1 Características de H.323

- Originalmente fue diseñado para el transporte de video conferencia.
- Su especificación es compleja.
- H.323 es un protocolo relativamente seguro, ya que utiliza RTP.
- Tiene dificultades con NAT, por ejemplo para recibir llamadas se necesita direccionar el puerto TCP 1720 al cliente, además de direccionar los puertos UDP para la media de RTP y los flujos de control de RTCP.
- Para más clientes detrás de un dispositivo NAT se necesita gatekeeper en modo proxy.

11.5.1.1.2 Componentes de la Arquitectura H.323

Los principales integrantes de una infraestructura H.323 típica como son los endpoints (terminales, gateways y Multipoint Control Units) y los gatekeepers.

a. Terminales

Son puntos finales de los clientes de la arquitectura y se encargan de las comunicaciones de dos vías en tiempo real. Todos los terminales H.323 tienen que soportar H.245 (para controlar el acceso a la utilización de los canales), Q.931 (se requiere para la señalización e inicio de las llamadas), Registration Admission Status (RAS, interactúa con el gatekeeper) y RTP. Como opción, los terminales pueden ser compatibles también con los protocolos de conferencia de datos T.120, codecs de vídeo y con MCU.

Los terminales H.323 se comunican con otros terminales H.323, con un gateway H.323 o con un MCU.

Como ejemplo de terminales típicos podemos señalar teléfonos hardware H.323,

Video teléfonos o softphones. Utilizaremos teléfonos convencionales conectados a una tarjeta de interfaz de voz con puertos FXS (Foreign Exchange Station) de un router Cisco 1761 como terminales H.323.

b. Gateways

Los gateways H.323 se ocupan de la comunicación en tiempo real entre dos terminales de la red IP o entre otros terminales ITU de una red de conmutación de paquetes o con otro Gateway H.323. Implementan la función de “traductores”, ya que favorecen la transición entre formatos de transmisión distintos y entre codecs de audio o vídeo según la capacidad del enlace. Son una parte muy importante de cualquier arquitectura de VoIP porque proveen la interfaz entre la Red

Telefónica Básica y la red pública de Internet. Así mismo, para las llamadas entre terminales de una misma LAN no son necesarios. Los gateways H.323 se usan en entornos en los que se necesita salida a la RTB o para comunicar diferentes redes con otros gateways usando los protocolos H.245 y Q.931.

c. Gatekeepers

Los más importantes componentes de un sistema H.323 son los gatekeepers porque ejercen de administradores de la arquitectura. Cada gatekeeper es el punto central de todas las llamadas de su zona (suma de un gatekeeper y todos los terminales que tenga suscritos). Las funciones más comunes de los gatekeepers son:

- **Traducción de direcciones:** Se traducen direcciones “alias” a direcciones de transporte usando una tabla de traducción actualizada mediante los mensajes de registro.
- **Control de admisión:** Pueden ofrecer o denegar acceso basándose en políticas de autorización, en las direcciones origen y destino o cualquier otro criterio.
- **Señalización de la llamada:** El gatekeeper elige si compartir la señalización de la llamada con los endpoints o procesarla íntegra por su cuenta. También puede dejar conectado el canal de la señalización de la llamada directamente entre los terminales origen y destino.
- **Autorización de llamada:** Haciendo uso de la señalización H.225, el gatekeeper es capaz de rechazar la llamada si el terminal origen o destino poseen algún tipo de restricción por período de tiempo o por destino.

- **Administración del ancho de banda:** Gracias al uso del protocolo H.225, es posible controlar el número de terminales H.323 que tienen acceso permitido a la red. El gatekeeper también puede denegar llamadas por falta de ancho de banda.
- **Administración de llamadas:** El gatekeeper mantiene una lista de las llamadas H.323 entrantes y salientes. Esto se usa para conocer el estado de cada terminal y recoger información para la función de administración del ancho de banda.

d. Multipoint Control Unit (MCU)

La unidad de control multipunto se encarga de hacer posible la participación de tres o más terminales y gateways en una conferencia. Es un componente opcional que consiste en un Multipoint Controller (MC) y Multipoint Processors (MP) opcionales. El MC o controlador multipunto determina las capacidades comunes de los terminales por la señalización del protocolo H.245 y controla a los MP o procesadores multipunto, que se encargan de la multiplexación y conmutación de los datos, vídeo y audio.

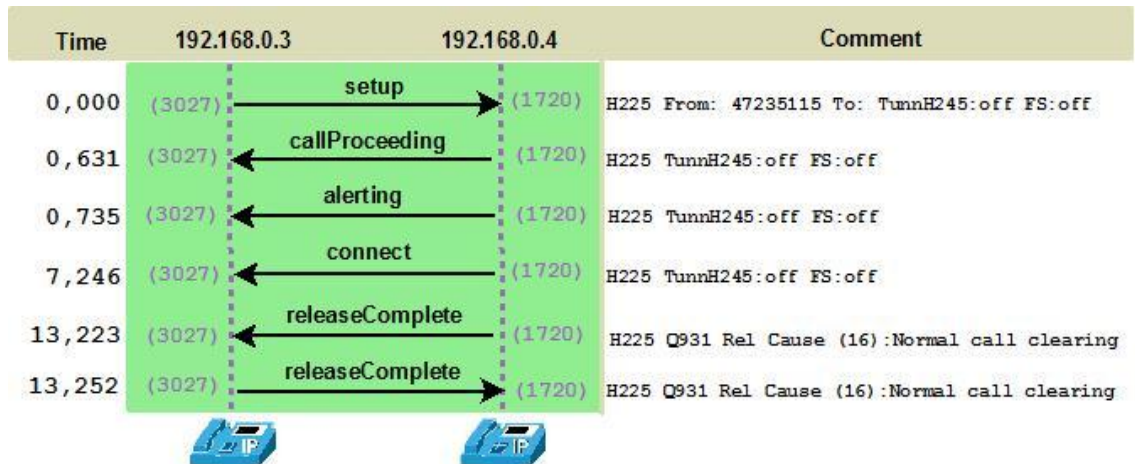


Figura 47. Intercambio de mensajes con MCU

11.5.1.2 Protocolo SIP (Session Initiation Protocol)

El Session Initiation Protocol (protocolo de inicio de sesión) es el estándar del IETF para la comunicación VoIP. Posee una arquitectura cliente/servidor parecida al conocido HTTP y es un protocolo de control de la capa de aplicación del modelo OSI que se usa para crear, mantener y finalizar sesiones entre dos o más usuarios. Además provee fiabilidad, no teniendo que recurrir al protocolo TCP y depende del protocolo SDP (Session Description Protocol, protocolo descriptor de sesión) que negocia entre los conferenciantes el grupo de codecs que se van a utilizar en la llamada. SIP soporta bien el redireccionamiento gracias a los proxies, que ubican y destinan el tráfico hacia los terminales.

11.5.1.2.1 Características de SIP

- Acrónimo de "Session Initiation Protocol".
- Este protocolo considera a cada conexión como un par y se encarga de negociar las capacidades entre ellos.
- Tiene una sintaxis simple, similar a HTTP o SMTP.

- Posee un sistema de autenticación de pregunta/respuesta.
- Tiene métodos para minimizar los efectos de DoS (Denial of Service o Denegación de Servicio), que consiste en saturar la red con solicitudes de invitación.
- Utiliza un mecanismo seguro de transporte mediante TLS.
- No tiene un adecuado direccionamiento de información para el funcionamiento con NAT.

11.5.1.2.2 Componentes de una Arquitectura SIP

Los elementos principales de un sistema SIP son:

- a. **Agentes de usuario:** Realizan las funciones de manejo de las llamadas entrantes o salientes. También filtrado de llamadas, localización del usuario y reintento de llamadas fallidas. En el proyecto usaremos Linphone agente de usuario software y las extensiones de la tarjeta FXO del router Cisco 1761 como agentes de usuario hardware.
- b. **Servidores:** Incluyen varias funciones, como son:
 - Proxy: actúa como servidor en un lado y como cliente en el otro y transmite los paquetes con o sin cambios.
 - Registrar: permite el registro y cancelación de agentes de usuario.
 - Redirigir: devuelve la dirección del siguiente servidor (siguiente salto).

Para comunicar los elementos de una arquitectura SIP se usan búsquedas DNS (*Domain Name System*) o direcciones IP configuradas y específicamente direcciones multicast para el servicio de registro del servidor.

11.5.1.2.3 Mensajes SIP

Los mensajes SIP se codifican con HTTP v1.1 y caracteres ISO 10646 y ocupan mayor ancho de banda que los mensajes H.323. Hay dos tipos de mensajes: peticiones y respuestas.

- a. **Peticiones:**
 - ACK: enviado por un cliente para confirmar que ha recibido la respuesta.
 - BYE: enviado para finalizar una llamada.
 - Cancel: aborta una solicitud enviada previamente.
 - Invite: se utiliza para iniciar una llamada.
 - Options: permite a un cliente conocer los métodos de un servidor.
 - Register: para que los agentes de usuario puedan registrar su localización actual.
- b. **Códigos de respuesta:**
 - 1xx: Información (180 Ringing)
 - 2xx: Éxito (200 OK)
 - 3xx: Redirección (301 Moved permanently)
 - 4xx: Error del Cliente (400 Bad Request, 406 Not Acceptable)
 - 5xx: Error del Servidor (502 Bad Gateway)
 - 6xx: Fallo General (600 Busy Everywhere)

11.5.1.2.4 Funciones de SIP

- a. **Resolución de direcciones:** Es una de las principales funciones de SIP y la efectúan tanto los agentes de usuario como los servidores aunque suele ser tarea del *proxy*. La resolución implica la mayoría de las ocasiones una traducción de una dirección en formato URI (*Universal Resource Identifier*, identificador universal de recurso) a una dirección IP. Normalmente se manda un mensaje DNS SRV, seguido de un ENUM *Lookup* (búsqueda ENUM) y de un *Location Server Lookup* (búsqueda de localización de servidor).
- b. **Establecimiento de la sesión:** Se realiza a través del proxy cuando está presente, si no, entre agentes de usuario directamente.
- c. **Negociación de contenidos:** La negociación se realiza con parte de las peticiones Invite y ACK con el protocolo SDP. Este protocolo también utiliza un sistema cliente/servidor, está diseñado para la arquitectura multimedia de Internet y está descrito en la RFC 2327. Con esta negociación se acuerda entre las dos partes los codecs que se van a usar, el tipo de contenido transportado, etc.
- d. **Modificación de la sesión:** Sólo se puede renegociar la sesión después de un primer establecimiento, no siendo indispensable cortar la transmisión en curso y con posibilidad de cambiar tanto el tipo de sesión como el codec usado o la dirección IP y el puerto.
- e. **Terminación y cancelación de la sesión:** Hay una pequeña diferencia entre la terminación y la cancelación que se basa en la petición utilizada, es decir, si se termina la sesión es que ésta ha llegado a establecerse y si se cancela es que la sesión todavía está en la etapa de establecimiento.
- f. **Control de llamada con REFER:** Gracias a esta funcionalidad de SIP podemos estar manteniendo una conversación con otro agente de usuario y antes de desconectarnos dejarlo llamando a un tercero con el mensaje REFER.
- g. **Sesiones con QoS:** Es una aproximación a la implementación de la calidad de servicio como funcionalidad aprovechando que SDP transporta información QoS y usando tres extensiones de SIP como son:
 - 183 Session Progress con SDP
 - Confiabilidad a 183
 - PreCondition MET COMET
- h. **Creación de servicios con SIP:** Una de las mayores ventajas de SIP, ya que se pueden implementar de una manera flexible y sencilla servicios más avanzados que la PSTN debido a la ingente cantidad de información que se transporta en el establecimiento de la llamada usando las redes IP. Además, estas nuevas funcionalidades pueden residir en todos los componentes de la arquitectura o sólo en algunos, garantizando la compatibilidad.

Las herramientas que hacen posible esta extensión de servicios son SIP CGI (RFC 3050), SIP Servlets o el lenguaje CPL (*Call Processing Language*, lenguaje de procesado de llamadas).

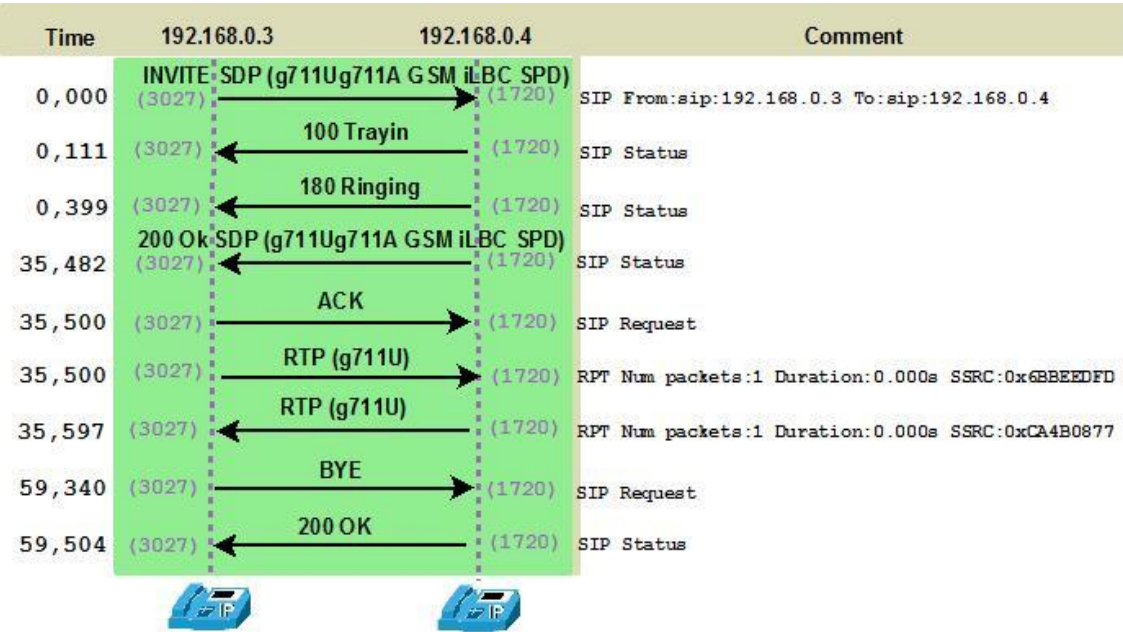


Figura 48. Intercambios de mensajes con el protocolo SIP

11.5.1.3 MGCP-MEGACO (Media Gateway Control Protocol)

Media Gateway Control Protocol (MGCP) es otro estándar de señalización paraVoIP desarrollado por la IETF. MGCP está basado en un modelo maestro/esclavo donde el Call Agent (servidor) es el encargado de controlar al gateway. De esta forma se consigue separar la señalización de la transmisión de la información, simplificando la integración con el protocolo SS7.

Esta importante ventaja propició la colaboración conjunta entre el IETF y la ITU para el desarrollo de una nueva especificación basada en MGCP que fuera complementaria a SIP y H.323. El resultado fue MEGACO aunque la ITU se refiere a este protocolo como H.248. En definitiva, SIP y H.323 se utiliza para la señalización en los extremos, mientras que MEGACO es óptimo para los grandes operadores de telefonía.

11.5.1.4 IAX (Inter-Asterisk eXchange protocol)

Inter-Asterisk eXchange protocol (IAX) fue desarrollado por Digium para la comunicación entre centralitas basadas en Asterisk aunque actualmente se ha implementado clientes que también soportan este protocolo.

El principal objetivo de IAX es minimizar el ancho de banda utilizado en la transmisión de voz y vídeo a través de la red IP y proveer un soporte nativo para ser transparente a los NATs (Network Address Translation). La estructura básica de IAX se fundamenta en la multiplexación de la señalización y del flujo de datos sobre un simple puerto UDP, generalmente el 4569.

El protocolo original ha quedado obsoleto en favor de su segunda versión conocida como IAX2. Se caracteriza por ser robusto y simple en comparación con otros protocolos. Permite manejar una gran cantidad de códecs y transportar cualquier tipo de datos.

12. PROTOCOLOS DE ENRUTAMIENTO CON IPv4

Nota: en el documento titulado **“Prácticas de laboratorios para la asignatura de Redes de Ordenadores II”**, se podrá encontrar con más amplitud el desarrollo teórico de los protocolos que a continuación estaremos abordando

12.1 RIP (Media Gateway Control Protocol)

El Protocolo de información de enrutamiento (RIP) fue descrito originalmente en el RFC 1058. Sus características principales son las siguientes:

- Es un protocolo de enrutamiento por vector-distancia.
- Utiliza el número de saltos como métrica para la selección de rutas.
- Si el número de saltos es superior a 15, el paquete es desechado.
- Por defecto, se envía un broadcast de las actualizaciones de enrutamiento cada 30 segundos

RIP ha evolucionado a lo largo de los años desde el Protocolo de enrutamiento con definición de clases, RIP Versión 1 (RIP v1), hasta el Protocolo de enrutamiento sin clase, RIP Versión 2 (RIP v2). Las mejoras en RIP v2 incluyen:

- Capacidad para transportar mayor información relativa al enrutamiento de paquetes.
- Mecanismo de autenticación para la seguridad de origen al hacer actualizaciones de las tablas.
- Soporta enmascaramiento de subredes de longitud variable (VLSM).

RIP evita que los bucles de enrutamiento se prolonguen en forma indefinida, mediante la fijación de un límite en el número de saltos permitido en una ruta, desde su origen hasta su destino. El número máximo de saltos permitido en una ruta es de 15. Cuando un Router recibe una actualización de enrutamiento que contiene una entrada nueva o cambiada, el valor de la métrica aumenta en 1, para incluir el salto correspondiente a sí mismo. Si este aumento hace que la métrica supere la cifra de 15, se considera que es infinita y la red de destino se considera fuera de alcance. RIP incluye diversas características las cuales están presentes en otros protocolos de enrutamiento. Por ejemplo, RIP implementa los mecanismos de espera y horizonte dividido para prevenir la propagación de información de enrutamiento errónea.

12.2 OSPF (Media Gateway Control Protocol)

El protocolo público conocido como "Primero la ruta más corta" (OSPF) es un protocolo de enrutamiento de estado del enlace no patentado. Las características clave del OSPF son las siguientes:

- Es un protocolo de enrutamiento de estado del enlace.
- Es un protocolo de enrutamiento público (open standard), y se describe en el RFC 2328.
- Usa el algoritmo SPF para calcular el costo más bajo hasta un destino.
- Las actualizaciones de enrutamiento producen un gran volumen de tráfico al ocurrir cambios en la

OSPF es uno de los protocolos del estado de enlace más importantes. OSPF se basa en las normas de código abierto, lo que significa que muchos fabricantes lo pueden desarrollar y mejorar. Es un protocolo complejo cuya implementación en redes más amplias representa un verdadero desafío. Los principios básicos de OSPF se tratan en este módulo.

La configuración de OSPF en un Router Cisco es parecido a la configuración de otros protocolos de enrutamiento. De igual manera, es necesario habilitar OSPF en un Router e identificar las redes que serán publicadas por OSPF. OSPF cuenta con varias funciones y procedimientos de configuración únicos. Estas funciones aumentan las capacidades de OSPF como protocolo de enrutamiento, pero también complican su configuración.

En grandes redes, OSPF se puede configurar para abarcar varias áreas y distintos tipos de área. La capacidad para diseñar e implementar OSPF en las grandes redes comienza con la capacidad para configurar OSPF en una sola área

12.3 IS-IS

IS-IS es un protocolo de IGP desarrollado por DEC, Y suscrito por la ISO en los años 80 como protocolo de routing para OSI. El desarrollo de IS-IS fue motivado por la necesidad de:

- Protocolo no propietario
- Esquema de direccionamiento grande.
- Esquema de direccionamiento jerárquico.
- Protocolo eficiente permitiendo una convergencia rápida y precisa y poca sobre carga en la red
- IS-IS es utilizado por el gobierno de los EEUU y actualmente está emergiendo.

El nuevo interés se basa en que IS-IS es un estándar que proporciona independencia del protocolo, capacidad de escalado y capacidad de definir Routing basado en ToS.

12.4 EIGRP (Media Gateway Control Protocol)

EIGRP es un protocolo de enrutamiento propietario de Cisco basado en IGRP. Admite CIDR y VLSM, lo que permite que los diseñadores de red maximicen el espacio de direccionamiento. En comparación con IGRP, que es un protocolo de enrutamiento con clase, EIGRP ofrece tiempos de convergencia más rápidos, mejor escalabilidad y gestión superior de los bucles de enrutamiento. Además, EIGRP puede reemplazar al Protocolo de Mantenimiento de Tablas de Enrutamiento (RTMP) AppleTalk y Novell RIP. EIGRP funciona en las redes IPX y AppleTalk con potente eficiencia.

Con frecuencia, se describe EIGRP como un protocolo de enrutamiento híbrido que ofrece lo mejor de los algoritmos de vector-distancia y del estado de enlace.

EIGRP es un protocolo de enrutamiento avanzado que se basa en las características normalmente asociadas con los protocolos del estado de enlace. Algunas de las mejores funciones de OSPF, como las actualizaciones parciales y la detección de vecinos, se usan de forma similar con EIGRP. Sin embargo, EIGRP es más fácil de configurar que OSPF.

EIGRP es una opción ideal para las grandes redes multiprotocolo construidas principalmente con Routers Cisco.

Cuando un componente de red no funciona correctamente, puede afectar toda la red. En todo caso, los administradores de red deben identificar y diagnosticar los problemas rápidamente cuando se produzcan. A continuación se presentan algunas de las razones por las que surgen problemas en la red.

- Se introducen comandos de forma incorrecta
- Se construyen o colocan las listas de acceso de forma incorrecta
- Los Routers, Switch u otros dispositivos de red están configurados de forma incorrecta
- Las conexiones físicas son de mala calidad

Un administrador de red debe realizar el diagnóstico de fallas de forma metódica, mediante un modelo general de resolución de problemas. A menudo es útil verificar si hay problemas de la capa física en primer lugar y luego ir subiendo por las capas de forma organizada. Aunque este módulo se concentra en la forma de diagnosticar las fallas de los protocolos de Capa 3, es importante diagnosticar y eliminar los problemas existentes en las capas inferiores.

12.4.1 Comparación de IGRP e EIGRP

Cisco lanzó EIGRP en 1994 como una versión escalable y mejorada de su protocolo propietario de enrutamiento por vector-distancia, IGRP. En esta sección se explican las diferencias y similitudes existentes entre EIGRP e IGRP. La información de distancia y la tecnología de vector-distancia que se usan en IGRP también se utilizan en EIGRP.

EIGRP mejora las propiedades de convergencia y opera con mayor eficiencia que IGRP. Esto permite que una red tenga una arquitectura mejorada y pueda mantener las inversiones actuales en IGRP.

Las comparaciones entre EIGRP e IGRP se pueden dividir en las siguientes categorías principales:

- Modo de compatibilidad
- Cálculo de métrica
- Número de saltos
- Redistribución automática de protocolos
- Etiquetado de rutas

IGRP y EIGRP son compatibles entre sí. Esta compatibilidad ofrece una interoperabilidad transparente con los Routers IGRP. Esto es importante, dado que los usuarios pueden aprovechar los beneficios de ambos protocolos. EIGRP ofrece compatibilidad multiprotocolo, mientras que IGRP no lo hace.

EIGRP e IGRP usan cálculos de métrica diferentes. EIGRP multiplica la métrica de IGRP por un factor de 256. Esto ocurre porque EIGRP usa una métrica que tiene 32 bits de largo, e IGRP usa una métrica de 24 bits. La información EIGRP puede multiplicarse o dividirse por 256 para un intercambio fácil con IGRP.

IGRP tiene un número de saltos máximo de 255. El límite máximo para el número de saltos en EIGRP es 224. Esto es más que suficiente para admitir los internetworks grandes y correctamente diseñadas.

Se requiere una configuración avanzada para permitir que protocolos de enrutamiento diferentes como OSPF y RIP compartan información. La redistribución, o la capacidad para compartir rutas, es automática entre IGRP e EIGRP, siempre y cuando ambos procesos usen el mismo número AS. En la Figura siguiente, RTB distribuye de forma automática las rutas aprendidas de EIGRP al AS de IGRP, y viceversa. EIGRP rotula como externas las rutas aprendidas de IGRP o cualquier otra fuente externa porque no se originan en los Routers EIGRP. IGRP no puede diferenciar entre rutas internas y externas.

Observe que en el resultado del comando `show ip route` de los Routers de la Figura 49, las rutas de EIGRP se rotulan con una D, mientras que las rutas externas se rotulan con EX. RTA identifica la diferencia entre la red 172.16.0.0, que se aprendió mediante EIGRP, y la red 192.168.1.0 que se redistribuyó desde IGRP. En la tabla RTC, el protocolo IGRP no indica este tipo de diferencia. RTC, que usa solamente IGRP, sólo ve las rutas IGRP, a pesar de que tanto 10.1.1.0 como 172.16.0.0 se redistribuyeron desde EIGRP.

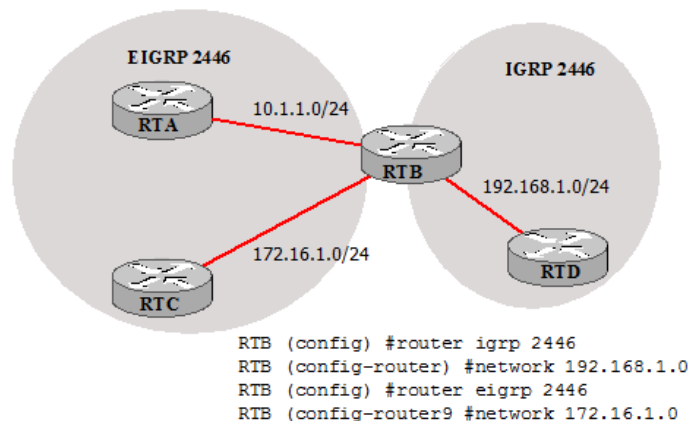


Figura 49. Comparación de EIGRP con IGRP

EIGRP Y IGRP redistribuyen automáticamente las rutas entre sistemas autónomos con el mismo número.

```

RTA# show ip route
<output omitted>
C      10.1.1.0 is directly connected serial0
D      172.16.1.0 [90/2681856] via 10.1.1.1, serial0
D EX   192.168.1.0 [170/2681856] via 10.1.1.1,
00:00:04, serial0

```

```

RTB# show ip route
<output omitted>
C      192.168.1.0 is directly connected serial0
I      10.0.0.0 [100/10476] via 192.168.1.1,
00:00:04, serial0
I      172.16.1.0 [100/10476] via 192.168.1.1,
00:00:04, serial0

```

12.4.2 Conceptos y terminologías de EIGRP

Los Routers EIGRP mantienen información de ruta y topología a disposición en la RAM, para que puedan reaccionar rápidamente ante los cambios. Al igual que OSPF, EIGRP guarda esta información en varias tablas y bases de datos.

EIGRP guarda las rutas que se aprenden de maneras específicas. Las rutas reciben un estado específico y se pueden rotular para proporcionar información adicional de utilidad.

EIGRP mantiene las siguientes tres tablas:

- Tabla de vecinos
- Tabla de topología
- Tabla de enrutamiento

La tabla de vecinos es la más importante de EIGRP. Cada Router EIGRP mantiene una tabla de vecinos que enumera a los Routers adyacentes. Esta tabla puede compararse con la base de datos de adyacencia utilizada por OSPF. Existe una tabla de vecinos por cada protocolo que admite EIGRP, similar a como se muestra en el Cuadro 28.

```
Router #show ip route eigrp neighbors
```

IP-EIGRP vecinos del proceso 100							
M	Dirección	Interfaces	Hold uptime (sec)	SRTT (ms)	RTO	Q CNT	SEQ NUM
2	200.10.10.10	Se1	13 00:19:09	26	200	0	10
1	200.10.10.5	Se0	12 03:31:36	50	300	0	39
0	199.55.32.10	Et0	11 03:31:40	10	200	0	40

Cuadro 28. Vecinos del proceso IP-EIGRP

Al conocer nuevos vecinos, se registran la dirección y la interfaz del vecino. Esta información se guarda en la estructura de datos del vecino. Cuando un vecino envía un paquete hello, publica un tiempo de espera. El tiempo de espera es la cantidad de tiempo durante el cual un Router considera que un vecino se puede alcanzar y que funciona. Si un paquete de salutación (hello) no se recibe dentro del tiempo de espera, entonces vence el tiempo de espera. Cuando vence el tiempo de espera, se informa al Algoritmo de Actualización Difusa (DUAL), que es el algoritmo de vector-distancia de EIGRP, acerca del cambio en la topología para que recalcule la nueva topología.

La tabla de topología se compone de todas las tablas de enrutamiento EIGRP en el sistema autónomo. DUAL toma la información proporcionada en la tabla de vecinos y la tabla de topología y calcula las rutas de menor costo hacia cada destino. EIGRP rastrea esta información para que los Routers EIGRP puedan identificar y conmutar a rutas alternativas rápidamente. La información que el Router recibe de DUAL se utiliza para determinar la ruta del sucesor, que es el término utilizado para identificar la ruta principal o la mejor. Esta información también se introduce a la tabla de

topología. Los Routers EIGRP mantienen una tabla de topología por cada protocolo configurado de red. La tabla de enrutamiento mantiene las rutas que se aprenden de forma dinámica.

12.4.3 Campos que conforman la tabla de enrutamiento

- **Distancia factible (FD):** Ésta es la métrica calculada más baja hacia cada destino. Por ejemplo, la distancia factible a 32.0.0.0 es 2195456.
- **Origen de la ruta:** Número de identificación del Router que publicó esa ruta en primer lugar. Este campo se llena sólo para las rutas que se aprenden de una fuente externa a la red EIGRP. El rotulado de rutas puede resultar particularmente útil con el enrutamiento basado en políticas. Por ejemplo, el origen de la ruta a 32.0.0.0 es 200.10.10.10 a 200.10.10.10.
- **Distancia informada (RD):** La distancia informada (RD) de la ruta es la distancia informada por un vecino adyacente hacia un destino específico. Por ejemplo, la distancia informada a 32.0.0.0 es 2195456 tal como lo indica (90/2195456).
- **Información de interfaz:** La interfaz a través de la cual se puede alcanzar el destino.
- **Estado de ruta:** El estado de una ruta. Una ruta se puede identificar como pasiva, lo que significa que la ruta es estable y está lista para usar, o activa, lo que significa que la ruta se encuentra en el proceso de recálculo por parte de DUAL.

La tabla de enrutamiento EIGRP contiene las mejores rutas hacia un destino. Esta información se recupera de la tabla de topología. Los Routers EIGRP mantienen una tabla de enrutamiento por cada protocolo de red. Un sucesor es una ruta seleccionada como ruta principal para alcanzar un destino. DUAL identifica esta ruta en base a la información que contienen las tablas de vecinos y de topología y la coloca en la tabla de enrutamiento. Puede haber hasta cuatro rutas de sucesor para cada destino en particular. Éstas pueden ser de costo igual o desigual y se identifican como las mejores rutas sin bucles hacia un destino determinado. Un sucesor factible (FS) es una ruta de respaldo. Estas rutas se identifican al mismo tiempo que los sucesores, pero sólo se mantienen en la tabla de topología. Los múltiples sucesores factibles para un destino se pueden mantener en la tabla de topología, aunque no es obligatorio.

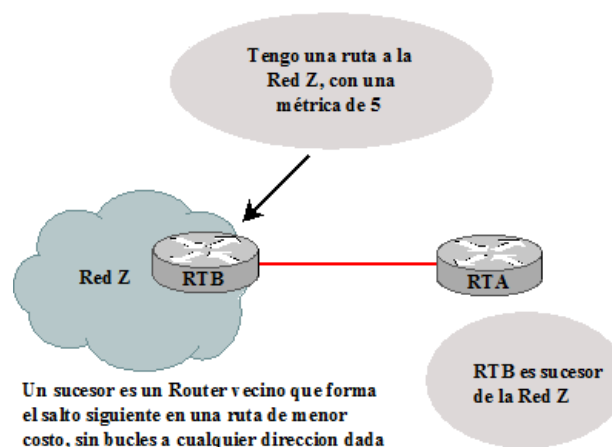


Figura 50. Identificación de sucesor

Un Router visualiza los sucesores factibles como vecinos corriente abajo, o más cerca del destino que él. El costo del sucesor factible se calcula a base del costo publicado del Router vecino hacia el destino. Si una ruta del sucesor colapsa, el Router busca un sucesor factible identificado. Esta ruta se promoverá al estado de sucesor. Un sucesor factible debe tener un costo publicado menor que el costo del sucesor actual hacia el destino. Si es imposible identificar un sucesor factible en base a la información actual, el Router coloca un estado Activo en una ruta y envía paquetes de consulta a todos los vecinos para recalcular la topología actual. El Router puede identificar cualquier nuevo sucesor o sucesor factible a partir de los nuevos datos recibidos de los paquetes de respuesta que responden a los pedidos de consulta. Entonces, el Router establecerá el estado de la ruta en Pasivo.

Es posible registrar información adicional acerca de cada ruta en la tabla de topología. EIGRP clasifica a las rutas como internas o externas. EIGRP agrega un rótulo de ruta a cada ruta para identificar esta clasificación. Las rutas internas se originan dentro del AS EIGRP se muestran en la Figura 51.

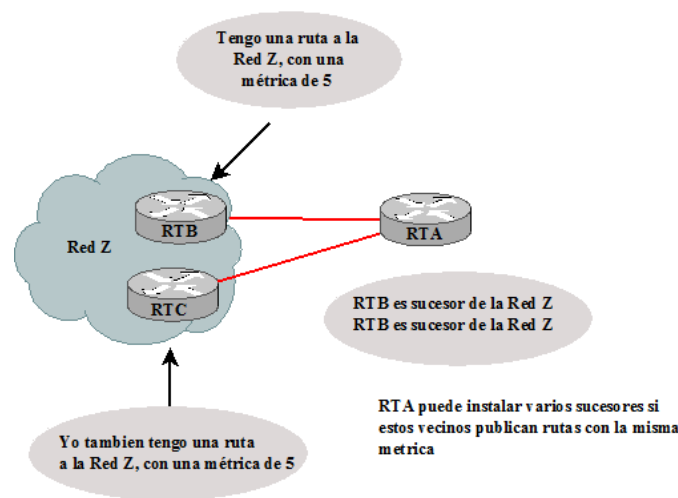


Figura 51. Instalación de varios sucesores

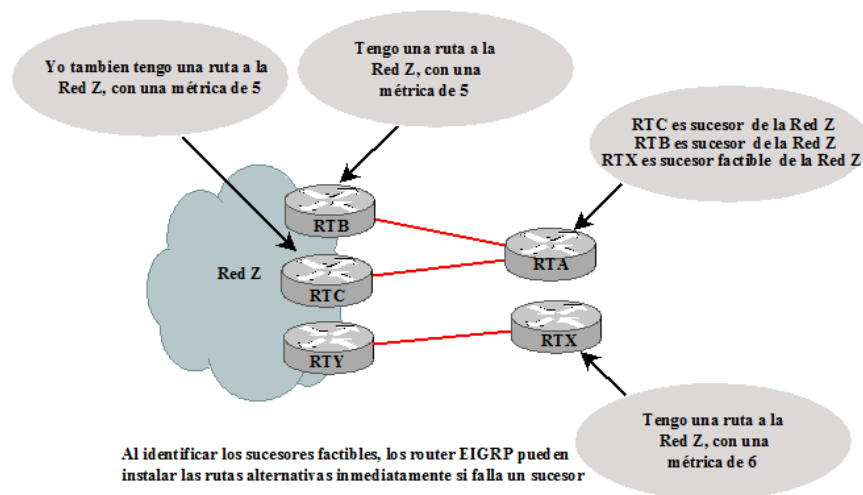


Figura 52. Identificación de sucesor factible

Las rutas externas se originan fuera del AS EIGRP. Las rutas aprendidas o redistribuidas desde otros protocolos de enrutamiento como RIP, OSPF e IGRP son externas. Las rutas estáticas que se originan fuera del AS EIGRP son externas. El rótulo puede establecerse en un número entre 0-255 para adaptar el rótulo

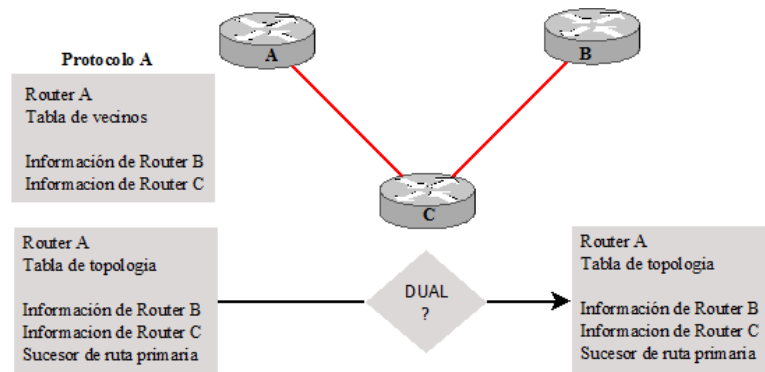


Figura 53. Tabla de vecinos y topología

12.4.4 Características de diseño de EIGRP

EIGRP opera de una manera bastante diferente de IGRP. EIGRP es un protocolo de enrutamiento por vector-distancia avanzado, pero también actúa como protocolo del estado de enlace en la manera en que actualiza a los vecinos y mantiene la información de enrutamiento. A continuación se presentan algunas de las ventajas de EIGRP sobre los protocolos de vector-distancia simples:

- Convergencia rápida
- Uso eficiente del ancho de banda
- Compatibilidad con VLSM y CIDR
- Compatibilidad con capas de varias redes

➤ Independencia de los protocolos enrutados

Los módulos dependientes de protocolo (PDM) protegen a EIGRP de las revisiones prolongadas. Es posible que los protocolos enrutados en evolución, como IP, requieran un nuevo módulo de protocolo, pero no necesariamente una reelaboración del propio EIGRP. Los Routers EIGRP convergen rápidamente porque se basan en DUAL. DUAL garantiza una operación sin bucles durante todo el cálculo de rutas, lo que permite la sincronización simultánea de todos los Routers involucrados en cambio de topología. EIGRP envía actualizaciones parciales y limitadas, y hace un uso eficiente del ancho de banda. EIGRP usa un ancho de banda mínimo cuando la red es estable. Los Routers EIGRP no envían las tablas en su totalidad, sino que envían actualizaciones parciales e incrementales. Esto es parecido a la operación de OSPF, salvo que los Routers EIGRP envían estas actualizaciones parciales sólo a los Routers que necesitan la información, no a todos los Routers del área. Por este motivo, se denominan actualizaciones limitadas. En vez de enviar actualizaciones de enrutamiento temporizadas, los Routers EIGRP usan pequeños paquetes hello para mantener la comunicación entre sí. Aunque se intercambian con regularidad, los paquetes hello no usan una cantidad significativa de ancho de banda.

EIGRP admite IP, IPX y AppleTalk mediante los PDM. EIGRP puede redistribuir información de IPX, RIP y SAP para mejorar el desempeño general. En efecto, EIGRP puede reemplazar estos dos protocolos. Los Routers EIGRP reciben actualizaciones de enrutamiento y de servicio y actualizan otros Routers sólo cuando se producen cambios en las tablas de enrutamiento o de SAP. En las redes EIGRP, las actualizaciones de enrutamiento se realizan por medio de actualizaciones parciales. EIGRP también puede reemplazar el RTMP de AppleTalk. Como protocolo de enrutamiento por vector-distancia, RTMP se basa en intercambios periódicos y completos de información de enrutamiento. Para reducir el gasto, EIGRP usa actualizaciones desencadenadas por eventos para redistribuir la información de enrutamiento AppleTalk. EIGRP también usa una métrica compuesta configurable para determinar la mejor ruta a una red AppleTalk. RTMP usa el número de saltos, lo que puede dar como resultado un enrutamiento por debajo del óptimo. Los clientes AppleTalk esperan información RTMP desde los Routers locales, de manera que EIGRP para AppleTalk sólo debe ejecutarse en una red sin clientes, como un enlace WAN.

12.4.5 Tecnologías EIGRP

Los Routers de vector-distancia simples no establecen ninguna relación con sus vecinos. Los Routers RIP e IGRP simplemente envían las actualizaciones en broadcast o multicast por las interfaces configuradas. En cambio, los Routers EIGRP establecen relaciones activamente con los vecinos, tal como lo hacen los Routers OSPF.

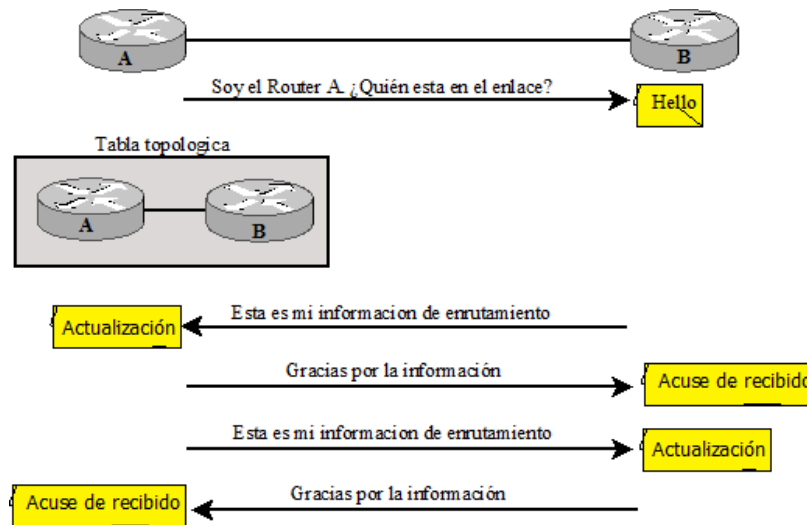


Figura 54. Establecimiento de relación en Routers vecinos

Tal y como se muestra en la Figura 54, los Routers EIGRP establecen adyacencias. Los Routers EIGRP lo logran mediante paquetes hello pequeños. Los hellos se envían por defecto cada cinco segundos. Un Router EIGRP supone que, siempre y cuando reciba paquetes hello de los vecinos conocidos, estos vecinos y sus rutas seguirán siendo viables o pasivas. Lo siguiente puede ocurrir cuando los Routers EIGRP forman adyacencias: Aprender de forma dinámica las nuevas rutas que se unen a la red.

- Identificar los Routers que llegan a ser inalcanzables o inoperables
- Redetectar los Routers que habían estado inalcanzables anteriormente

El Protocolo de Transporte Confiable (RTP) es un protocolo de capa de transporte que garantiza la entrega ordenada de paquetes EIGRP a todos los vecinos. En una red IP, los hosts usan TCP para secuenciar los paquetes y asegurarse de que se entreguen de manera oportuna. Sin embargo, EIGRP es independiente de los protocolos. Esto significa que no se basa en TCP/IP para intercambiar información de enrutamiento de la forma en que lo hacen RIP, IGRP y OSPF. Para mantenerse independiente de IP, EIGRP usa RTP como su protocolo de capa de transporte propietario para garantizar la entrega de información de enrutamiento. EIGRP puede hacer una llamada a RTP para que proporcione un servicio confiable o no confiable, según lo requiera la situación. Por ejemplo, los paquetes hello no requieren el gasto de la entrega confiable porque se envían con frecuencia y se deben mantener pequeños. La entrega confiable de otra información de enrutamiento puede realmente acelerar la convergencia porque entonces los Routers EIGRP no tienen que esperar a que un temporizador expire antes de retransmitir.

Con RTP, EIGRP puede realizar envíos en multicast y en unicast a diferentes pares de forma simultánea. Esto maximiza la eficiencia.

El núcleo de EIGRP es DUAL, que es el motor de cálculo de rutas de EIGRP. El nombre completo de esta tecnología es máquina de estado finito DUAL (FSM). Una FSM es una máquina de algoritmos, no un dispositivo mecánico con

piezas que se mueven. Las FSM definen un conjunto de los posibles estados de algo, los acontecimientos que provocan esos estados y los eventos que resultan de estos estados. Los diseñadores usan las FSM para describir de qué manera un dispositivo, programa de computador o algoritmo de enrutamiento reaccionará ante un conjunto de eventos de entrada. La FSM DUAL contiene toda la lógica que se utiliza para calcular y comparar rutas en una red EIGRP.

DUAL rastrea todas las rutas publicadas por los vecinos. Se comparan mediante la métrica compuesta de cada ruta. DUAL también garantiza que cada ruta esté libre de bucles. DUAL inserta las rutas de menor costo en la tabla de enrutamiento. Estas rutas principales se denominan rutas del sucesor. Una copia de las rutas del sucesor también se coloca en la tabla de enrutamiento.

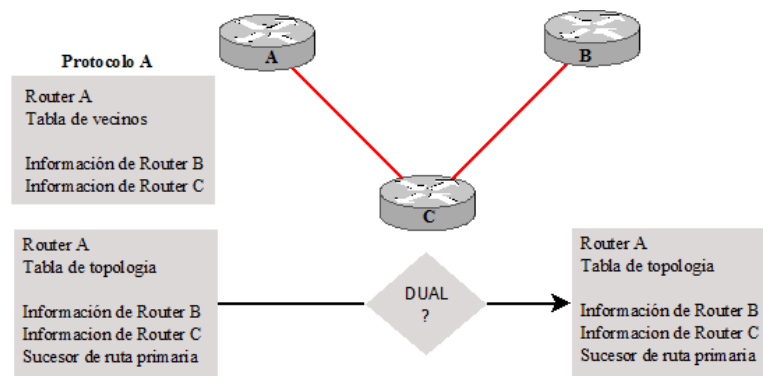


Figura 55. Tabla de vecinos y topología

EIGRP mantiene información importante de ruta y topología a disposición en una tabla de vecinos y una tabla de topología. Estas tablas proporcionan información detallada de las rutas a DUAL en caso de problemas en la red. DUAL usa la información de estas tablas para seleccionar rápidamente las rutas alternativas. Si un enlace se desactiva, DUAL busca una ruta alternativa, o sucesor factible, en la tabla de topología.

Una de las mejores características de EIGRP es su diseño modular. Se ha demostrado que los diseños modulares o en capas son los más escalables y adaptables. EIGRP logra la compatibilidad con los protocolos enrutados, como IP, IPX y AppleTalk, mediante los PDM. En teoría, EIGRP puede agregar PDM para adaptarse fácilmente a los protocolos enrutados nuevos o revisados como IPv6.

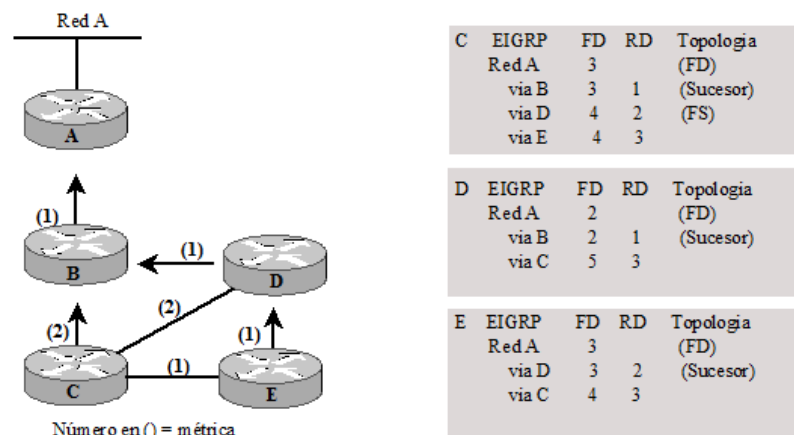


Figura 56. Construcción de tabla topológica

Leyenda	
EIGRP	Tipo de protocolo
FD	Distancia factible
RD	La distancia informada como la publico el Router vecino
Sucesor	Ruta primaria al destino
FS	Sucesor factible - Ruta de respaldo al destino

Cada PDM es responsable de todas las funciones relacionadas con su protocolo enrutado específico. El módulo IP-EIGRP es responsable de las siguientes funciones:

- Enviar y recibir paquetes EIGRP que contengan datos IP
- Avisar a DUAL una vez que se recibe la nueva información de enrutamiento IP
- Mantener de los resultados de las decisiones de enrutamiento DUAL en la tabla de enrutamiento IP
- Redistribuir la información de enrutamiento que se aprendió de otros protocolos de enrutamiento capacitados para IP

12.4.6 Estructura de datos EIGRP

Al igual que OSPF, EIGRP depende de diferentes tipos de paquetes para mantener sus tablas y establecer relaciones con los Routers vecinos.

A continuación se presentan los cinco tipos de paquetes EIGRP:

- Hello
- Acuse de recibo
- Actualización
- Consulta

➤ Respuesta

EIGRP depende de los paquetes hello para detectar, verificar y volver a detectar los Routers vecinos. La segunda detección se produce si los Routers EIGRP no intercambian hellos durante un intervalo de tiempo de espera pero después vuelven a establecer la comunicación

Los Routers EIGRP envían hellos con un intervalo fijo pero configurable que se denomina el intervalo hello. El intervalo hello por defecto depende del ancho de banda de la interfaz. En las redes IP, los Routers EIGRP envían hellos a la dirección IP multicast 224.0.0.10.

Ancho de banda	Enlace de ejemplo	Intervalo de Hello por defecto	Tiempo de espera Por defecto
1.544 Mbps o menos	Frame Relay Multipunto	60 segundos	180 segundos
Mayor que 1.544 Mbps	T1 Ethernet	5 segundos	15 segundos

Los Routers EIGRP almacenan la información sobre los vecinos en la tabla de vecinos. La tabla de vecinos incluye el campo de Número de Secuencia (Seq No) para registrar el número del último paquete EIGRP recibido que fue enviado por cada vecino. La tabla de vecinos también incluye un campo de Tiempo de Espera que registra el momento en que se recibió el último paquete. Los paquetes deben recibirse dentro del período correspondiente al intervalo de Tiempo de Espera para mantenerse en el estado Pasivo. El estado Pasivo significa un estado alcanzable y operacional.

Si EIGRP no recibe un paquete de un vecino dentro del tiempo de espera, EIGRP supone que el vecino no está disponible. En ese momento, interviene DUAL para reevaluar la tabla de enrutamiento. Por defecto, el tiempo de espera es equivalente al triple del intervalo hello, pero un administrador puede configurar ambos temporizadores según lo desee.

OSPF requiere que los Routers vecinos tengan los mismos intervalos hello e intervalos muertos para comunicarse. EIGRP no posee este tipo de restricción. Los Routers vecinos conocen el valor de cada uno de los temporizadores respectivos de los demás mediante el intercambio de paquetes hello. Entonces, usan la información para forjar una relación estable aunque los temporizadores no sean iguales. Los paquetes hello siempre se envían de forma no confiable. Esto significa que no se transmite un acuse de recibo. Los Routers EIGRP usan paquetes de acuse de recibo para indicar la recepción de cualquier paquete EIGRP durante un intercambio confiable. RTP proporciona comunicación confiable entre hosts EIGRP. El receptor debe enviar acuse de recibo de un mensaje recibido para que sea confiable. Los paquetes de acuse de recibo, que son paquetes hello sin datos, se usan con este fin. Al contrario de los hellos multicast, los paquetes de acuse de recibo se envían en unicast. Los acuses de recibo pueden adjuntarse a otros tipos de paquetes EIGRP, como los paquetes de respuesta.

Los paquetes de actualización se utilizan cuando un Router detecta un nuevo vecino. Los Routers EIGRP envían paquetes de actualización en unicast a ese nuevo vecino para que pueda aumentar su tabla de topología. Es posible que se necesite más de un paquete de actualización para transmitir toda la información de topología al vecino recientemente detectado.

Los paquetes de actualización también se utilizan cuando un Router detecta un cambio en la topología. En este caso, el Router EIGRP envía un paquete de actualización en multicast a todos los vecinos, avisándolos del cambio. Todos los paquetes de actualización se envían de forma confiable.

Un Router EIGRP usa paquetes de consulta siempre que necesite información específica de uno o de todos sus vecinos. Se usa un paquete de respuesta para contestar a una consulta.

Si un Router EIGRP pierde su sucesor y no puede encontrar un sucesor factible para una ruta, DUAL coloca la ruta en el estado Activo. Entonces se envía una consulta en multicast a todos los vecinos con el fin de ubicar un sucesor para la red destino. Los vecinos deben enviar respuestas que suministren información sobre sucesores o indiquen que no hay información disponible. Las consultas se pueden enviar en multicast o en unicast, mientras que las respuestas siempre se envían en unicast. Ambos tipos de paquetes se envían de forma confiable.

12.4.7 Algoritmo EIGRP

Esta sección describe el algoritmo DUAL, al que se debe la convergencia excepcionalmente rápida de EIGRP. Para comprender mejor la convergencia con DUAL, vea el ejemplo de la Figura 57. Cada Router ha construido una tabla de topología que contiene información acerca de la manera de enrutar al destino Red A.

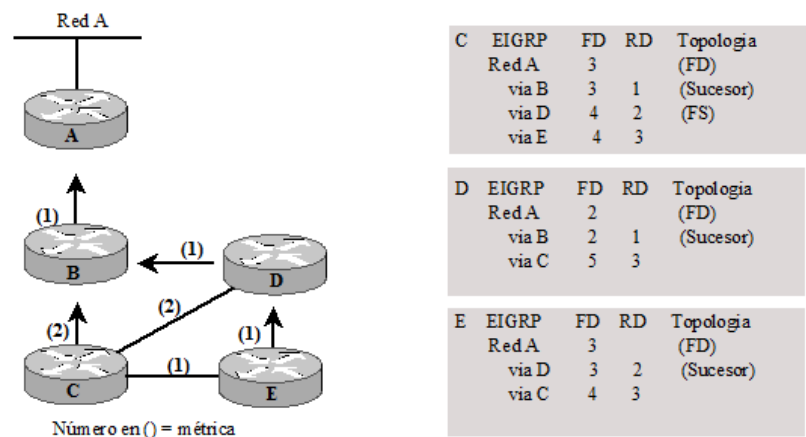


Figura 57. Convergencia DUAL algoritmo EIGRP

Leyenda	
EIGRP	Tipo de protocolo

FD	Distancia factible
RD	La distancia informada como la publica el Router vecino
Sucesor	Ruta primaria al destino
FS	Sucesor factible - Ruta de respaldo al destino

Cuadro 29. Leyenda de convergencia DUAL algoritmo EIGRP

- Cada tabla de topología identifica la siguiente información: El protocolo de enrutamiento o EIGRP
- El costo más bajo de la ruta, denominado distancia factible (FD)
- El costo de la ruta, según lo publica el Router vecino, denominado distancia informada (RD)

La columna de Topología identifica la ruta principal denominada ruta del sucesor (sucesor), y, cuando se identifica, la ruta de respaldo denominada sucesor factible (FS). Observe que no es necesario contar con un sucesor factible identificado.

La red EIGRP sigue una secuencia de acciones para permitir la convergencia entre los Routers, que actualmente tienen la siguiente información de topología:

- El Router C tiene una ruta del sucesor a través del Router B.
- El Router C tiene una ruta del sucesor factible a través del Router D.
- El Router D tiene una ruta del sucesor a través del Router B.
- El Router D no tiene una ruta del sucesor factible.
- El Router E tiene una ruta del sucesor a través del Router D.
- El Router E no tiene un sucesor factible.

Las normas para la selección de la ruta del sucesor factible se especifican a continuación

La del sucesor factible es una ruta de respaldo alternativa en caso de que la ruta del sucesor se desconecte.

La distancia informada (RD) al destino, tal como la publica el Router vecino, debe ser menor que la distancia factible (FD) de la ruta del sucesor primario.

Si se cumple este criterio, y no existe ningún bucle de enrutamiento, la ruta se puede seleccionar como la ruta de sucesor factible.

- La ruta del sucesor factible ahora se puede promover al estado de ruta de sucesor.
- Si el RD de ruta alternativa es igual a, o superior al FD del sucesor original, se rechaza la ruta como ruta de sucesor factible.
- El Router debe de recalcular la topología de la red reuniendo información de todos los vecinos.
- El Router envía un paquete de consulta a todos los vecinos para solicitar las rutas de enrutamiento disponibles a la red de destino y su costo de métrica correspondiente.

- Todos los Routers vecinos deben enviar un paquete de respuesta para contestar el pedido de paquete de consulta
- Los datos de respuestas se describen en la tabla de topología del Router que hace la consulta.
- Dual ahora pueden identificar nuevas rutas de sucesor y, donde resulte apropiado, nuevas rutas de sucesor factible a base de esta nueva información.

El siguiente ejemplo demuestra la forma en que cada Router de la topología aplica las normas de selección del sucesor factible cuando se desactiva la ruta del Router D al Router B:

- En el Router D
 - La ruta que pasa por el Router B se elimina de la tabla de topología.
 - Ésta es la ruta del sucesor. El Router D no cuenta con un sucesor factible identificado.
 - El Router D debe realizar un nuevo cálculo de ruta.

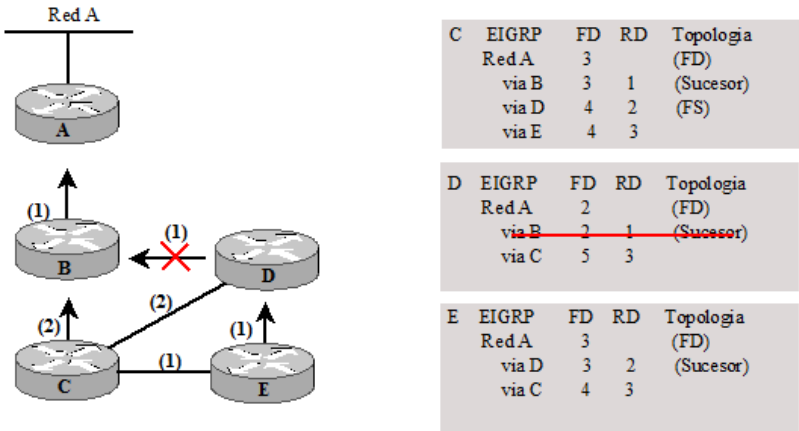


Figura 58. Selección de sucesor factible en la ruta de D a B

Leyenda	
EIGRP	Tipo de protocolo
FD	Distancia factible
RD	La distancia informada como la publico el Router vecino
Sucesor	Ruta primaria al destino
FS	Sucesor factible - Ruta de respaldo al destino

Cuadro 30. Leyenda de selección de sucesor factible en la ruta de D a B

- En el Router C
 - La ruta a la Red A, a través del Router D está deshabilitada.
 - La ruta que pasa por el Router D se elimina de la tabla.
 - Ésta es la ruta del sucesor factible para el Router C.

➤ En el Router D

- El Router D no tiene un sucesor factible. Por lo tanto, no puede cambiarse a una ruta alternativa identificada de respaldo.
- El Router D debe recalcular la topología de la red. La ruta al destino Red A se establece en Activa.
- El Router D envía un paquete de consulta a todos los Routers vecinos conectados para solicitar información de topología.
- El Router C tiene una entrada anterior para el Router D.
- El Router D no tiene una entrada anterior para el Router E.

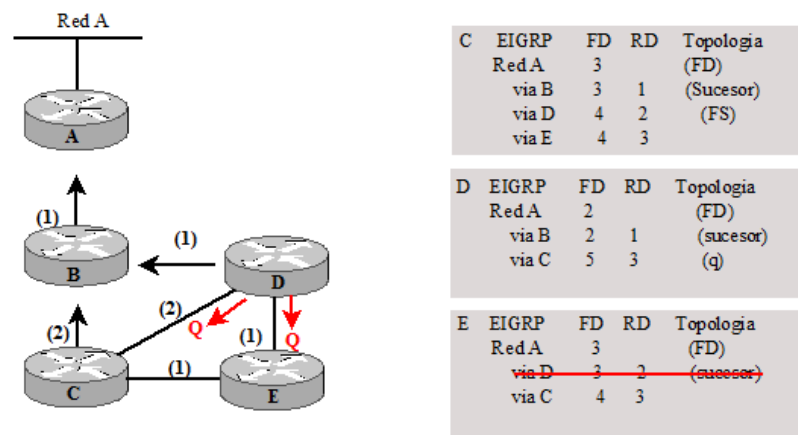


Figura 59. Sucesor factible de D

Leyenda	
EIGRP	Tipo de protocolo
FD	Distancia factible
RD	La distancia informada como la publico el Router vecino
Sucesor	Ruta primaria al destino
FS	Sucesor factible - Ruta de respaldo al destino

Cuadro 31. Leyenda sucesor factible de D

➤ En el Router E

- La ruta a la Red A a través del Router D está deshabilitada.
- La ruta que pasa por el Router D se elimina de la tabla.
- Ésta es la ruta del sucesor para el Router E.
- El Router E no tiene una ruta factible identificada.

- Observe que el costo RD de enrutar a través del Router C es 3. Este costo es igual al de la ruta del sucesor a través del Router D.

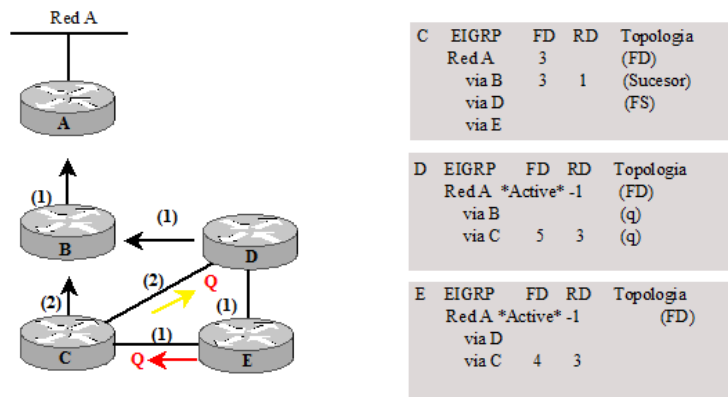


Figura 60. Sucesor factible de E

Leyenda	
C	Destino
EIGRP	Tipo de protocolo
FD	Distancia factible
RD	La distancia informada como la publico el Router vecino
Sucesor	Ruta primaria al destino
FS	Sucesor factible - Ruta de respaldo al destino

Cuadro 32. Leyenda sucesor factible de E

- En el Router C
 - El Router E envía un paquete de consulta al Router C.
 - El Router C elimina el Router E de la tabla.
 - El Router C responde al Router D con una nueva ruta a la Red A.
- En el Router D
 - La ruta al destino Red A sigue en estado Activa. El cálculo aún no se ha terminado.
 - El Router C ha respondido al Router D para confirmar que hay una ruta disponible al destino Red A con un costo de 5.
 - El Router D sigue esperando respuesta del Router E.
- En el Router E

- El Router E no tiene un sucesor factible para alcanzar el destino Red A.
- Por lo tanto el Router E rotula la ruta a la red destino como Activa.
- El Router E tiene que recalcular la topología de red.
- El Router E elimina de la tabla la ruta que pasa por el Router D.
- El Router D envía una consulta al Router C, para solicitar información de topología.
- El Router E ya tiene una entrada a través del Router C. Tiene un costo de 3, igual que la ruta del sucesor.

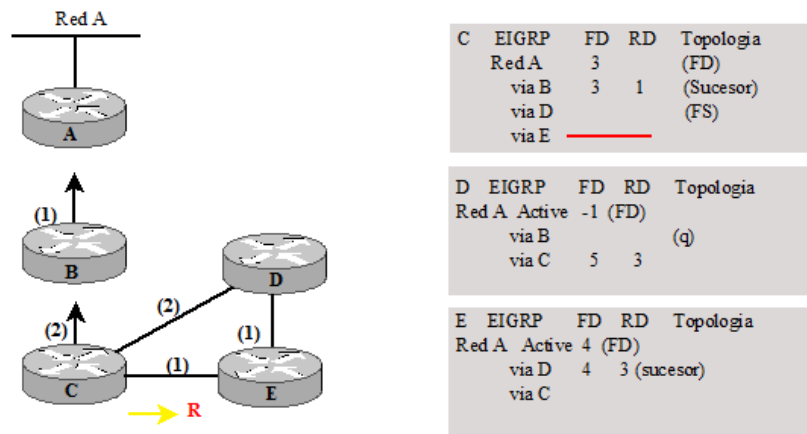


Figura 61. Elección final de sucesor factible de C, D y E

Leyenda	
C	Destino
EIGRP	Tipo de protocolo
FD	Distancia factible
RD	La distancia informada como la publico el Router vecino
Sucesor	Ruta primaria al destino
FS	Sucesor factible - Ruta de respaldo al destino

Cuadro 33. Leyenda elección final de sucesor factible de C, D y E

12.4.8 Configuración de EIGRP

A pesar de la complejidad de DUAL, la configuración de EIGRP puede ser relativamente sencilla. Los comandos de configuración de EIGRP varían según el protocolo que debe enrutarse. Algunos ejemplos de estos protocolos son IP, IPX y AppleTalk. Esta sección describe la configuración de EIGRP para el protocolo IP.

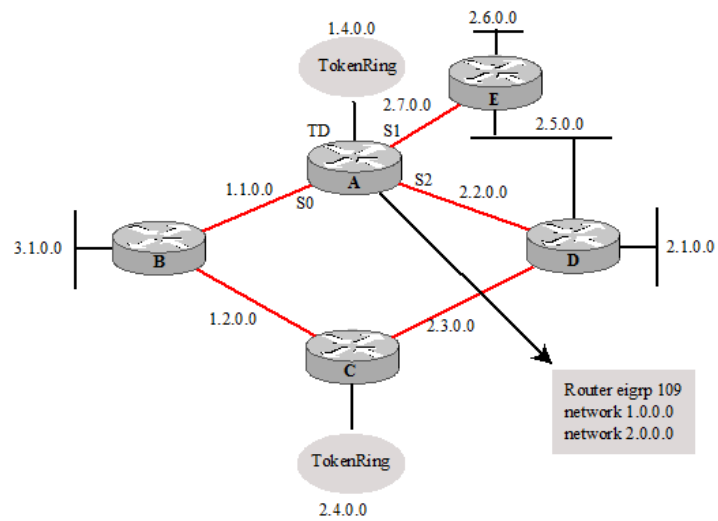


Figura 62. Topología de configuración de EIGRP

- Para habilitar EIGRP y definir el sistema autónomo utilice los siguientes comandos:

```
Router (config) #Router emigro autonomous-system-number
```

El número de sistema autónomo se usa para identificar todos los Routers que pertenecen a la internetwork. Este valor debe coincidir para todos los Routers dentro de la internetwork.

Indique cuáles son las redes que pertenecen al sistema autónomo EIGRP en el Router local

El siguiente comando indica cuáles son las redes que pertenecen al sistema autónomo EIGRP en el Router local

```
Router (config-Router) #network network-number Network-number
```

Es el número de red que determina cuáles son las interfaces del Router que participan en EIGRP y cuáles son las redes publicadas por el Router. El comando network configura sólo las redes conectadas. Por ejemplo, la red 3.1.0.0, que se encuentra en el extremo izquierdo de la Figura principal, no se encuentra directamente conectada al Router A. Como consecuencia, esa red no forma parte de la configuración del Router A.

Al configurar los enlaces seriales mediante EIGRP, es importante configurar el valor del ancho de banda en la interfaz. Si el ancho de banda de estas interfaces no se modifica, EIGRP supone el ancho de banda por defecto en el enlace en lugar del verdadero ancho de banda. Si el enlace es más lento, es posible que el Router no pueda convergir, que se pierdan las actualizaciones de enrutamiento o se produzca una selección de rutas por debajo de la óptima. Para establecer el ancho de banda para la interfaz, aplique la siguiente sintaxis:

```
Router (config-if) #bandwidth kilobits
```

Sólo el proceso de enrutamiento utiliza el comando bandwidth y es necesario configurar el comando para que coincida con la velocidad de línea de la interfaz.

Cisco también recomienda agregar el siguiente comando a todas las configuraciones EIGRP:

```
Router (config-Router) #eigrp log-neighbor-changes
```

Este comando habilita el registro de los cambios de adyacencia de vecinos para monitorear la estabilidad del sistema de enrutamiento y para ayudar a detectar problemas.

12.4.9 Configuración del resumen de EIGRP

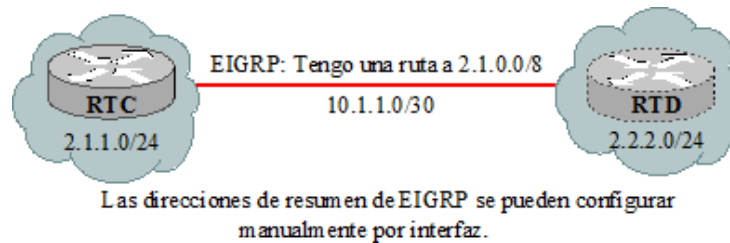


Figura 63. Configuración manual en la interfaz

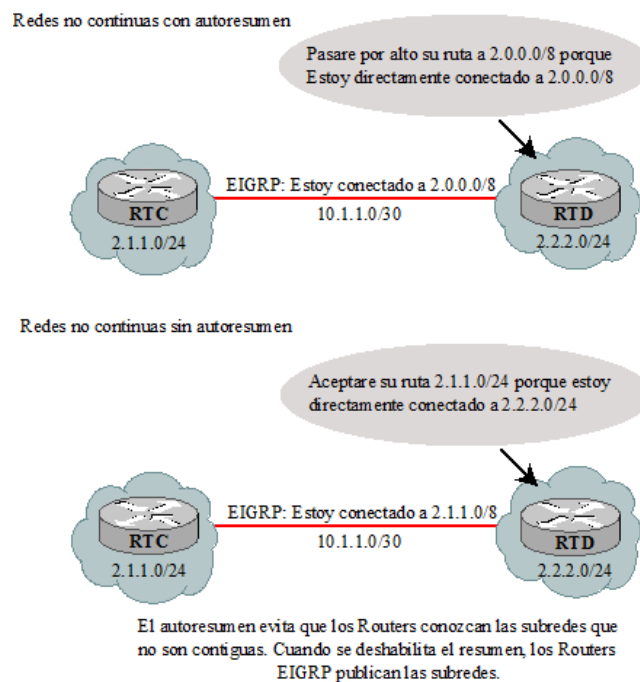


Figura 64. Auto resumen en redes no continuas

Sin embargo, es posible que el resumen automático no sea la mejor opción en ciertos casos. Por ejemplo, si existen subredes no contiguas el resumen automático debe deshabilitarse para que el enrutamiento funcione correctamente. Para desconectar el resumen automático, use el siguiente comando:

```
router(config-router)#no auto-summary
```

Con EIGRP, una dirección de resumen se puede configurar manualmente al configurar una red prefijo. Las rutas de resumen manuales se configuran por interfaz, de manera que la interfaz que propagará el resumen de ruta se debe seleccionar primero. Entonces, la dirección de resumen se puede definir con el comando

```
router(config-if)#ip summary-address eigrp autonomous-system-number ip-  
address mask administrativedistance
```

Las rutas de resumen EIGRP tienen una distancia administrativa por defecto de 5. De manera opcional, se pueden configurar con un valor entre 1 y 255.

RTC se puede configurar mediante los comandos que aparecen a continuación:

```
RTC(config)#router eigrp 2446  
RTC(config-router)#no auto-summary  
RTC(config-router)#exit  
RTC(config)#interface serial 0/0  
RTC(config-if)#ip summary-address eigrp 2446 2.1.0.0 255.255.0.0
```

Por lo tanto, RTC agrega una ruta a esta tabla de la siguiente manera:

```
D 2.1.0.0/16 is a summary, 00:00:22, Null0
```

Observe que la ruta de resumen se obtiene a partir de Null0 y no de una interfaz real. Esto ocurre porque esta ruta se usa para fines de publicación y no representa una ruta que RTC puede tomar para alcanzar esa red. En RTC, esta ruta tiene una distancia administrativa de 5.

RTD no es consciente del resumen pero acepta la ruta. A la ruta se le asigna la distancia administrativa de una ruta EIGRP normal, que es 90 por defecto.

En la configuración de RTC, el resumen automático se desactiva con el comando **no auto-summary**. Si no se desactivara el resumen automático, RTD recibiría dos rutas, la dirección de resumen manual, que es 2.1.0.0 /16, y la dirección de resumen automática con clase, que es 2.0.0.0 /8.

En la mayoría de los casos, cuando se hace el resumen manual, se debe ejecutar el comando **no autosummary**.

13. PROTOCOLO IPv6 (Direccionamiento)

13.1 IPv6

El protocolo Internet (IP) ha sido el fundamento de Internet y que virtualmente de todas las redes privadas de múltiples proveedores. Este protocolo está alcanzando el fin de su vida útil y se ha definido un nuevo protocolo conocido como IPv6 (IP versión 6) para, en última instancia, reemplazar a IP.

En primer lugar, examinaremos la motivación para desarrollar una nueva versión de IP y después analizaremos algunos de sus detalles.

13.1.1 IP de nueva generación

El motivo que ha conducido a la adopción de una nueva versión ha sido las limitaciones impuesta por el campo de dirección de 32 bits en IPv4. Con un campo de direcciones de 32 bits, en principio es posible asignar 2^{32} direcciones de diferentes, alrededor de 4.000 millones de direcciones posibles. Se podría pensar que este número de direcciones era más adecuado para satisfacer las necesidades en Internet. Sin embargo, a finales de la década de los ochentas se percibía que habría un problema y este problema empezó a manifestarse a comienzos de la década de los noventas. Algunas de las razones por las que es inadecuado utilizar estas direcciones de 32 bits son las siguientes:

- La estructura en dos niveles de la dirección IP (número de red, número de computador) es conveniente pero también es una forma poco económica utilizar el espacio de direcciones. Una vez que se le asigna un número de red a una red, todos los números de computador de ese número de red se asignan a esa red. El espacio de esa red podría estar poco usado, pero en lo que concierne a la efectividad del espacio de direcciones, si se usa un número de red entonces se consumen todas las direcciones dentro de la red.
- El modelo de direccionamiento IP requiere que se le asigne un número de red único a cada red IP independientemente de si la red está realmente conectada a Internet.
- Las redes están proliferando rápidamente. La mayoría de las organizaciones establecen LAN múltiples, no un único sistema LAN. Las redes inalámbricas están adquiriendo un mayor protagonismo. Internet mismo ha crecido explosivamente durante años.
- El uso creciente de TCP/IP en áreas nuevas producirá un crecimiento rápido en la demanda de direcciones únicas IP (por ejemplo, el uso de TCP/IP para interconectar terminales electrónicos de puntos de venta y para los receptores de televisión por cables).
- Normalmente, se asigna una dirección única a cada computador. Una disposición más flexible es permitir múltiples direcciones IP a cada computador. Esto, es por supuesto, incrementa la demanda de direcciones.

Por tanto, la necesidad de incremento en el espacio de direcciones ha impuesto la necesidad de una nueva versión de IP. Además, IP es un protocolo muy viejo y se han definido nuevos requisitos en las áreas de configuración de red, flexibilidad en el enrutamiento y funcionalidades para el tráfico.

En respuesta de estas necesidades, el grupo de Trabajo de Ingeniería de Internet (IETF) emitió las siguientes propuestas para una nueva generación de IP (IPng) en julio de 1992. Se recibieron varias propuestas y en 1994 emergió el diseño final del IPng. Uno de los hechos destacados fue el desarrollo de la publicación del RFC 1752, el

cual describe los requisitos de IPng, especifica el formato de la PDU y señala las técnicas de IPng en las áreas de direccionamiento, encaminamiento y seguridad. Existen otros documentos que definen los detalles del protocolo, ahora llamado oficialmente IPv6; estos incluyen una especificación general de IPv6 (RFC 2373), un RFC que trata sobre la estructura de direccionamiento de IPv6 (RFC 2373), y una larga lista adicional.

13.1.2 Mejoras de IPv6 sobre IPv4

- **Un espacio de direcciones ampliado:** IPv6 utiliza direcciones de 128 bits en un lugar de las direcciones de 32 bits de IPv4.
- **Un mecanismo de opciones mejorado:** las opciones de IPv6 se encuentran en cabeceras opcionales separadas situadas entre la cabecera IPv6 y la cabecera de la capa de transporte. La mayoría de estas cabeceras opcionales no se examinan ni procesan por ningún dispositivo de encaminamiento en la trayectoria del paquete. Esto simplifica y acelera el procesamiento que realiza un dispositivo de encaminamiento sobre IPv6 en comparación de los datagramas IPv4. Esto hace que sea más fácil incorporar opciones adicionales.
- **Autoconfiguración de direcciones:** Esta capacidad proporciona una asignación dinámica de direcciones IPv6.
- **Aumento de la flexibilidad en el direccionamiento:** IPv6 incluye el concepto de una dirección modo difusión (anycast), mediante la cual un paquete se entrega solo a un nodo seleccionando de entre un conjunto de nodos. Se mejora la escalabilidad del encaminamiento multidistribución con la incorporación de un campo de ámbito a las direcciones multidistribución.
- **Funcionalidad para la asignación de recursos:** en lugar del campo tipo de servicio de IPv4, IPv6 habilita el etiquetado de los paquetes como pertenecientes a un flujo de tráfico particular para que el emisor solicita un tratamiento especial. Esto ayuda al tratamiento de tráfico especializado como el de video en tiempo real.

13.1.3 Estructura de IPv6

Una unidad de datos (conocida como paquete) tiene el formato general mostrado en la Figura 65:

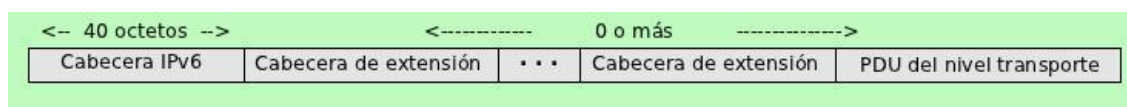


Figura 65. Formato de paquetes IPv6

La única cabecera que se requiere se denomina simplemente cabecera IPv6. Esta una longitud fija de 40 octetos, comparados con los 20 octetos de la parte obligatoria de la cabecera IPv4, se han definido las siguientes cabeceras de extensión:

13.1.3.1 Cabeceras de opciones salto a salto

Define opciones especiales que requieren procesamiento en cada salto.

13.1.3.2 Cabecera de encaminamiento

Proporciona un encaminamiento ampliado, similar al encaminamiento en el origen de IPv4.

13.1.3.3 Cabecera de fragmentación

Contiene información de fragmentación y reensablado.

13.1.3.4 Cabecera de autenticación

Proporciona la integridad del paquete y la autenticación.

13.1.3.5 Cabecera de encapsulamiento de la carga de seguridad

Proporciona privacidad.

13.1.3.6 Cabeceras de las opciones para el destino

Contiene información opcional para que sea examinada en el nodo destino.

El estándar de IPv6 recomienda que, en el caso de que se usen varias cabeceras de extensión, las cabeceras IPv6 aparezcan el siguiente orden:

- Cabecera IPv6: obligatoria, debe aparecer siempre primero.
- Cabeceras de las opciones salto a salto.
- Cabeceras de opciones para el destino: para opciones a procesar por el primer destino que aparece en el campo dirección IPv6 de destino y por los destinos subsecuentes indicados en la cabecera de encaminamiento.
- Cabecera de encaminamiento.
- Cabecera de fragmentación.
- Cabecera de autenticación.
- Cabecera de encapsulado de la carga de seguridad.
- Cabeceras de opciones para el destino: para opciones a procesar por el destino final del paquete.

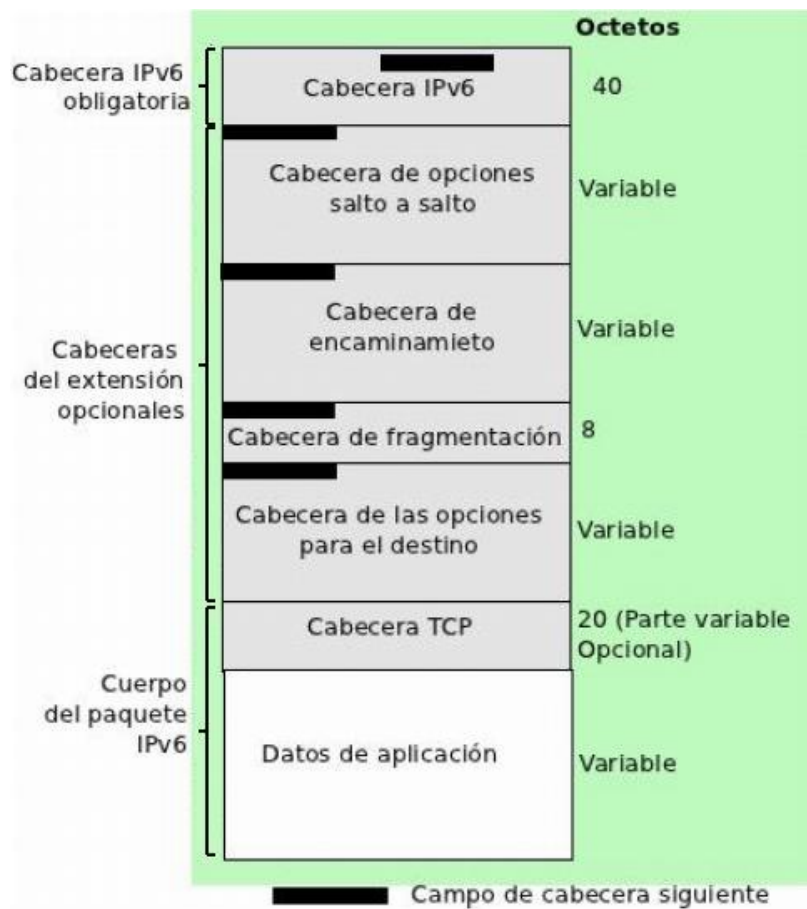


Figura 66. Cabeceras de opciones para el destino

En la Figura 66 se muestra un ejemplo de un paquete IPv6 que incluye un ejemplar de cada cabecera, excepto aquellas relacionadas con la seguridad. Obsérvese que la cabecera IPv6 y cada cabecera de extensión incluyen el campo cabecera siguiente. Este campo identifica el tipo de cabecera que viene a continuación. Si la siguiente cabecera es de extensión, entonces este campo contiene el identificador de tipo de esa cabecera. En caso contrario, este campo contiene el identificador del protocolo de la capa superior que se está usando en IPv6 (normalmente un protocolo de la capa de transporte), utilizando el mismo valor que el campo del protocolo IPv4. En la Figura 66 se muestra el protocolo de la capa superior es TCP; por tanto, los datos de la capa superior transportados por el paquete IPv6 constan de una cabecera TCP seguida por un bloque de datos de aplicación.

A continuación, se examina la cabecera principal de IPv6 y después se examinan cada una de las extensiones.

13.1.4 Cabecera IPv6

La cabecera de IPv6 tiene una longitud fija de 40 octetos, que consta con los siguientes campos (véase la Figura 67):

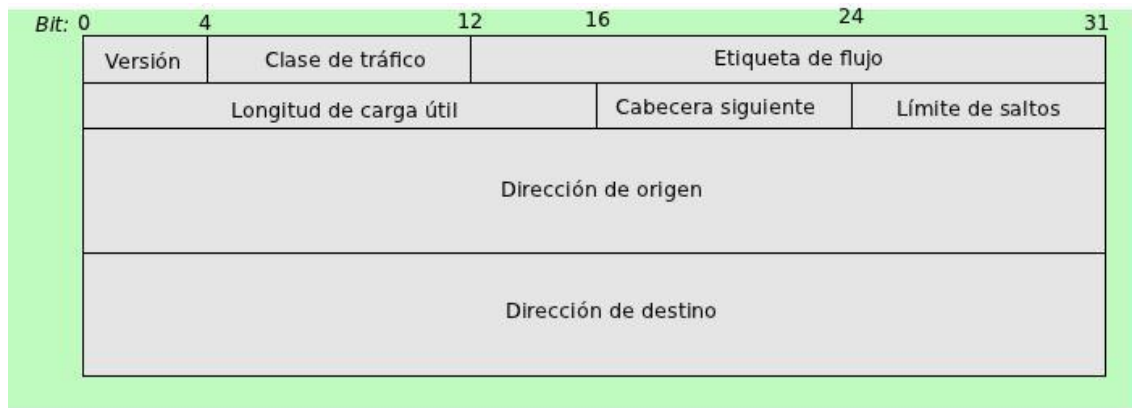


Figura 67. Cabecera IPv6

Versión 4 bits

Número de la versión del protocolo Internet; el valor es 6

- **Clase de tráfico (8 bits):** Disponible para su uso por el nodo origen y/o los dispositivos de encaminamiento para identificar y distinguir entre clases o prioridades de paquetes IPv6. Este campo se usa actualmente para los campos de ceros en y ECN, como se describió para el campo tipo de servicio en IPv4.
- **Etiqueta de flujos (20 bits):** Se puede utilizar por un computador para etiquetar a aquellos paquetes para los que requieren un tratamiento especial en los dispositivos de encaminamiento dentro de la red.
- **Longitud de la carga útil (16 bits):** Longitud del resto de paquetes IPv6 excluidas la cabecera, en octetos. En otras palabras, representa la longitud de todas las cabeceras de extensión más de la PDU de la capa de transporte.
- **Cabecera siguiente (8 bits):** Identifica el tipo de cabecera que sigue inmediatamente a la cabecera IPv6; se puede tratar como una cabecera de extensión IPv6 como de una cabecera de la capa superior, como TCP o UDP.
- **Límite de saltos (8 bits):** El número restante de saltos permitidos para este paquete. El límite de saltos se establece por la fuente o algún valor máximo deseado y se decrementa en 1 en cada nodo que reenvía el paquete. El paquete se descarta si el límite de saltos se hace cero. Esto es una simplificación del procesamiento requerido por el campo tiempo de vida de IPv4. El consenso fue el esfuerzo de extra de contabilizar los intervalos de tiempo en IPv4 no añadía el valor significativo al protocolo. De hecho, y como regla general, los dispositivos de encaminamiento IPv4 tratan el campo tiempo de vida como un límite de saltos.
- **Dirección origen (128 bits):** Dirección del productor del paquete.

- **Dirección destino (128 bits):** Dirección de destino deseado del paquete. Puede este no sea en realidad el último destino deseado si está presente la cabecera de encaminamiento.

13.1.5 Direcciones IPv6

Las direcciones IPv6 tienen una longitud de 128 bits. Las direcciones se asignan a interfaces individuales en los nodos, no a los nodos. Una única interfaz puede tener múltiples direcciones únicas. Cualquiera de las direcciones asociadas a las interfaces de los nodos se puede utilizar para identificar de forma única el nodo.

La combinación de direcciones largas y de direcciones múltiples por interfaz permite una eficiencia mejorada del encaminamiento con respecto a IPv4. En IPv4, generalmente las direcciones no tienen una estructura que ayude al encaminamiento y, por tanto, un dispositivo de encaminamiento necesita tener una gran tabla con rutas de encaminamiento. Una dirección internet más grande permite agrupar las direcciones por jerarquías de red, por proveedores de acceso, por proximidad geográfica, por institución, etc. Estas agrupaciones deben conducir a tablas de encaminamientos más pequeñas y a consultas más rápidas. El primer múltiples de direcciones por interfaz posibilita a un abandono que utiliza varios proveedores de acceso a través de la misma interfaz, tener direcciones distintas agrupadas bajo el espacio de direcciones de cada proveedor.

La arquitectura de las direcciones IPv6 se encuentra descrita en RFC4291.

128 bits = 16 bytes e agrupan los bytes de 2 en 2, se separan por ":" y se representan en hexadecimal de la siguiente manera:

2001 : 0DB8 : 3C4D : 5E6F : 0216 : CBFF : FEB2 : 7474

Simplificación:

:: representa uno o varios grupos de 2 bytes a 0. Sólo puede usarse una vez en una dirección IPv6.

2001 : 0DB8 : 0000 : 0000 : 0008 : 0800 : 200C : 417A

2011 : DB8 : 0: 0: 800 : 200C : 417A

2001 : DB8 :: 8 : 800 : 200C : 417A

Las direcciones IPv6 están divididas en 3 campos:

- **Prefijo de red:** Conjunto de direcciones que se le asignan a una organización. Los ISPs suelen tener prefijos /32. Las grandes organizaciones normalmente tienen /48.
- **Identificador de subred:** identifica una determinada subred dentro de una organización.
- **Identificador de máquina (64 bits):** identifica a una interfaz de una máquina dentro de una subred.

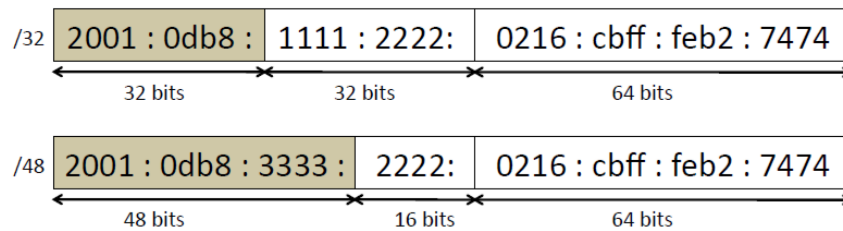


Figura 68. Divisiones de los campos de direcciones IPv6

13.1.5.1 Tipos de direcciones IPv6

IPv6 permite tres tipos de direcciones:

- **Unidifusión (unicast):** Un identificador para una interfaz individual. Un paquete enviado a una dirección de este tipo se entrega a la interfaz identificada por esa dirección.
 - Utilizadas dentro de un mismo enlace o la misma red local.
 - Los paquetes enviados a este tipo de dirección no van a ser encaminados por ningún router.
 - Necesarias para "Neighbor Discovery"
 - Se configuran automáticamente.

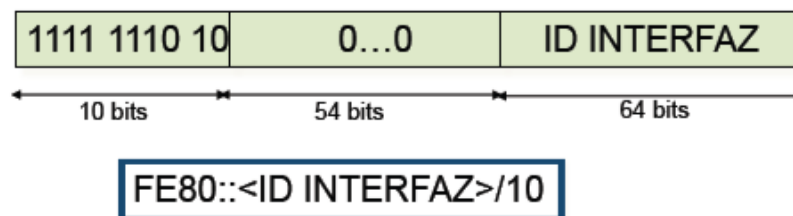


Figura 69. Direcciones IPv6 UNICAST

- Se construye el identificador de interfaz utilizando la dirección MAC de la tarjeta Ethernet (48 bits).

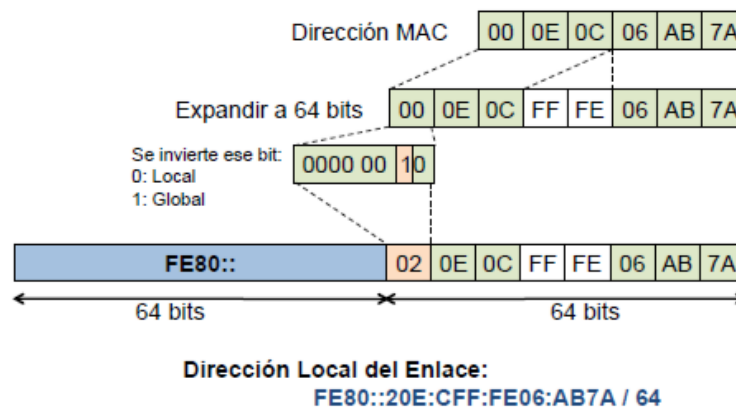


Figura 70. Identificador de Interfaz UNICAST

- **Monodifusión (anycast):** Un identificador para un conjunto de interfaces (normalmente pertenecientes a diferentes nodos). Un paquete enviado a una dirección monodifusión se entrega a una de las interfaces identificadas por esa dirección (la más cercana, de acuerdo a la medida de distancia de los protocolos de encaminamiento).
 - Son direcciones que se asignan a más de una interfaz en diferentes máquinas.
 - Un paquete enviado a una dirección anycast llegaría a una de la máquinas que tenga configurada dicha dirección anycast, a la más cercana.
 - Son indistinguibles de las direcciones unicast.
 - Si todas las máquinas a las que se quiere asignar una dirección anycast se encuentran en la misma organización, la dirección anycast tendrá el mismo prefijo que las direcciones unicast de ese sitio.
 - La dirección anycast de los routers que están conectados a una subred:

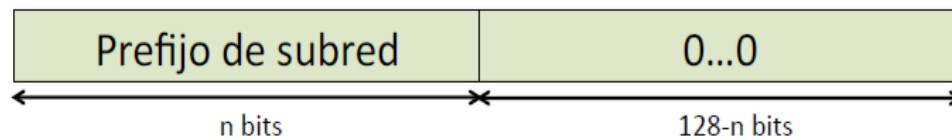


Figura 71. Monodifusión (anycast)

- **Multidifusión (multicast):** Un identificador para un conjunto de interfaces (normalmente perteneciente a diferentes nodos). Un paquete enviado a una dirección multidifusión se entrega en todas las interfaces identificadas por esa dirección.
 - Son direcciones que se utilizan para direccionar un conjunto de interfaces, también se denominan grupos de multicast.

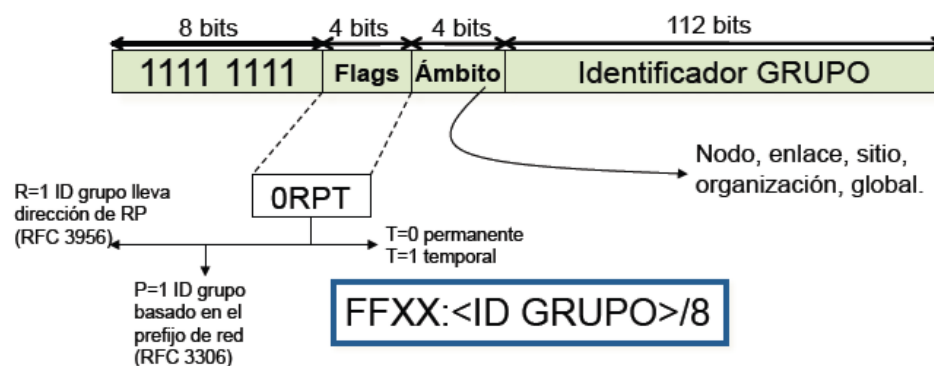


Figura 72. Multidifusión IPv6 (Multicast)

Algunas de las direcciones multicast reservadas:

Todos los nodos	FF02:0:0:0:0:0:1	link-local
Todos los routers	FF02:0:0:0:0:0:2	link-local
De nodo solicitado	FF02:0:0:0:1:FFXX:XXXX	link-local

FF02:0:0:0:1:FF00::/104

13.1.5.2 Direcciones IPv6 que deber reconocer una máquina/router como suyas**➤ Maquinas**

- Dirección local de enlace por cada interfaz.
- Dirección Unicast y Anycast que se desee configurar.
- Dirección de loopback .
- Dirección multicast de todos los nodos (FF02::1).
- Dirección multicast de nodo solicitado, por cada dirección unicast y anycast que tiene configurada.
- Direcciones de multicast que tenga configuradas.

➤ Routers

- Por cada una de las interfaces del router, debe reconocer la dirección anycast de los routers que están conectados a una subred.
- Dirección de multicast de todos los routers (FF02::2).

En el Cuadro 34 muestra la distribución de los tipos de direcciones IPv6

Notación IPv6	Prefijo Binario	Tipo de Dirección
0::0/128	00...0 (128 bits)	Dirección sin especificar
0::1/128	00...1 (128 bits)	Dirección de loopback
FC00::/7	1111 110 ...	Unique Local Unicast (RFC4193)
FE80::/10	1111 1110 10	Link Local Unicast (RFC4291)
FF00::/8	1111 1111	Multicast (RFC4291)
Resto de direcciones: Global Unicast (RFC4291)		
2000::/3	001	Prefijo que está asignando el IANA
2001:DB8::/32		Documentación (RFC3849)

Cuadro 34. Distribución de los tipos de direcciones IPv6

13.1.5.3 Direcciones IPv4 contra direcciones IPv6

IPv4	IPv6
32 bits	128 bits
Notación decimal: a.b.c.d	Notación hexadecimal (cada X son 4 dígitos hexadecimales): X:X:X:X:X:X:X:X
Organización CIDR: Prefijo + subred + id_máquina	Organización: Prefijo + subred + id_máquina

Tipos de direcciones:	Tipos de direcciones:
Unicast	Unicast
Multicast	Multicast
Broadcast	Broadcast

Cuadro 35. Comparaciones de direcciones IPv4 e IPv6

13.1.6 Cabeceras de opciones salto a salto

Transporta información opcional que, si está presente, debe ser examinada por cada dispositivo de encaminamiento a lo largo de la ruta. Esta cabecera contiene los siguientes campos:

- **Cabecera siguiente (8 bits):** Identifica el tipo de la cabecera que sigue inmediatamente a ésta.
- **Longitud de la cabecera de extensión (8 bits):** Longitud de la cabecera en unidades de 64 bits, sin incluir los primero 64 bits.
- **Opciones:** Campo de longitud variable que consta de una o más definiciones de opción. Cada definición se expresa mediante tres subcampos: tipo de opción (8 bits), que identifica la opción; longitud (8 bits) que especifica la longitud en octetos del campo de datos de la opción; y datos de opción, que es una especificación de la opción de longitud variable.

En realidad, se utilizan los cinco bits menos significativos del campo de tipo opción para especificar una opción particular. Los bits más significativos indican la acción que tiene que realizar un nodo que no reconoce el tipo de opción, de acuerdo a:

00—ignorar esta opción y continuar procesando la cabecera.

01—descartar el paquete.

10—descartar el paquete y enviar un mensaje ICMP de problema de parámetro a la dirección de origen del paquete, indicando el tipo de opción no reconocida.

11—descartar el paquete y, solamente si la dirección destino del paquete no es una dirección multidifusión, enviar un mensaje ICMP de problema de parámetro a la dirección origen del paquete, indicando el tipo de opción no reconocida.

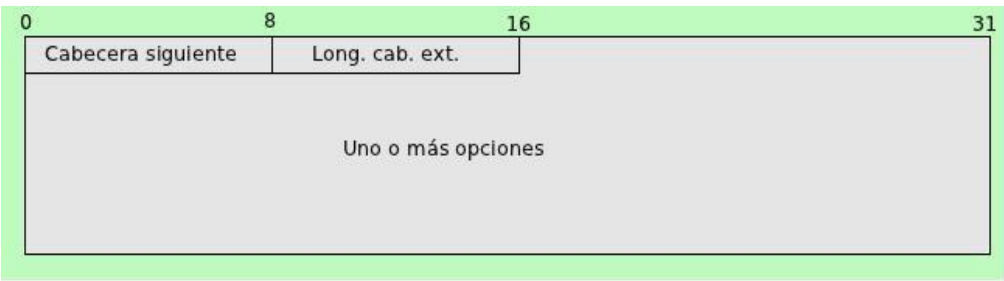
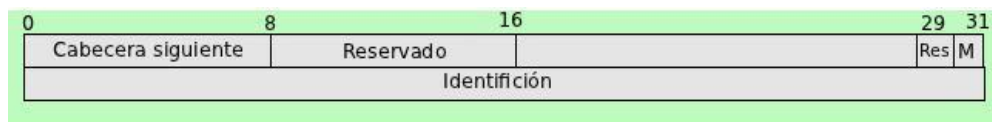
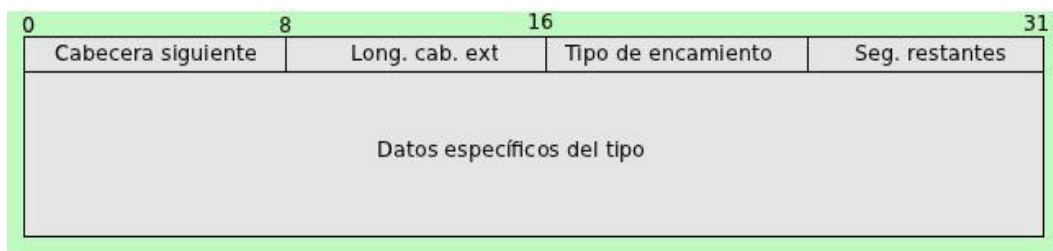
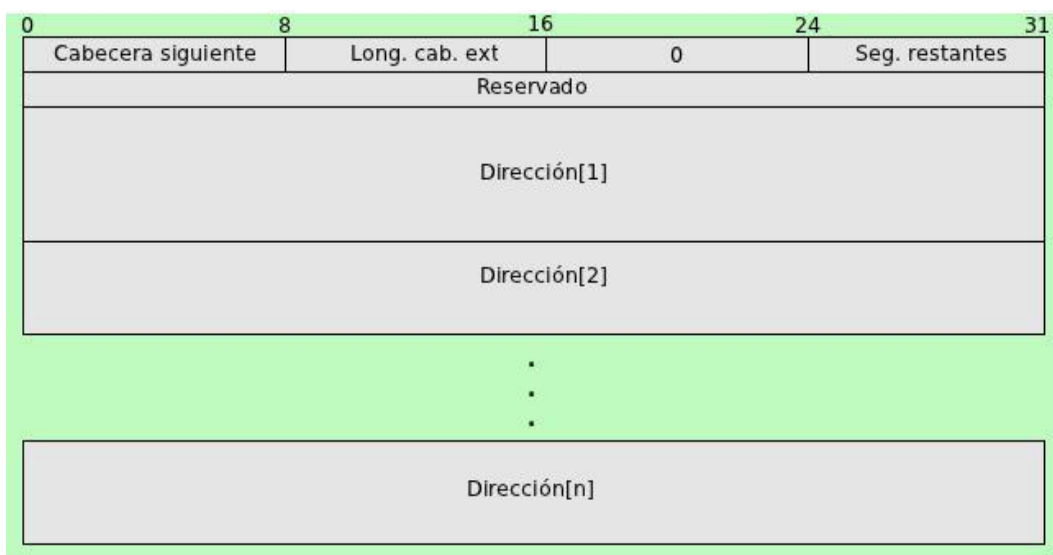


Figura 73. Cabeceras de opciones salto a salto para el destino

**Figura 74. Cabecera de fragmentación****Figura 75. Cabecera de encaminamiento genérica****Figura 76. Cabecera de encaminamiento tipo 0**

El tercer bit indica si el campo de datos de la opción no cambia (0) o si puede cambiar en (1) en el camino desde el origen al destino. Los datos que pueden cambiar se deben excluir de los cálculos de autenticación.

Estas convenciones para el campo del tipo de opción también se aplican a la cabecera de opciones del destino.

Hasta ahora se han especificado cuatro opciones salto a salto:

- **Relleno1:** Utilizada para insertar un byte de relleno dentro de la zona de opciones de la cabecera.

- **RellenoN:** Utilizada para insertar N bytes ($N \geq 2$) de relleno dentro de la zona de opciones de la cabecera. Las dos opciones de relleno aseguran que la cabecera tiene una longitud múltiplo de 8 bits.
- **Carga útil Jumbo:** Se utiliza para enviar paquetes con una carga útil mayor de 65.535 octetos. El campo de datos de esta opción tiene una longitud de 32 bits y da la longitud del paquete en octetos, excluyendo la cabecera IPv6 debe estar en cero y no puede haber cabecera de fragmentación. Con esta opción, IPv6 permite tamaños de paquetes de hasta 4.000 millones de octetos. Esto facilita la transmisión de paquetes de video grandes y posibilita que IPv6 hasta el mejor uso de la capacidad disponible sobre cualquier medio de transmisión.
- **Alerta al dispositivo de encaminamiento:** informa al dispositivo de encaminamiento que el contenido de este paquete es de interés para el dispositivo de encaminamiento y para tratar adecuadamente cualquier información de control. La ausencia de esta opción en un datagrama IPv6 informa al dispositivo de encaminamiento que el paquete no contiene información necesaria para el dispositivo de encaminamiento y, por tanto, puede eliminarlo de forma segura sin ningún análisis adicional.

13.1.7 Cabecera de fragmentación

La fragmentación solo puede ser realizada por el nodo origen, no por los dispositivos de encaminamiento a lo largo del paquete. Para obtener todas las ventajas del entorno de interconexión, un nodo debe ejecutar un algoritmo de obtención de la ruta, que le permita conocer la unidad máxima de transferencia (MTU) permitida por cada red en la ruta. Con este conocimiento, el nodo origen fragmentará el paquete, según se requiera, para cada dirección de destino dada. Si no se ejecuta este algoritmo, el origen debe limitar todos los paquetes a 1.280 octetos, que deben ser la mínima MTU que permita las redes.

La cabecera de fragmentación contiene los siguientes campos:

- **Cabecera siguiente (8 bits):** Identifica el tipo de cabecera que sigue inmediatamente a ésta.
- **Reservado (8 bits):** reservados para uso futuros.
- **Desplazamiento del fragmento (13 bits):** indica dónde se sitúa en el paquete original la carga útil de este fragmento. Se mide en unidades de 64 bits. Esto implica que los fragmentos (excepto el último), deben contener un campo de datos con una longitud de múltiplo de 64 bits.
- **Reservado (2 bits):** reservados para usos futuros.
- **Indicador M (1 bits):** 1 = más fragmentos; 0 = último fragmento.
- **Identificación (32 bits):** utilizado para identificar de forma única el paquete original. El identificador debe ser único para la dirección origen y direccionamiento destino durante el tiempo que el paquete permanece en Internet. Todos los fragmentos con el mismo identificador, dirección origen y dirección destino son reensamblados para recuperar el paquete original.

13.1.8 Cabecera de encaminamiento

Contiene una lista de uno o más nodos intermedios por lo que se pasa por el camino del paquete a su destino. Todas las cabeceras de encaminamiento comienzan con un bloque de 32 bits consistentes en 4 campos de 8 bits, seguido por datos de encaminamiento específicos al tipo de encaminamiento dado. Los cuatro campos de 8 bits son los siguientes:

- **Cabecera siguiente (8 bits):** identifica el tipo de cabecera que sigue inmediatamente a ésta.
- **Longitud de la cabecera de extensión (8 bits):** longitud de esta cabecera en unidades de 64 bits, sin incluir los primeros 64 bits.
- **Tipo de encaminamiento (8 bits):** identifica una variante particular de cabecera de encaminamiento. Si un dispositivo de encaminamiento no reconoce el valor del tipo de encaminamiento, debe descartar el paquete.
- **Segmento restante (8 bits):** número de segmentos en la ruta de quedan; esto es, el número de nodos intermedios explícitamente contenidos en la lista que se visitarán todavía antes de alcanzar el destino.

13.1.9 Cabecera de opciones del destino

Lleva información opcional que, si está presente, se examina por el nodo destino del paquete. El formato de esta cabecera es el mismo que el de la cabecera de opciones salto a salto.

14. PROTOCOLOS DE ENRUTAMIENTO CON IPv6

14.1 OSPF3

OSPFv3 es el protocolo de enrutamiento OSPF para IPv6. En 1999, OSPFv3 para IPv6 se publicó en RFC 2740. John Moy, Rob Coltun y Dennis Ferguson desarrollaron RFC 2740. Es la versión de OSPF para IPv6 basada en OSPFv2, con varias adiciones usada para distribuir prefijos de IPv6, utiliza IPv6 como transporte aunque tiene el mismo nombre que OSPFv2, son dos protocolos diferentes.

- OSPFv3: La versión 3 de OSPF fue creada para que, a diferencia de la versión 2, pueda soportar direccionamiento IPv6.
- La mayoría de las características son las mismas en una versión que en la otra
- En OSPF para IPv6, el "routing process" no necesita ser explícitamente creado
- Habilitando OSPF para IPv6 en la interfaz, el proceso será creado.
- En OSPF para IPv6, cada interfaz debe ser habilitada con un comando en modo de configuración de interfaz.
- Esto lo diferencia de OSPFv2, donde las interfaces quedan automáticamente habilitadas con un comando de configuración global.
- Al mismo tiempo, se pueden configurar varios prefijos en una única interfaz.
- Cuando hacemos esto, en OSPF para IPv6, todos los prefijos de la interfaz serán anunciados por OSPF.
- No se podrá elegir qué prefijos serán importados dentro del OSPF (todos o ninguno)

14.1.1 Configuración de OSPF para IPv6

Habilitando el ruteo para IPv6:

```
Router> enable
Router# configure terminal
Router(config)# ipv6 unicast-routing
```

Habilitando OSPF para IPv6

```
Router> enable
Router# configure terminal
Router(config)# interface <type> <number>
Router(config-if)# ipv6 ospf <process-id> area <area-id>
```

Definiendo Area Range

```
Router> enable
Router# configure terminal
Router(config)# ipv6 router ospf <process-id>
Router(config-rtr)# area <area-id> range <ipv6-prefix/ prefix-length>
```

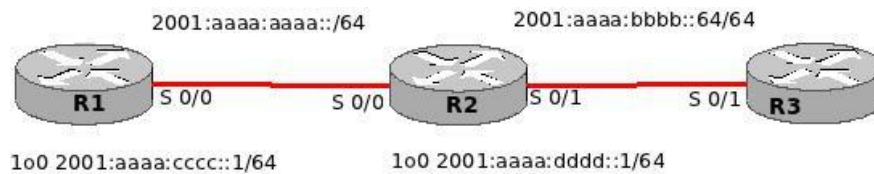


Figura 77. Topología de muestra para implementación de OSPF3 con IPv6

Siguiendo el esquema de a continuación vamos a configurar OSPF con IPv6. Vamos a modificar las link local address que se auto configuran en el emulador porque todos los routers tienen la misma MAC y el protocolo DAD Duplicate Address Detection no te deja usar la interfaz con IP duplicada. Primero configuramos las interfaces de cada router sin olvidar ipv6 unicast-routing en el modo de configuración global.

R1

```
interface Loopback0
no ip address
ipv6 address 2001:AAAA:CCCC::1/64
end
interface Serial0/0/0
no ip address
ipv6 address 2001:AAAA:AAAA::1/64
ipv6 address FE80::1 link-local
serial restart-delay 0
end
```

R2

```
interface Serial0/0
no ip address
ipv6 address 2001:AAAA:AAAA::2/64
ipv6 address FE80::2 link-local
serial restart-delay 0
end
interface Serial0/1
no ip address
ipv6 address 2001:AAAA:BBBB::2/64
ipv6 address FE80::3 link-local
serial restart-delay 0
end
```

R3

```
interface Loopback0
no ip address
ipv6 address 2001:AAAA:DDDD::1/64
end
interface Serial0/1
no ip address
ipv6 address 2001:AAAA:BBBB::1/64
ipv6 address FE80::4 link-local
serial restart-delay 0
end
```

Para verificar como quedan las interfaces después de la configuración realizada anteriormente, usar el siguiente comando:

```
show ipv6 int brie | excl admin
```

Puesto que los routers no tienen ninguna IP en formato IPv4 hay que crear obligatoriamente un ID de 32 bits para cada uno.

R1

```
ipv6 router ospf 1  
router-id 0.0.0.1
```

R2

```
ipv6 router ospf 1  
router-id 0.0.0.2
```

R3

```
ipv6 router ospf 1  
router-id 0.0.0.3
```

Para habilitar el enrutamiento, hay que decidir en cada interfaz que habla ipv6, que área va a enrutar, el comando network queda desfasado en RIP OSPF y EIGRP.

R1

```
interface loopback 0  
  ipv6 ospf 1 area 0  
interface serial0/0  
  ipv6 ospf 1 area 0
```

R2

```
interface serial0/0  
  ipv6 ospf 1 area 0  
interface serial0/1  
  ipv6 ospf 1 area 0
```

R3

```
interface serial0/1  
  ipv6 ospf 1 area 0
```

Verificar que interfaz habla por OSPF

```
show ipv6 protocols
```

Verificar enrutamiento IPv6 con OSPF

```
show ipv6 route ospf
```

14.2 IS-IS

- Las características de IPv6 para IS-IS 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.

14.2.1 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 IS-IS 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.

14.2.2 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 family.
- Los routers construyen una topología por cada protocolo de capa 3, por lo que, pueden encontrar el camino óptimo (SPF) aun si algún link soporta solo uno de estos protocolos.

14.2.3 Configuración de IS-IS para IPv6

- Comprende 2 pasos:
 - Antes que nada, crear un proceso IS-IS routing process, esto independiente del protocolo.
 - La segunda actividad es configurar el protocolo IS-IS en la
- interfaz
 - Pre-requisito:
 - Tener IPv6 unicast-routing habilitado

Configurando el proceso IS-IS

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

Configurando las interfaces

```
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 Multi-topology

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

```
Router(config-router)# metric-style wide [level-1 | level-2
| level-1-2]
Router(config-router)# address-family ipv6 [unicast |
multicast]
Router(config-router-af)# multi-topology
```

14.3 RIPng

RIPng es un protocolo de enrutamiento vector distancia con un límite de 15 saltos que usa actualizaciones de envenenamiento en reversa y horizonte dividido para evitar routing loops. Su simplicidad proviene del hecho de que no requiere ningún conocimiento global de la red. Sólo los routers vecinos intercambian mensajes locales, debe ser implementado solo en routers, sigue implementando la misma métrica que RIPv1, las tablas de enrutamiento presentes en los routers contienen entradas con la siguiente información: El prefijo Ipv6 de destino, la métrica o número de saltos para llegar a este destino, la dirección del siguiente salto (esta dirección debe ser Ipv6), una bandera que indica los cambios recientes en el estado de la ruta y los temporizadores asociados a la entrada.

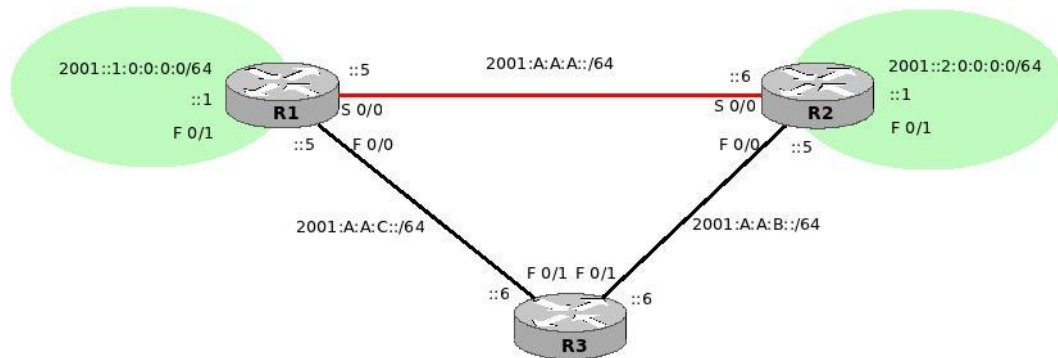


Figura 78. Topología de muestra para implementación de RIPng con IPv6

En la Figura 78 existen 3 Routers y dos redes LAN que se unirán mediante direccionamiento IPv6. Los bloques asignados están escritos y para simplificar la configuración se han dejado todos en /64.

Configuraciones de las direcciones IP en cada interfaz de cada router.

R1

```
R1(config)#int s0/0
R1(config-if)#ipv6 address 2001:A:A:A::5/64
R1(config-if)#no shutdown
R1(config-if)#int f0/0
R1(config-if)#ipv6 address 2001:A:A:C::5/64
R1(config-if)#no shutdown
R1(config-if)#int f0/1
R1(config-if)#ipv6 address 2001:0:0:1::1/64
R1(config-if)#no shutdown
```

R2

```

R2(config)#int s0/0
R2(config-if)#ipv6 address 2001:A:A:A::6/64
R2(config-if)#no shutdown
R2(config-if)#int f0/0
R2(config-if)#ipv6 address 2001:A:A:B::5/64
R2(config-if)#no shutdown
R2(config-if)#int f0/1
R2(config-if)#ipv6 address 2001::2:0:0:0:1/64
R2(config-if)#no shutdown

```

R3

```

R3(config)#int f0/0
R3(config-if)#ipv6 address 2001:A:A:C::6/64
R3(config-if)#no shutdown
R3(config-if)#int f0/1
R3(config-if)#ipv6 address 2001:A:A:B::6/64
R3(config-if)#no shutdown

```

Indispensable es activar en todos los routers el enrutamiento IPv4 que viene desactivado de manera predeterminada. El comando de modo global es `IPv6 unicast-routing` (similar a IP routing).

```
ipv6 unicast-routing
```

Para habilitar RIPng solamente se debe ingresar a la interfaz de router que se desea publicar en el proceso RIP e ingresar el comando `ipv6 rip IDENTIFICADOR enable` donde "IDENTIFICADOR" es un ID de proceso al más puro estilo OSPF. Este valor puede ser un número o una palabra. A continuación ingresaremos en todas las interfaces de R1, R2 y R3 para ingresar este comando. Note que la interfaz f0/1 de R1 y R3 no conectan con ningún otro router, pero sin embargo en ellas también se debe habilitar RIPng para que esas redes se publiquen.

R1

```

R1(config)#int f0/0
R1(config-if)#ipv6 rip IDENTIFICADOR1 enable
R1(config-if)#int f0/1
R1(config-if)#ipv6 rip IDENTIFICADOR1 enable
R1(config-if)#int s0/0
R1(config-if)#ipv6 rip IDENTIFICADOR1 enable
R1(config-if)#end

```

R2

```

R2(config)#int f0/0
R2(config-if)#ipv6 rip IDENTIFICADOR1 enable
R2(config-if)#int f0/1
R2(config-if)#ipv6 rip IDENTIFICADOR1 enable
R2(config-if)#int s0/0

```

```
R2(config-if)#ipv6 rip IDENTIFICADOR1 enable  
R2(config-if)#end
```

R3

```
R3(config)#int f0/0  
R3(config-if)#ipv6 rip IDENTIFICADOR1 enable  
R3(config-if)#int f0/1  
R3(config-if)#ipv6 rip IDENTIFICADOR1 enable  
R3(config-if)#end
```

NOTA: 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.

El siguiente comando se usa para verificar si la tabla de enrutamiento se ha actualizado con las redes aprendidas por RIPv6.

```
show ipv6 route
```

15. POLÍTICAS DE SEGURIDAD EN REDES DE COMUNICACIONES

15.1 Protocolo SSH (Secure Shell)

SSH, o "*Secure Shell*" es tanto una aplicación como un protocolo, que permite conectar dos ordenadores a través de una red, ejecutar comandos de manera remota y mover ficheros entre los mismos. Proporciona autenticación fuerte y comunicaciones seguras sobre canales no seguros y pretende ser un reemplazo seguro para aplicaciones tradicionales no seguras, como telnet, rlogin, rsh y rcp. Su versión 2 (SSH2) proporciona también sftp, un reemplazo seguro para FTP.

Además de esto, SSH permite establecer conexiones seguras a servidores XWindows, realizar conexiones TCP seguras y realizar acciones como sincronización de sistemas de ficheros (rsync) y copias de seguridad por red, de manera segura.

SSH permite, además, solventar problemas de seguridad que pueden derivarse del hecho de que usuarios tengan acceso como *root* a ordenadores de la red, o acceso al cable de comunicaciones, de modo que podrían interceptar contraseñas que se transmiten por la red. Esto no puede ocurrir con SSH, ya que SSH nunca envía texto plano, sino que la información siempre viaja cifrada, la diferencia entre una conexión segura con SSH y una conexión insegura se muestra en la Figura 79.

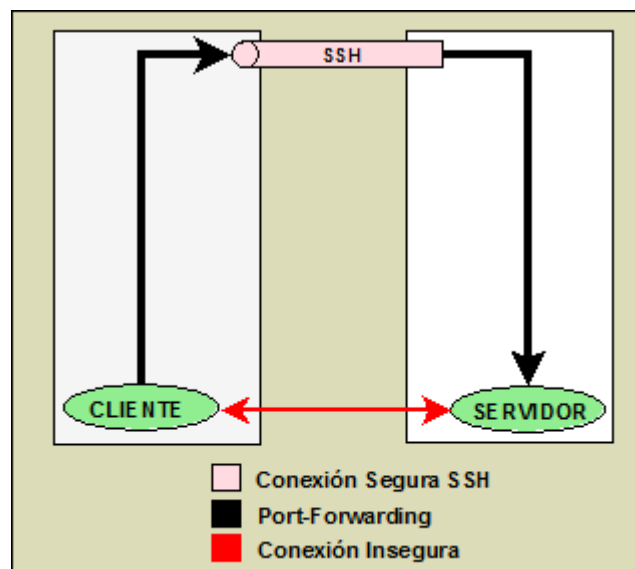


Figura 79. Transporte de información vía SSH

15.1.1 Cifrado de SSH

SSH, dependiendo de su versión utiliza diferentes algoritmos de cifrado.

- SSH1: DES, 3DES, IDEA, Blowfish
- SSH2: 3DES, Blowfish, Twofish, Arcfour, Cast128-cbc

El cifrado utilizado para cuestiones de autenticación utiliza RSA para SSH1 y DSA para SSH2.

15.1.2 Autentificación

SSH permite autenticarse utilizando uno o varios de los siguientes métodos:

- Password
- Sistema de clave pública
- Kerberos (SSH1)
- Basado en el cliente (por relaciones de confianza en SSH1 y sistema de clave pública en SSH2)

15.1.3 Protección proporcionada por SSH

SSH protege contra los siguientes tipos de ataques:

- **IP Spoofing:** Un ordenador trata de hacerse pasar por otro ordenador (en el que confiamos) y nos envía paquetes "procedentes" del mismo. SSH es incluso capaz de proteger el sistema contra un ordenador de la propia red que se hace pasar por el router de conexión con el exterior.
- **Enrutamiento de la IP de origen:** Un ordenador puede cambiar la IP de un paquete procedente de otro, para que parezca que viene desde un ordenador en el que se confía.
- **DNS Spoofing:** Un atacante compromete los registros del servicio de nombres.
- Intercepción de passwords y datos a través de la red.
- Manipulación de los datos en ordenadores intermediarios.
- Ataques basados en escuchar autenticación contra servidores X-Windows remotos.

SSH considera la red como un medio hostil en el que no se puede confiar. Un atacante, podrá forzar a que SSH se desconecte, pero no podrá descifrar o reproducir el tráfico, o introducirse en la conexión.

Todo esto, de todos modos, es sólo posible si se utiliza algún tipo de cifrado. SSH permite no usar cifrado (cifrado de tipo **"none"**), pero esta opción no debería utilizarse. Hay que tener en cuenta que SSH no sirve de nada si el atacante consigue hacerse con el control de la máquina como **root**, ya que podrá hacer que SSH quede inutilizado. Si algún atacante tiene acceso al directorio **"/home"** de un usuario de la máquina, la seguridad es ya inexistente.

15.1.4 Versiones de SSH

Existen dos versiones de SSH (SSH1 y SSH2) con diferentes funcionalidades e incompatibles entre sí (SSH2 soluciona problemas de seguridad de SSH1). A pesar de estas mejoras, existen razones a favor del uso de una y otra versión:

- SSH1 tiene fallos estructurales que le hacen vulnerable a ciertos tipos de ataques
- SSH1 puede sufrir un ataque del tipo "Man In The Middle" (intercepción de las claves públicas en el intercambio de las mismas)
- SSH1 está soportado en más plataformas que SSH2
- SSH1 permite autenticación de tipo .rhosts (propuesta en SSH2)
- SSH1 permite métodos de autenticación más diversos (AFS, Kerberos, etc.)
- SSH1 tiene mejor rendimiento que SSH2

El uso de SSH está muy extendido, y cuenta con más de 2 millones de usuarios en más de 60 países.

15.1.4.1 SSH1

Como se ha comentado anteriormente, el protocolo SSH1:

- Evita ciertos fallos de seguridad.
- Proporciona diferentes métodos de autenticación (uso de autenticación del cliente utilizando “.rhosts” junto con RSA, o autenticación utilizando sólo RSA).
- Todas las comunicaciones se cifran de manera automática y transparente. Este cifrado también se utiliza para proteger la integridad.
- El forwarding de conexiones X11 proporciona sesiones X11 seguras.
- Cualquier puerto TCP/IP puede ser redireccionado a un canal cifrado en ambos sentidos.
- El cliente autentica al servidor al comienzo de cada sesión mediante RSA, y el servidor autentica al cliente mediante RSA antes de aceptar autenticación mediante “.rhosts” o “/etc/hosts.equiv”.
- Un agente de autenticación, en la máquina del usuario local, puede utilizarse para gestionar las claves RSA del usuario.

El propósito consiste en crear un software sencillo de utilizar por los usuarios normales y a su vez, un protocolo seguro.

15.1.4.2 SSH2

SSH2 es un protocolo que permite servicios de red y conexiones seguras sobre una red insegura. Consta de tres componentes:

- Capa de protocolo de transporte: proporciona autenticación del servidor, confidencialidad e integridad y, opcionalmente, compresión. Generalmente se ejecutará sobre una conexión TCP/IP.
- Capa de protocolo de autenticación: autentica al cliente contra el servidor. Se ejecuta sobre la capa del protocolo de transporte.
- Capa del protocolo de conexión: multiplexa el canal cifrado en múltiples canales lógicos. Se ejecuta sobre la capa de autenticación.

El cliente envía una petición de servicio una vez que se ha establecido una conexión segura en la capa de transporte, y una segunda petición tras la autenticación del usuario. Esto permitiría definir nuevos protocolos que coexistieran con los anteriormente definidos.

El protocolo de conexión proporciona canales que pueden utilizarse para diferentes propósitos, como abrir sesiones interactivas, redireccionar puertos TCP/IP (tunneling) y conexiones X11.

15.2 Servidor RADIUS

RADIUS (acrónimo en inglés de *Remote Authentication Dial-In User Server*). Es un protocolo de autenticación y autorización para aplicaciones de acceso a la red o movilidad IP. Utiliza el puerto 1812 UDP para establecer sus conexiones.

Cuando se realiza la conexión con un ISP mediante módem, DSL, cable módem, Ethernet o Wi-Fi, se envía una información que generalmente es un nombre de usuario y una contraseña. Esta información se transfiere a un dispositivo Network Access Server (NAS) sobre el protocolo PPP, quien redirige la petición a un servidor RADIUS sobre el protocolo RADIUS. El servidor RADIUS comprueba que la información es correcta utilizando esquemas de autenticación como PAP, CHAP o EAP. Si es aceptado, el servidor autorizará el acceso al sistema del ISP y le asigna los recursos de red como una dirección IP, y otros parámetros como L2TP, etc.

Una de las características más importantes del protocolo RADIUS es su capacidad de manejar sesiones, notificando cuando comienza y termina una conexión, así que al usuario se le podrá determinar su consumo y facturar en consecuencia; los datos se pueden utilizar con propósitos estadísticos.

Es un protocolo ampliamente usado en el ambiente de redes, para dispositivos tales como routers, servidores y switches entre otros. Es utilizado para proveer autenticación centralizada, autorización y manejo de cuentas para redes de acceso dial-up, redes privadas virtuales (VPN) y, recientemente, para redes de acceso inalámbrico.

Puntos importantes:

- Los sistemas embebidos generalmente no pueden manejar un gran número de usuarios con información diferente de autenticación. Requiere una gran cantidad de almacenamiento.
- RADIUS facilita una administración centralizada de usuarios. Si se maneja una enorme cantidad de usuarios, continuamente cientos de ellos son agregados o eliminados a lo largo del día y la información de autenticación cambia continuamente. En este sentido, la administración centralizada de usuarios es un requerimiento operacional.
- Debido a que las plataformas en las cuales RADIUS es implementado son frecuentemente sistemas embebidos, hay oportunidades limitadas para soportar protocolos adicionales. Algún cambio al protocolo RADIUS deberá ser compatible con clientes y servidores RADIUS pre-existentes.

Un cliente RADIUS envía credenciales de usuario e información de parámetros de conexión en forma de un mensaje RADIUS al servidor. Éste autentica y autoriza la solicitud del cliente y envía de regreso un mensaje de respuesta. Los clientes RADIUS también envían mensajes de cuentas a servidores RADIUS.

Los mensajes RADIUS son enviados como mensajes UDP (User Datagram Protocol). El puerto UDP 1812 es usado para mensaje de autenticación RADIUS y, el puerto UDP 1813, es usado para mensajes de cuentas RADIUS. Algunos servidores usan el puerto UDP 1645 para mensajes de autenticación y, el puerto 1646, para mensajes de cuentas. Esto último debido a que son los puertos que se usaron inicialmente para este tipo de servicio.

15.2.1 Formato de Paquetes

Los datos entre el cliente y el servidor son intercambiados en paquetes RADIUS. Los campos en un paquete RADIUS son:

- **Code (Código):** Un octeto que contiene el tipo de paquete.

Valor	Descripción
1	Access-Request
2	Access-Accept
3	Access-Reject
4	Accounting-Request
5	Accounting-Response
11	Access-Challenge
12	Status-Server (experimental)
13	Status-Client (experimental)
255	Reserved

Cuadro 36. Tipo de paquetes del campo Code

- **Identifier (Identificador):** Un octeto que permite al cliente RADIUS relacionar una respuesta RADIUS con la solicitud adecuada.
- **Length:** Longitud del paquete (2 octetos).
- **Authenticator (Verificador):** Valor usado para autenticar la respuesta del servidor
- **RADIUS:** Es usado en el algoritmo de encubrimiento de contraseña.
- **Attributes (Atributos):** Aquí son almacenados un número arbitrario de atributos. Los únicos atributos obligatorios son el User-Name (usuario) y el User-Password (contraseña).
- **Access-Request:** Enviado por un cliente RADIUS para solicitar autenticación y autorización para conectarse a la red. Debe contener el usuario y contraseña (ya sea de usuario o CHAP); además del puerto NAS, si es necesario.
- **Access-Accept:** Enviado por un servidor RADIUS en respuesta a un mensaje de Access-Request: Informa que la conexión está autenticada y autorizada y le envía la información de configuración para comenzar a usar el servicio.
- **Access-Reject:** Enviado por un servidor RADIUS en respuesta a un mensaje de Access-Request. Este mensaje informa al cliente RADIUS que el intento de conexión ha sido rechazado. Un servidor RADIUS envía este mensaje ya sea porque las credenciales no son auténticas o por que el intento de conexión no está autorizado.
- **Access-Challenge:** Envío de un servidor RADIUS en respuesta a un mensaje de Access- Request. Este mensaje es un desafío para el cliente RADIUS. Si este tipo de paquete es soportado, el servidor pide al cliente que vuelva a enviar un paquete Access-Request para hacer la autenticación. En caso de que no sea soportado, se toma como un Access- Reject.
- **Accounting-Request:** Enviado por un cliente RADIUS para especificar información de cuenta para una conexión que fue aceptada.
- **Accounting-Response:** Enviado por un servidor RADIUS en respuesta a un mensaje de Accounting-Request. Este mensaje reconoce el procesamiento y recepción exitosa de un mensaje de Accounting-Response. La Figura 80 muestra la secuencia seguida cuando un cliente accede a la red y se desconecta de la misma.

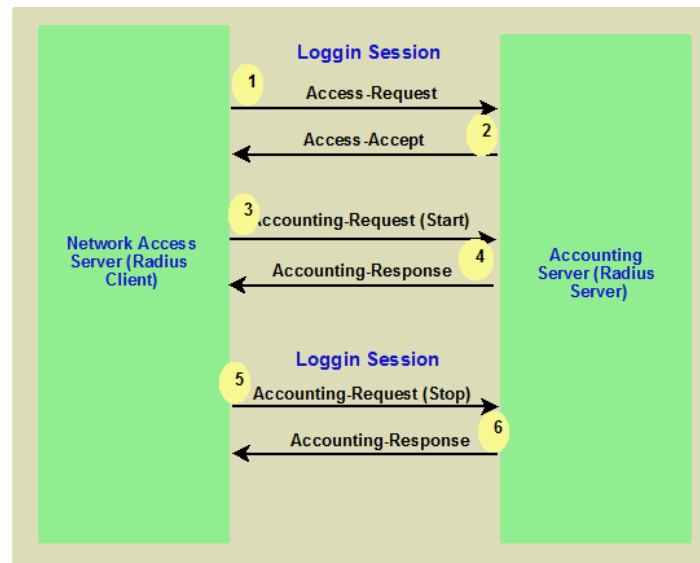


Figura 80. Secuencia de acceso a la red.

1. El cliente envía su usuario/contraseña, esta información es encriptada con una llave secreta y enviada en un Access-Request al servidor RADIUS (Fase de Autenticación).
2. Cuando la relación usuario/contraseña es correcta, entonces el servidor envía un mensaje de aceptación, Access-Accept, con información extra (Por ejemplo: dirección IP, máscara de red, tiempo de sesión permitido, etc.) (Fase de Autorización).
3. El cliente ahora envía un mensaje de Accounting-Request (Start) con la información correspondiente a su cuenta y para indicar que el usuario está reconocido dentro de la red (Fase de Accounting).
4. El servidor RADIUS responde con un mensaje Accounting-Response, cuando la información de la cuenta es almacenada.
5. Cuando el usuario ha sido identificado, éste puede acceder a los servicios proporcionados. Finalmente, cuando desee desconectarse, enviará un mensaje de Accounting-Request (Stop) con la siguiente información:
 - **Delay Time:** Tiempo que el cliente lleva tratando de enviar el mensaje.
 - **Input Octets:** Número de octetos recibido por el usuario.
 - **Output Octets:** Número de octetos enviados por el usuario.
 - **Session Time:** Número de segundos que el usuario ha estado conectado.
 - **Input Packets:** Cantidad de paquetes recibidos por el usuario.
 - **Output Packets:** Cantidad de paquetes enviados por el usuario.
 - **Reason:** Razón por la que el usuario se desconecta de la red.
6. El servidor RADIUS responde con un mensaje de Accounting-Response cuando la información de cuenta es almacenada.

15.3 IPSec

IPSec es un Grupo de Tareas de Ingeniería de Internet (IETF) conjunto estándar de protocolos que proporciona datos de **autenticación, integridad y confidencialidad** que los datos se transfieren entre los puntos de comunicación a través de redes IP. IPSec proporciona seguridad de datos a nivel de paquete IP. Un paquete es un paquete de datos que está organizado para la transmisión a través de una red, y que incluye una cabecera y la carga útil (los datos en el paquete). IPSec surgió como un estándar de seguridad de red viable porque las empresas querían asegurarse de que los datos pueden ser transmitidos con seguridad a través de Internet. IPSec protege contra posibles riesgos de seguridad mediante la protección de los datos en tránsito.

15.3.1 Características de Seguridad en IPSec

IPSec es el método más seguro disponible comercialmente para conectar sitios de la red. IPSec fue diseñado para proporcionar las siguientes funciones de seguridad al transferir paquetes a través de las redes:

- **Autenticación:** Verifica que el paquete recibido es realmente del remitente reclamado.
- **Integridad:** asegura que el contenido del paquete no se modificaron durante el transporte.
- **Confidencialidad:** Oculta el contenido del mensaje a través del cifrado.

15.3.2 Componentes de IPSec

IPSec contiene los siguientes elementos:

15.3.2.1 Carga de seguridad encapsuladora (ESP)

ESP proporciona autenticación, integridad y confidencialidad, que protege contra la alteración de datos y, sobre todo, proporcionar protección de contenido del mensaje.

ESP también proporciona los servicios de cifrado de IPSec. Cifrado traduce un mensaje legible en un formato ilegible para ocultar el contenido del mensaje. El proceso inverso, llamado descifrado, traduce el contenido del mensaje de un formato ilegible para un mensaje legible.

Cifrado/descifrado permite sólo el remitente y el receptor autorizado a leer los datos. Además, el ESP tiene una opción para realizar la autenticación, llamado autenticación ESP proporcionando autenticación y la integridad de la carga y no para el encabezado IP.

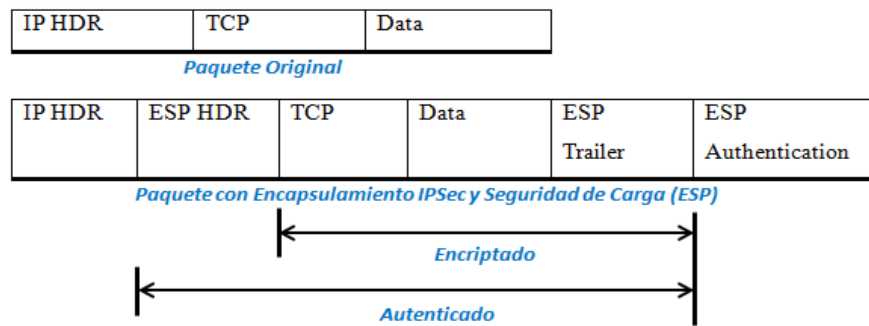


Figura 81. Carga de seguridad encapsuladora

La cabecera ESP se inserta en el paquete entre el encabezado IP y cualquier contenido de los paquetes subsiguientes. Sin embargo, debido a que ESP cifra los datos, se cambia la carga útil. ESP no cifra la cabecera ESP, ni cifra la autenticación ESP.

15.3.2.2 Encabezado de autenticación (AH)

AH proporciona autenticación e integridad, que protegen contra la alteración de datos, utilizando los mismos algoritmos como ESP. AH también proporciona protección anti-replay opcional, que protege contra la retransmisión no autorizada de paquetes. El encabezado de autenticación se inserta en el paquete entre el encabezado IP y cualquier contenido de los paquetes subsiguientes. La carga no se toca.

Aunque AH protege el origen del paquete, el destino, y el contenido de ser manipulado, la identidad del remitente y el receptor es conocida. Además, AH no protege la confidencialidad de los datos. Si se interceptan los datos y sólo AH se utiliza, el contenido del mensaje se puede leer. ESP protege la confidencialidad de los datos. Para una protección adicional en ciertos casos, AH y ESP se pueden utilizar juntos. En la Figura 82, IP HDR representa la cabecera IP e incluye tanto las direcciones IP de origen y destino.

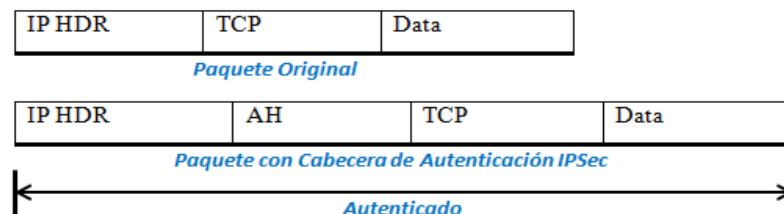


Figura 82. Encabezado de autenticación

15.3.2.3 Intercambio de claves de Internet (IKE)

Proporciona la gestión de claves y la Asociación de Seguridad (SA) de gestión. IPSec utiliza el protocolo Internet Key Exchange (IKE) para facilitar y automatizar la configuración de SA y el intercambio de claves entre las partes en la transferencia de datos. Utilización de las claves asegura que sólo el remitente y el receptor de un mensaje pueden acceder a él.

IPSec requiere que las claves pueden volver a crearse o actualizarse con frecuencia, para que las partes puedan comunicarse de forma segura con los demás.

15.3.3 Modo Operacional de IPSec

IPSec se puede utilizar en el modo de túnel o modo de transporte. Por lo general, el modo de túnel se utiliza para a-gateway protección túnel IPSec, pero el modo de transporte se utiliza para la protección del túnel IPSec host-to-host. Una pasarela es un dispositivo que monitoriza y gestiona el tráfico y las rutas de tráfico en consecuencia red entrante y saliente. Un host es un dispositivo que envía y recibe tráfico de la red.

15.3.3.1 Modo de transporte

El modo de implementación IPSec transporte encapsula sólo la carga útil del paquete. La cabecera IP no se cambia. Después de que el paquete es procesado con IPSec, el nuevo paquete IP contiene la edad de la cabecera IP (con la fuente y el destino de direcciones IP sin cambios) y la carga útil del paquete procesado. El modo de transporte no protege la información en el encabezado IP, por lo que un atacante puede obtener en el paquete viene ni a dónde va.

15.3.3.2 Modo Túnel

El modo de implementación IPSec túnel encapsula todo el paquete IP. El paquete entero se convierte en la carga útil del paquete que se procesa con IPSec. Una nueva cabecera IP se crea que contiene las dos direcciones de puerta de enlace IPSec. Las pasarelas de realizar la encapsulación/desencapsulación en nombre de los anfitriones. El modo de túnel ESP evita que un atacante analizar los datos y descifrar la misma, así como saber que el paquete es, ni a dónde va.

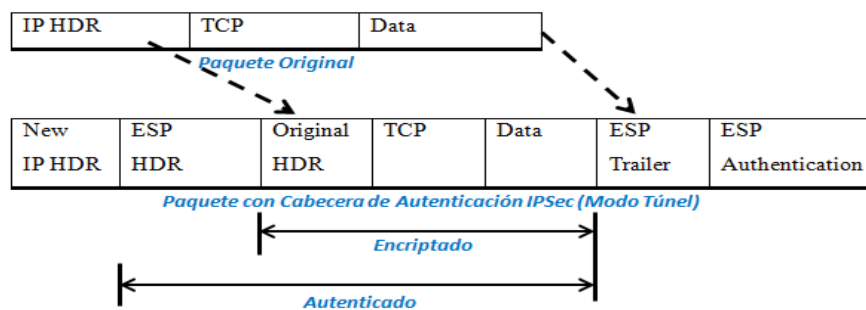


Figura 83. Operación modo Túnel

Nota: AH y ESP se puede utilizar tanto en el modo de transporte y el modo túnel

15.4 VPN (Virtual Private Network)

Ha habido muchas mejoras en la Internet, incluyendo calidad de servicio, el rendimiento de la red y las tecnologías de bajo costo. Pero uno de los avances más importantes ha sido en la red privada virtual (VPN), el protocolo de seguridad

de Internet (IPSec). IPSec es uno de los protocolos más completos, seguros y disponibles en el mercado, basados en las normas elaboradas para el transporte de datos.

Por tanto una VPN es una red compartida en la que los datos privados son segmentados a partir de otro tipo de tráfico de manera que sólo el destinatario tiene acceso. El término VPN fue originalmente utilizado para describir una conexión segura a través de Internet. Hoy, sin embargo, VPN también se utiliza para describir las redes privadas, tales como Frame Relay, modo de transferencia asíncrono (ATM) y Multiprotocol Label Switching (MPLS).

Un aspecto clave de la seguridad de datos, es que los datos que fluyen a través de la red están protegidos por tecnologías de encriptación. Las redes privadas carecen de seguridad en los datos, lo que puede permitir a los atacantes de datos conectarse directamente a la red y leer los datos. VPNs basadas en IPSec utilizan el cifrado para garantizar la seguridad de datos, lo que aumenta la resistencia de la red a la alteración o robo de datos.

VPNs basadas en IPSec se pueden crear a través de cualquier tipo de red IP, incluyendo Internet, Frame Relay, ATM y MPLS, pero sólo el Internet es ubicuo y barato.

15.4.1 Utilización de las VPN

Las VPN se utilizan tradicionalmente para:

15.4.1.1 Intranets

Intranets conectan ubicaciones de una organización. Estos lugares van desde las oficinas de la sede, a las sucursales, la casa de un empleado a distancia. A menudo, esta conectividad se utiliza para el correo electrónico y para aplicaciones de intercambio y los archivos. Mientras Frame Relay, ATM y MPLS hacen cumplir estos objetivos, las deficiencias de cada una conectividad tiene límites. El costo de conexión de los usuarios domésticos también es muy caro en comparación con las tecnologías de acceso a Internet, como DSL o cable. Debido a esto, las organizaciones se están moviendo sus redes a Internet, que es de bajo costo, y el uso de IPSec para crear estas redes.

15.4.1.2 Acceso remoto

El acceso remoto permite a los teletrabajadores y los trabajadores móviles para acceder al correo electrónico y aplicaciones empresariales. Una conexión de acceso telefónico al conjunto de módems de una organización es un método de acceso de los trabajadores a distancia, pero es caro porque la organización tiene que pagar el teléfono de larga distancia asociados y los costes del servicio.

15.4.1.3 Extranets

Son conexiones seguras entre dos o más organizaciones. Los usos más comunes para extranets incluyen la gestión de la cadena de suministro, asociaciones de desarrollo, y los servicios de suscripción. Estas empresas pueden ser difíciles de utilizar las tecnologías de red existentes, debido a los costos de conexión, retrasos, y la disponibilidad de acceso. VPNs basadas en IPSec son ideales para las conexiones de extranet. Dispositivos IPSec se pueden instalar de forma rápida y económica en las conexiones de Internet existentes.

15.4.2 Descripción general del proceso VPN

A pesar de que IPSec se basa en estándares, cada proveedor tiene su propio conjunto de términos y procedimientos para la aplicación de la norma. Debido a estas diferencias, puede ser una buena idea revisar algunos de los términos y los procesos genéricos para la conexión de dos puertas antes de sumergirse en los detalles.

15.4.2.1 Interfaces de Red y Direcciones

La puerta de enlace VPN es bien llamado, ya que funciona como un “guardián” para cada uno de los ordenadores conectados a la red local detrás de él.

En la mayoría de los casos, cada pasarela tendrá una dirección “pública” frente (lado WAN) y una dirección “privada” frente (lado LAN). Estas direcciones se refieren como la “interfaz de red”.

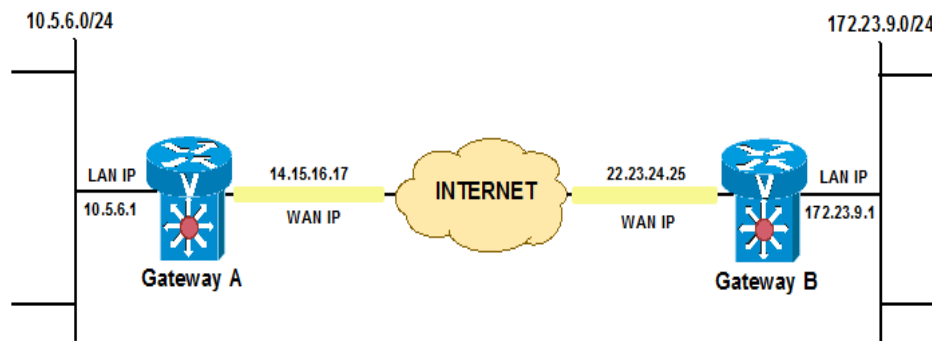


Figura 84. Puerta de enlace VPN

También es importante asegurarse de que las direcciones no se superpongan o conflicto. Es decir, cada conjunto de direcciones deben ser separadas y distintas.

15.4.3 Configuración de un túnel VPN entre puertas de enlace

Una SA, con frecuencia llamado túnel, es el conjunto de información que permite que dos entidades (redes, ordenadores, routers, firewalls, gateways) a “confiar en los demás”, y comunicarse de forma segura a medida que pasan la información a través de Internet.

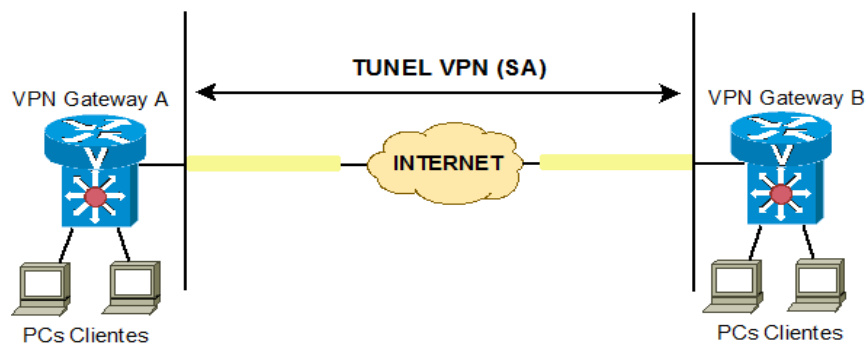


Figura 85. Túnel VPN

El SA contiene toda la información necesaria para la entrada de A, a negociar un flujo de comunicación segura y cifrada con la puerta de enlace B. Esta comunicación se refiere a menudo como “túnel”. Las puertas de enlace contienen esta información de modo que no tiene que ser cargado en todos los ordenador conectado a las puertas de enlace.

Cada puerta de enlace debe negociar su asociación de seguridad con otra puerta de enlace con los parámetros y procesos establecidos por IPSec. Como se ilustra continuación, el método más común de la realización de este proceso es a través del protocolo de Intercambio de Claves de Internet (IKE), que automatiza algunos de los procedimientos de negociación. Como alternativa, puede configurar las puertas de enlace mediante el intercambio manual de claves, lo que implica la configuración manual de cada parámetro en ambas puertas de enlace.

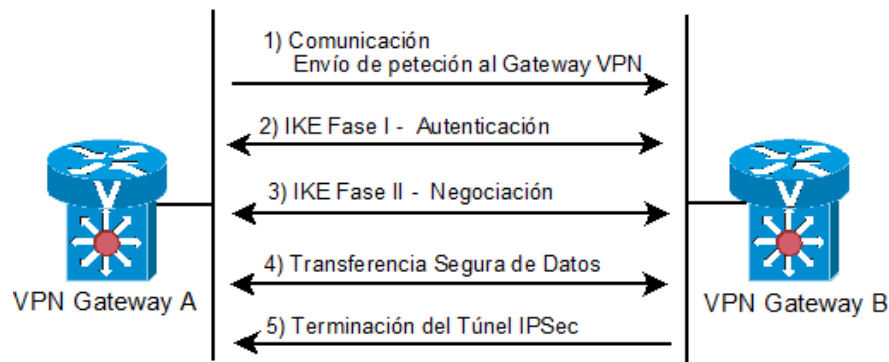


Figura 86. Negociación de asociación de seguridad

El software IPSec en el host A inicia el proceso de IPSec en un intento de comunicarse con el host B. A continuación, los dos ordenadores comienzan el proceso de intercambio de claves de Internet (IKE).

➤ **IKE Fase I**

- Las dos partes negocian el cifrado y algoritmos de autenticación a usar en la SA IKE.
- Las dos partes se autentican mutuamente mediante un mecanismo predeterminado, tal como las claves previamente compartidas o certificados digitales.
- Una llave maestra compartida es generada por el algoritmo de clave pública Diffie-Hellman en el marco IKE para las dos partes. La clave maestra se utiliza también en la segunda fase de la derivación de claves de IPSec para la SA.

➤ **IKE Fase II**

- Las dos partes negocian los algoritmos de cifrado y autenticación que se utilizará en el IPSec.
- La clave principal se utiliza para obtener las claves de IPSec para las SA. Una vez que las llaves SA se crean e intercambian, el IPSec están listos para proteger los datos del usuario entre las dos puertas de enlace VPN.

15.4.4 Transferencia de datos

Los datos se transfieren entre interlocutores IPSec basado en los parámetros de IPsec y las claves almacenadas en la base de datos SA.

Terminación del túnel IPSec. IPSec terminan por eliminación o por el tiempo de espera.

15.4.5 VPNC IKE parámetros de seguridad

Es importante recordar que ambas pasarelas deben tener los mismos parámetros establecidos para el proceso funcione correctamente.

15.4.6 VPNC IKE Fase I Parámetros

La Fase 1 de IKE parámetros utilizados:

- El modo principal
- TripleDES
- SHA-1
- Grupo 1 MODP
- Secreta pre-compartida "hr5xb84l6aa9r6"
- Vida SA de 28.800 segundos (de ocho horas)

15.4.7 VPNC parámetros IKE Fase II

Los parámetros de la Fase 2 IKE utilizados son los siguientes:

- Triple DES
- SHA-1
- El modo de túnel ESP
- Grupo 1 MODP
- Confidencialidad directa perfecta para cambio de claves
- Vida SA de 28.800 segundos (una hora)

16. SERVICIOS Y GESTIÓN DE REDES

En el mundo de las redes existen diferentes protocolos de aplicación de Internet que ofrecen maneras para que las aplicaciones se comuniquen en un entorno IP para lograr una tarea común.

Este capítulo intenta dar un estudio de los estándares de protocolos de la capa de aplicación e ilustrar cómo funcionan y como hacer uso de los servicios de transporte IP mediante el análisis de algunos de los protocolos muy comunes y más importantes.

El protocolo de Sistema de Nombre de Dominio (DNS) es un facilitador importante. Este es usado para asociar nombres de host y de dominios con las direcciones IP específicas que permite, por ejemplo, un localizador uniforme de recursos (URL), como <http://www.telematica.com> que se convierte en una dirección IP (165.193.122.135), para permitir que el tráfico IP sea enviado a su destino correcto.

El protocolo de transferencia de archivo (FTP) es usado ampliamente en Internet para copiar archivos de una estación de trabajo hacia otra. Esta transacción es orientada y enfocada en la transferencia de bloques de datos usando TCP como protocolo de transporte.

El protocolo de transferencia de hipertexto (HTTP) podría ser descrito como el protocolo que se ha hecho de Internet un éxito popular. Este protocolo es usado para publicar y leer documentos hipertextos a través de la World Wide Web (WWW), es decir que el protocolo permite publicar y leer páginas WEB.

El protocolo de transferencia simple de correo electrónico (SMTP), es un protocolo de la capa de aplicación. Este protocolo es utilizado para el intercambio de mensajes de correo electrónico entre estaciones de trabajos. Está definido en el RFC 2821 y es un estándar oficial de Internet.

En este capítulo estaremos abordando los protocolos de aplicación antes mencionados como los más usados o los más populares en la red. De igual manera en este capítulo se pretende abarcar temas de gestión de los dispositivos en la red, por lo tanto se estudiara un poco algunos el protocolo de gestión de dispositivos en la red a como protocolos que nos ayudan a gestionar y estar a la perspectiva de cualquier anomalía en la red. Algunos de estos protocolos mencionados a continuación.

El protocolo simple de administración de red (SNMP) definido en el RFC 1157, es un protocolo de capa de aplicación que facilita el intercambio de información de administración de dispositivos de red. Permite a los administradores supervisar el funcionamiento de la red, buscar y resolver sus problemas y planear su crecimiento.

El protocolo de tiempo de red (NTP) definido en el RFC 1305, es un protocolo de Internet para sincronizar los relojes de los sistemas informáticos a través del enrutamiento de paquetes en redes con latencia variable. Está diseñado para resistir los efectos de la latencia variable.

Syslog es un estándar de facto para el envío de mensajes de registro en una red informática IP. Por Syslog se conoce tanto al protocolo de red como a la aplicación o biblioteca que envía los mensajes de registro. Un mensaje de registro

suele tener información sobre la seguridad del sistema, aunque puede tener cualquier información, junto con cada mensaje se incluye la fecha y hora de envío.

De esta manera comenzaremos a estudiar, analizar y comprender cada uno de los protocolos antes mencionados, en lo que respecta tanto al funcionamiento y modo de operación de cada uno de ellos al igual que su importancia en la red.

16.1 DNS

Hasta ahora, toda la discusión de la comunicación entre los host y enrutadores en Internet se ha centrado en el uso de direcciones IP. Las direcciones IP son, sin embargo, un inconveniente para los usuarios que encuentran recordar números de hasta 12 dígitos. El sistema de nombres de dominios (**DNS RFC 1034 y 1035**) define la convenciones de nomenclatura para los host y la manera de resolver los nombres en direcciones IP.

16.1.1 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. Para poder conciliar estas dos preferencias diferenciadas, se necesita un servicio de directorio que traduzca los nombres de host a direcciones IP. Esta es la principal tarea de los sistemas de nombres de dominio (**DNS Domain Name Service**) de Internet. El DNS es:

- 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. A modo de ejemplo, considérese lo que ocurre cuando un navegador (cliente HTTP) que se ejecuta en algún host de usuario, pide el URL `www.telematica.edu/index.html`. Para que el host del usuario pueda enviar un mensaje HTTP de petición al servidor web `www.telematica.edu`, esta debe de obtener la dirección IP de `www.telematica.edu`. Esto se hace de la siguiente forma.

- La propia máquina del usuario ejecuta el lado cliente de la aplicación DNS.
- El navegador extrae el nombre de host `www.Telematica.edu` del 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.

- El navegador abre entonces una conexión TCP con el proceso HTTP servidor localizado en esa dirección IP. Como podemos ver en este ejemplo el DNS introduce un retardo adicional a las aplicaciones de Internet que utilizan DNS. Afortunadamente a como se verá posteriormente, la dirección IP deseada está a menudo depositada en un servidor de nombres DNS cercano, lo que ayuda a reducir el tráfico de red DNS y el retardo DNS medio.

DNS proporciona otros servicios importantes, además de la traducción de nombres de host a direcciones IP:

16.1.2 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 `tele.comp.enterprice.com` podría tener dos alias, como `enterprice.com` y `www.enterprice.com`. En este caso, se dice que `tele.comp.enterprice.com` se dice que es el nombre canónico. Los alias de nombres de host, de existir, son más mnemotécnicos 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.

16.1.3 Alias de servidor de correo

Por razones obvias, es recomendable que las direcciones de correo electrónico sean mnemotécnicas. Por ejemplo, si Ruddy tiene una cuenta en yahoo, su dirección de correo podría ser tan sencilla como `ruddy@yahoo.com`. Sin embargo, el nombre de host del servidor de correo de yahoo es mucho más complicado y mucho menos mnemotécnico que simplemente `yahoo.com` (por ejemplo el nombre canónico podría ser `tele.comp.yahoo.com`). 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.

16.1.4 Distribución de carga

Cada vez más, el DNS es también utilizado para realizar una distribución de carga entre servidores replicados, como los son los servidores web replicados. Por ejemplo los sitios ocupados como `cnn.com`, están replicados en múltiples servidores, cada uno de los cuales se ejecuta en un sistema final diferente, y tiene una dirección IP distinta. Por tanto, para los servidores web replicados, el nombre canónico de host está asociado a un grupo de direcciones IP. La base de datos DNS contiene este conjunto de direcciones IP. 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 rota 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.

El DNS está especificado en los RFC 1034 y 1035, es un sistema complejo y en este capítulo solo trataremos aspectos claves de su operativa. El lector interesado deberá consultar los mencionados RFC.

16.1.5 Modo de operación del DNS

Se presenta en este momento un resumen de cómo opera un DNS. La discusión se centrará en la traducción de nombres de host a direcciones IP.

Supóngase que una aplicación (navegador WEB o lector de correo) que se ejecuta en un host de usuario necesita traducir un nombre de host a dirección IP. La aplicación invocará el lado cliente del DNS, especificando el nombre de host que necesita ser traducido. Es entonces cuando el DNS en el host del usuario entra en funcionamiento, enviando un mensaje de petición a través de la red. Después de un retardo, que puede ir de 1 milisegundo a decenas de segundos, el DNS en el host del usuario recibe un mensaje DNS de respuesta que proporciona la correspondencia deseada. Esta correspondencia es entonces pasada a la aplicación.

Un diseño sencillo de DNS tendría un único servidor de nombres de Internet que contuviese todas las correspondencias. En este diseño centralizado, los clientes simplemente dirigen todas las consultas al servidor de nombres único, y el servidor de nombres responde directamente a los clientes peticionarios. A pesar del atractivo de tan simple diseño, este es inapropiado en el Internet actual, con su creciente número de estaciones. Los problemas de un diseño centralizado son:

- **Un único punto de fallo:** Si el servidor de nombres falla, también lo hace toda la Internet.
- **Volumen de tráfico:** Un único servidor de nombres tendría que gestionar todas las consultas DNS (de todas las peticiones HTTP y mensajes de correo electrónico generados por cientos de millones de estaciones).
- **Base de datos centralizada distante:** Un único servidor de nombres no puede estar cerca de todos los clientes que lo consultan. Si ponemos el servidor único en Nueva York, entonces todas las consultas desde Australia deberían viajar al otro lado del planeta, posiblemente sobre enlaces lentos y muy congestionados. Esto llevaría a significativos retardos (incrementos de las esperas globales para la web y otras aplicaciones).
- **Mantenimiento:** Un único servidor de nombres debería registrar todos los host de Internet. No solo sería gigantesca esta base de datos centralizada, sino que tendría que ser frecuentemente actualizada para dar cuenta de cada nuevo host. Existe también un problema de autenticación y autorización asociado que permite que cualquier usuario registre un host en la base de datos centralizada.

En resumen, una base de datos centralizada en un único servidor de nombres simplemente no es escalable. En consecuencia, el DNS es distribuido por diseño. De hecho, el DNS es un excelente ejemplo de cómo en Internet se puede implementar una base de datos distribuida.

Para poder gestionar el tema de la escala, el DNS utiliza un gran número de servidores de nombres organizados de forma jerárquica y distribuida alrededor del mundo. No existe ningún servidor de nombres que contenga todas las correspondencias para todos los host de Internet, sino que las correspondencias están distribuidas por los servidores de nombres.

16.1.6 Tipos de servidores DNS

En una primera aproximación, existen tres tipos de servidores de nombres, estos servidores interactúan entre sí y con los host peticionarios de la siguiente forma.

16.1.6.1 Servidores locales de nombres

Cada ISP (universidad, un departamento académico, una compañía de empleados, o un ISP residencial) tiene un servidor local de nombres (también denominado servidor de nombre por defecto). Cuando un host emite un mensaje DSN de consulta, dicho mensaje es enviado, en primer lugar, al servidor local de nombres de host. Típicamente, el servidor local de nombres está cerca del cliente. En el caso de un ISP institucional, podría estar en la misma LAN que el host cliente; para el caso de un ISP residencial, el servidor de nombres normalmente está separado del host cliente por no más de unos pocos routers. Si el host pide una traducción de otro host que es parte de la misma ISP local, entonces el servidor local de nombres podrá inmediatamente proporcionar la dirección IP pedida. De esta manera cualquier host que pida la dirección IP de otro host, el servidor local de nombres, será capaz de proporcionar la dirección IP pedida sin necesidad de contactar a ningún otro servidor de nombres.

16.1.6.2 Servidor raíz de nombres

En Internet hay alrededor de una docena de servidores raíz de nombres, la mayoría de los cuales se encuentran actualmente en Norteamérica. Cuando un servidor de nombres local no puede satisfacer la consulta de host (porque no tiene un registro para el nombre del host que se ha pedido), el servidor local se comporta como un cliente DNS, y hace la consulta a uno de los servidores raíz de nombres. Si el servidor raíz tiene un registro del nombre del host, envía un mensaje DNS de respuesta al servidor local, y entonces el servidor local envía una respuesta DNS al host peticionario. Pero el servidor raíz de nombres podría no tener el registro del nombre del host, en cuyo caso conoce la dirección IP de un servidor autorizado de nombres que conoce la correspondencia para ese nombre de host particular.

16.1.6.3 Servidores autorizados de nombres

Todo host se registra en un servidor autorizado de nombres. Típicamente, el servidor autorizado es un servidor de nombres en el ISP local del host. (En realidad, todo host está obligado a tener al menos 2 servidores autorizados de nombres para el caso de fallos). Por definición, un servidor de nombres está autorizado para un host siempre que contenga un registro DNS que permita traducir el nombre del host a su respectiva dirección IP. Cuando un servidor autorizado de nombres es consultado por un servidor raíz, el primero responde con una respuesta DNS que contiene la correspondencia pedida. Entonces el servidor raíz reenvía la respuesta al host que hizo la petición. Muchos servidores de nombres se comportan al mismo tiempo como servidores locales y como autorizados.

Veamos un ejemplo sencillo. Supóngase que el host pc1.telematica.fr quiere la dirección IP de pc2.isi.edu. Además, supóngase que el servidor local de nombres de telematica se llama dns.telematica.fr, y que el servidor autorizado de nombres para pc2.isi.edu es dns.isi.edu. Como se muestra en la Figura 87, el host pc1.telematica.fr primero envía un mensaje DNS de consulta a su servidor local de nombres (dns.telematica.fr). El mensaje de consulta contiene el nombre del host a traducir, en concreto pc2.isi.edu. El servidor local reenvía el mensaje de consulta al servidor raíz. El

servidor raíz reenvía el mensaje de consulta al servidor de nombres autorizado para todos los host en el dominio isi.edu (dns.isi.edu). El servidor autorizado envía entonces la correspondencia deseada al host que hizo la consulta a través del servidor raíz y del servidor local. Observe que en la Figura 87 se muestra que para obtener una correspondencia del nombre de un host se han enviado 6 mensajes DNS: tres mensajes de consulta y tres de respuesta.

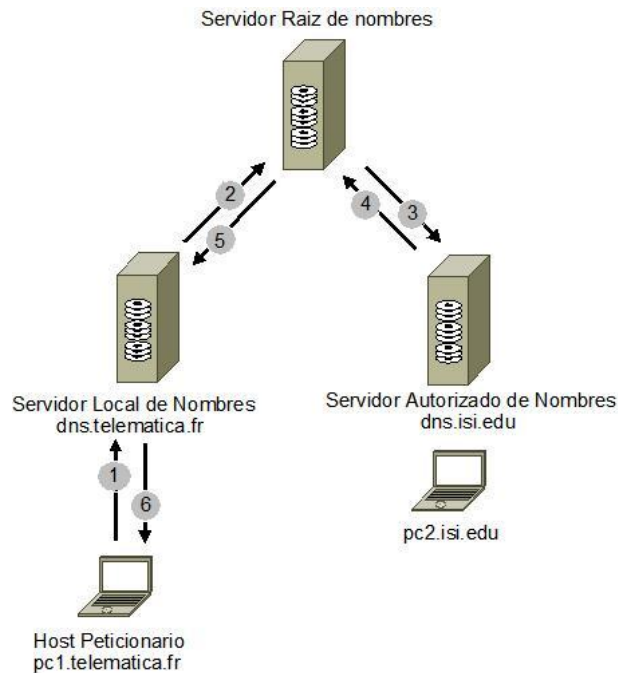


Figura 87. Consultas recursivas para obtener las correspondencias de pc2.isi.edu

La discusión precedente parte del hecho de que el servidor raíz de nombres conoce la dirección IP de un servidor autorizado para todo nombre de host. Esta suposición es incorrecta. Para un nombre de host dado, el servidor raíz puede solo conocer la dirección IP de un servidor de nombres intermedio, que, por su parte, conoce la dirección IP de un servidor autorizado para el nombre del host. Para ilustrar esto considérese una vez más el ejemplo anterior del host **pc1.telematica.fr** consultando la dirección IP de **pc2.isi.edu**. Supóngase también que la universidad de ingeniería en sistemas de información tiene un servidor de nombres para la universidad llamado **dns.isi.edu**. Supóngase ahora que cada uno de los departamentos de la universidad tiene su propio servidor de nombres, y que cada uno de los servidores de nombres departamentales está autorizado para cada uno de los host del departamento. Como se muestra en la Figura 88, cuando el servidor raíz de nombres recibe una consulta para un host cuyo nombre termina en **isi.edu**, reenvía la consulta al servidor de nombres llamado **dns.isi.edu**. Este servidor de nombres reenvía todas las consultas para nombres de host terminados en **.comp.isi.edu**. El servidor autorizado envía la correspondencia deseada al servidor de nombres intermedio, **dns.isi.edu**, que lo reenvía al servidor raíz, que lo reenvía al servidor local, **dns.telematica.fr**, que lo reenvía al host peticionario.

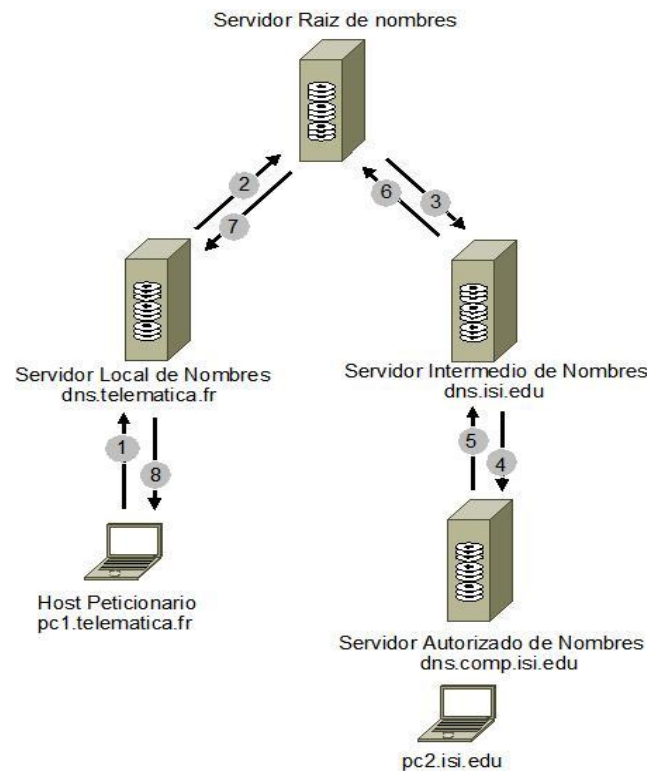


Figura 88. Consulta recursiva con un servidor de nombres intermedio

En el ejemplo anterior se realizan 8 mensajes DNS. Realmente pueden haber muchos más servidores intermediarios entre el servidor raíz y el servidor autorizado.

En los ejemplos vistos anteriormente, se asume que todas las consultas son de las denominadas **recursivas**. Se dice que una consulta es recursiva cuando un host **A** hace una consulta hacia un host **B**, y este obtiene la correspondencia pedida en representación de **A**, y después reenvía esta correspondencia a **A**. El protocolo DNS también permite realizar consultas **iterativas** en cualquier paso de la cadena entre el host petionario y el servidor autorizado de nombres.

Se dice que una consulta es iterativa cuando un servidor de nombres **A** hace una consulta a un servidor de nombres **B**, si el servidor **B** no dispone de la correspondencia pedida, inmediatamente envía a **A** una respuesta DNS que contienen la dirección IP del siguiente servidor de nombres en la cadena (Por ejemplo un servidor **C**). El servidor **A** entonces envía una consulta directamente al servidor **C**.

En la Figura 89, se muestra una combinación de consultas **recursivas** e **iterativas**. Típicamente todas las consultas de la cadena son recursivas, excepto la consulta desde el servidor local al servidor raíz (debido a que los servidores

raíces manejan grandes cantidades de consultas, es preferible usar consultas **iterativas**), que es iterativa.

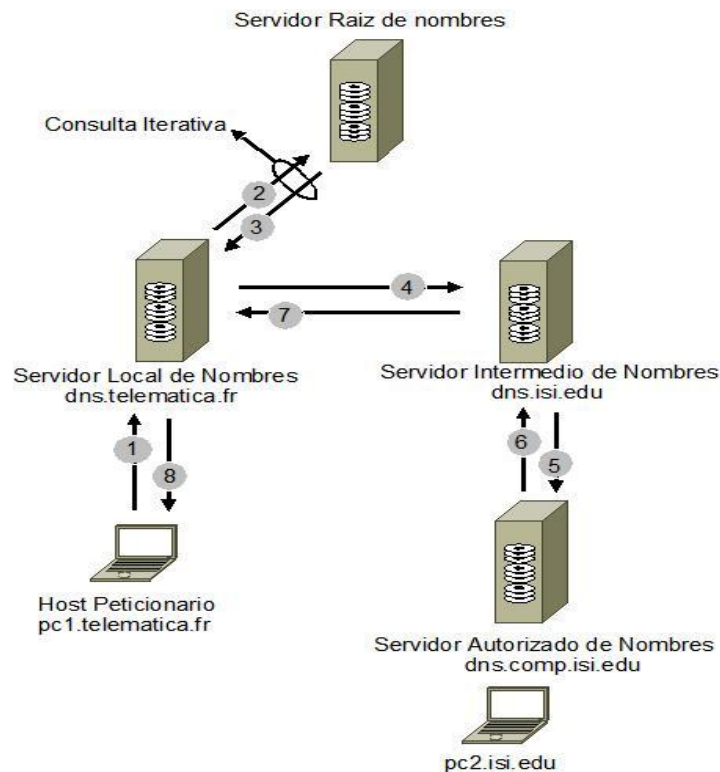


Figura 89. Combinación de consultas recursivas e iterativas.

La discusión mantenida hasta el momento no ha tratado una importante característica: la **cache DNS**. El protocolo DNS utiliza esta característica para mejorar el rendimiento de retardos y reducir el número de mensajes DNS en la red.

La idea de esta característica es muy simple. Cuando un servidor de nombres recibe una correspondencia DNS para algún nombre de host, hace una copia cache en memoria local (disco o RAM). Un registro de cache es desestimado después de un periodo de tiempo (normalmente establecido en 2 días).

16.1.7 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. En esta sección y en la siguiente se proporciona un breve resumen de los registros de recursos y los mensajes DNS. Para mayor información abocarse con los **RFC 1034, 1035, y el libro Redes de computadores un enfoque descendente**.

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:

- **Tipo A:** Si tipo es A, entonces nombre es un nombre de host, y valor es la dirección IP del host. A modo de ejemplo (**pc1.comp.telematica.com, 210.1.1.1, A**).
- **Tipo NS:** Si tipo es NS, entonces nombre es un dominio (**telematica.com**), y valor es el nombre de host de un servidor autorizado, que sabe cómo obtener las direcciones IP de host en el dominio. A modo de ejemplo tenemos (**telematica.com, dns.telematica.com, NS**).
- **Tipo CNAME:** Entonces el valor es un nombre canónico para el alias nombre. A manera de ejemplo tenemos (**telematica.com, pc1.comp.telematica.com, CNAME**).
- **Tipo MX:** Entonces valor es el nombre canónico de un servidor de correo que tiene el alias nombre. A manera de ejemplo tenemos que (**telematica.com, mail.bar.telematica.com, MX**).
- **Mensajes DNS**

En las secciones anteriores de hicieron alusiones a los mensajes DNS de consulta y respuesta. Estos son los dos únicos mensajes DNS que existen, además ambos mensajes tienen el mismo formato a como se muestra en la [Figura 90](#).

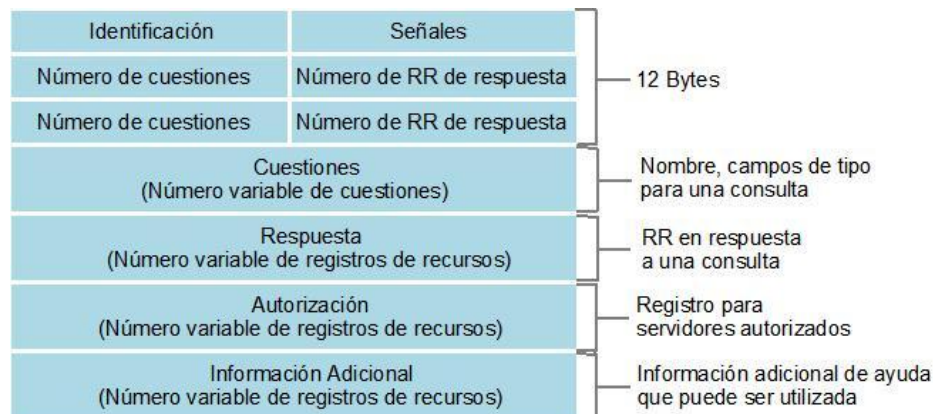


Figura 90. Formato de mensaje DNS.

Para una mayor información de los respectivos campos de este mensaje, consulte los RFC 1034, 1035 y el libro de redes de computadores un enfoque descendente.

16.2 HTTP

El protocolo de transferencia de hipertexto (HTTP), protocolo de la capa de aplicación de la WEB, está en el corazón de la web. Está definido en el RFC 1954 y 2616. HTTP esta implementado en 2 programas: un programa cliente y otro servidor, estos a su vez se ejecutan en sistemas finales diferentes. HTTP define la estructura de estos mensajes y como se realiza el intercambio entre el cliente y el servidor.

A lo largo de 1997, básicamente todos los navegadores y servidores web implementaban la versión **HTTP/1.0**, definida en el **RFC 1945**. A principios de 1998, algunos navegadores y servidores web comenzaron a implementar la versión **HTTP/1.1**, definida en el **RFC 2616**. HTTP/1.0 es compatible con la versión anterior **HTTP/1.0**, de manera que un navegador que ejecuta HTTP/1.0 puede hablar con un servidor que ejecuta HTTP/1.1 y viceversa.

En la Figura 91 se ilustra una conversación simple de dos clientes HTTP con un mismo servidor.

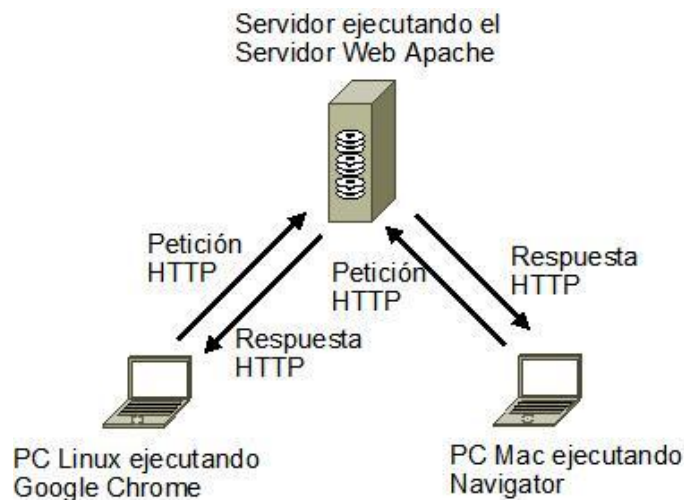


Figura 91. Comportamiento de petición-respuesta de HTTP

Tanto HTTP/1.0 como HTTP/1.1 utilizan **TCP** como protocolo de transporte (no utilizan **UDP**). En primer lugar, el cliente HTTP inicia una conexión TCP con el servidor. Una vez establecida la conexión, los procesos del navegador y el servidor acceden a TCP por medio de sus interfaces de socket. Esto es, en el lado del cliente la interfaz de socket es la puerta entre el proceso cliente y la conexión TCP; en el lado del servidor es la puerta entre el proceso servidor y la conexión TCP. Una vez que el cliente envía un mensaje a su interfaz de socket, éste está fuera de las manos del cliente, y está en las manos de TCP.

TCP proporciona a HTTP un servicio fiable de transferencia de datos. Esto implica que cada mensaje HTTP de petición emitido por el proceso cliente, eventualmente llega sin modificaciones al servidor; de igual manera, cada mensaje HTTP de respuesta emitido por el proceso del servidor, eventualmente llega intacto al cliente.

Es importante observar que el servidor envía los archivos sin guardar información de ningún cliente. Si un cierto cliente pide 2 veces el mismo objeto en un periodo de tiempo de unos pocos segundos, el servidor no responde que acaba de servir el objeto al cliente, sino que reenvía el objeto, ya que ha olvidado por completo lo que hizo anteriormente. Se dice que HTTP es un **protocolo sin estado**, esto debido a que el servidor no guarda información de los clientes.

16.2.1 Tipos de conexiones en HTTP

HTTP puede utilizar tanto conexiones no persistentes como persistentes. HTTP/1.0 utiliza conexiones no persistentes. HTTP/1.1 utiliza conexiones persistentes por defecto, sin embargo, los clientes y servidores HTTP/1.1 pueden ser configurados para utilizar conexiones no persistentes. En breve veremos con más detalles cada uno de estos métodos.

16.2.1.1 Conexiones no Persistentes

En este método de operación veremos que solo se puede transferir un único objeto web sobre una conexión TCP. Veamos un pequeño ejemplo de la transferencia de una página web desde el servidor al cliente para el caso de conexiones no persistentes. Supongamos que la página consta de un archivo HTML base y 10 imágenes JPEG, y que los 11 objetos residen en el mismo servidor. Sea la siguiente URL del archivo HTML base.

www.escuela.edu/departamento/home.index

- El cliente HTTP inicia la conexión TCP con el servidor **www.escuela.edu** sobre el número de puerto 80, que es el puerto por defecto para HTTP.
- El cliente HTTP envía al servidor un mensaje HTTP de petición a través del socket asociado con la conexión TCP establecida en el paso anterior. El mensaje de petición incluye el nombre de la ruta **/departamento/home.index**
- El servidor HTTP recibe el mensaje de petición a través del socket asociado con la conexión, recupera el objeto **/departamento/home.index** de su almacenamiento (disco, RAM), encapsula el objeto en el mensaje HTTP de respuesta, y envía el mensaje de respuesta al cliente a través de socket.
- El servidor HTTP dice q TCP que cierre la conexión (TCP realmente no cierra la conexión hasta que sepa con seguridad que el cliente ha recibido intacto el mensaje de respuesta).
- El cliente HTTP recibe el mensaje de respuesta. Finaliza la conexión TCP. EL mensaje indica que el objeto encapsulado es un archivo HTML. El cliente extrae el archivo del mensaje de respuesta, examina el archivo HTML, y encuentra la referencia a los diez objetos JPEG.
- Se repiten los cuatro primeros pasos para cada uno de los objetos JPEG referenciados.

Los pasos anteriores utilizan conexiones no persistentes, porque cada una de las conexiones TCP es cerrada después de que el servidor envíe el objeto (la conexión no persiste para los otros objetos). Por tanto, en este ejemplo, cuando un usuario demanda una página web, se generan 11 conexiones TCP.

Antes de continuar hagamos un cálculo estimado del tiempo que transcurre desde que el cliente pide el archivo HTML base hasta que éste es recibido completamente por el cliente. Para este propósito, se define el tiempo de ida y vuelta (RTT), que es el tiempo que tarda un pequeño paquete en ir desde el cliente hasta el servidor y después de vuelta al cliente. El RTT incluye los retardos de propagación de paquetes, los de las colas de los routers y en los switches

intermedios, y los de procesamiento de paquetes. A continuación se considera que es lo que ocurre cuando un usuario pulsa un hipervínculo. Como se muestra en la Figura 92, en donde el navegador inicia una conexión TCP entre el navegador y el servidor, lo que implica un acuerdo de tres vías.

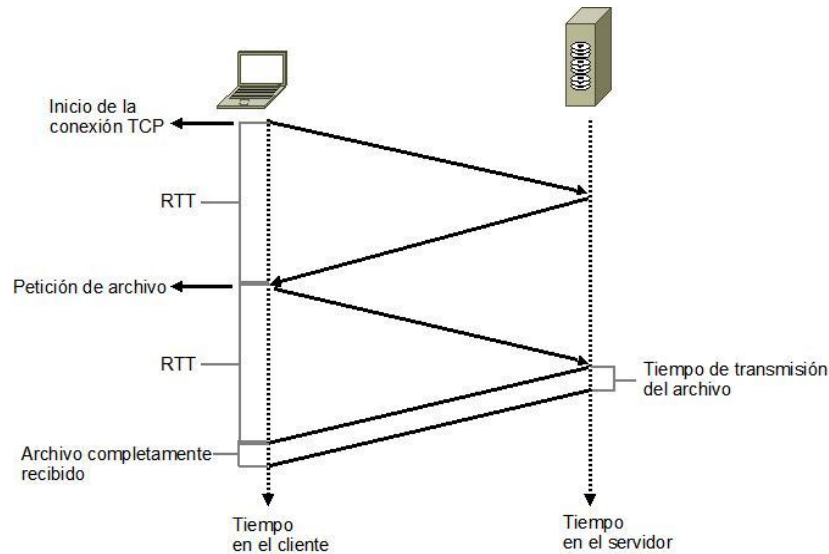


Figura 92. Cálculo estimado de la petición de un archivo HTML

Exponiendo un poco el ejemplo ilustrado en la Figura 92, se dice que el cliente HTTP envía un segmento TCP al servidor, el servidor reconoce y responde con un segmento TCP, y finalmente el cliente hace un reconocimiento de vuelta al servidor. Transcurre un RTT tras las dos primeras partes del acuerdo de tres vías. Después de completar las dos primeras partes del acuerdo, el cliente envía un mensaje HTTP de petición en combinación con la tercera parte del acuerdo (el reconocimiento) sobre la conexión TCP. Una vez que el mensaje de petición llega al servidor, éste envía el archivo HTML sobre la conexión TCP. Esta petición/respuesta HTTP consume otro RTT. Por tanto, aproximadamente, el tiempo total de respuesta es de dos RTT más el tiempo de transmisión del archivo HTML en el servidor.

16.2.1.2 Conexiones Persistentes

Las conexiones no persistentes presentan muchos defectos. En primer lugar se debe de mantener y establecer una conexión para cada uno de los objetos pedidos. Para cada una de estas conexiones se deben reservar búferes y variable TCP, que deben ser almacenados en el cliente como en el servidor.

Esto supone una carga importante sobre el servidor web, que puede estar sirviendo simultáneamente a cientos de clientes distintos. En segundo lugar, como se acaba de describir, cada objeto sufre un retardo de entrega de 2 RTT: uno para establecer la conexión TCP, y el otro para la petición y recepción del objeto. Con las conexiones persistentes, el servidor deja abierta la conexión TCP después de enviar una respuesta, de manera que las subsiguientes peticiones y respuesta entre el mismo cliente y servidor pueden ser enviadas sobre la misma conexión.

Existen dos versiones de conexiones persistentes: **con entubamiento** y **sin entubamiento**. Para la versión sin entubamiento, el cliente solo emite una nueva petición cuando la respuesta previa ha sido recibida, en este caso el cliente sufre de un RTT por cada uno de los objetos referenciados que son pedidos y recibidos. Otra desventaja de no hacer entubamiento es que después de que el servidor envía un objeto sobre la conexión TCP persistente, la conexión no hace nada más que esperar a que llegue otra conexión. Esta espera desaprovecha el recurso del servidor.

El modo por defecto de HTTP/1.1 utiliza conexiones persistentes con entubamiento. Con entubamiento, el cliente HTTP realiza una petición tan pronto encuentra una referencia. Por tanto, el cliente HTTP puede realizar peticiones seguidas para los objetos referenciados; esto es, puede realizar una nueva petición antes de recibir una respuesta para la petición anterior. Con entubamiento es posible que solo se emplee un único RTT para todos los objetos referenciados. La conexión TCP con entubamiento permanece inactiva durante una fracción de tiempo menor.

16.2.2 Formato del mensaje HTTP

Existen dos tipos de mensajes HTTP: los mensajes de petición, y los mensajes de respuesta; ambos se describen a continuación.

16.2.2.1 Mensaje HTTP de petición

```
GET /dir/pagina.html HTTP/1.1
Host: www.escuela.edu
Conection: close
User-agent: Mozilla/4.0
Accept-language:fr
```

A continuación se plantea un típico mensaje HTTP de petición.

En el ejemplo anterior, podemos observar que el mensaje consta de cinco líneas. Un mensaje de petición puede tener muchas más líneas o tan solo una línea. La primera línea del mensaje HTTP de petición se denomina **línea de petición**; las siguientes se denominan líneas de cabeceras. La línea de petición consta de tres campos: el campo de método, el campo de URL, y el campo de la versión HTTP. El campo del método puede tomar distintos valores, entre los que se encuentran **GET**, **POST**, y **HEAD**. La gran mayoría de mensajes HTTP de petición utilizan GET. El método GET es utilizado cuando el navegador demanda un objeto que identifica en el campo URL, en este ejemplo el navegador pide el objeto /dir/pagina.html. El campo versión indica la versión HTTP que está implementando el navegador.

Veamos ahora las líneas de cabecera del ejemplo. La línea de cabecera: **Host: www.escuela.edu** especifica el host en el que reside en objeto. La línea de cabecera: **Conection: close**, el navegador está indicando al servidor que no quiere utilizar conexiones persistentes; quiere que el servidor cierre la conexión después de enviar el objeto pedido. La línea de cabecera: **User-agent** especifica el agente de usuario, esto es. El tipo de navegador que está haciendo la petición al servidor. La cabecera **Accept-language** indica que el usuario prefiere recibir una versión en francés del objeto, si es que este objeto existe en el servidor.

Una vez visto un ejemplo, pasaremos a ver el formato general del mensaje de petición. El formato general que estaremos viendo a continuación es muy similar a nuestro ejemplo anterior. Sin embargo, podemos observar que tras las líneas de cabecera existe un **cuerpo de entidad**. El cuerpo de entidad está vacío con el método GET, aunque se utiliza con el método POST. A menudo, un cliente HTTP utiliza el método POST cuando el usuario rellena un formulario. Si el valor del campo método es POST, entonces el cuerpo de entidad contiene lo que el usuario introdujo en los campos del formulario.

El método **HEAD** es similar al método GET. Cuando un servidor recibe una petición con el método HEAD, responde con un mensaje HTTP, pero que no incluye en objeto pedido. Este método es utilizado a menudo por desarrolladores de aplicaciones para depuraciones.

Las especificaciones HTTP/1.0 solo permite 3 tipos de métodos, que son: GET, POST y HEAD. Además de estos 3 métodos HTTP/1.1 permite otros métodos adicionales, como **PUT** y **DELETE**. El método PUT se utiliza habitualmente con herramientas de publicación Web. Permite a un usuario cargar un objeto en una ruta específica en un servidor. El método **DELETE** permite que un usuario borre un objeto de un servidor Web.

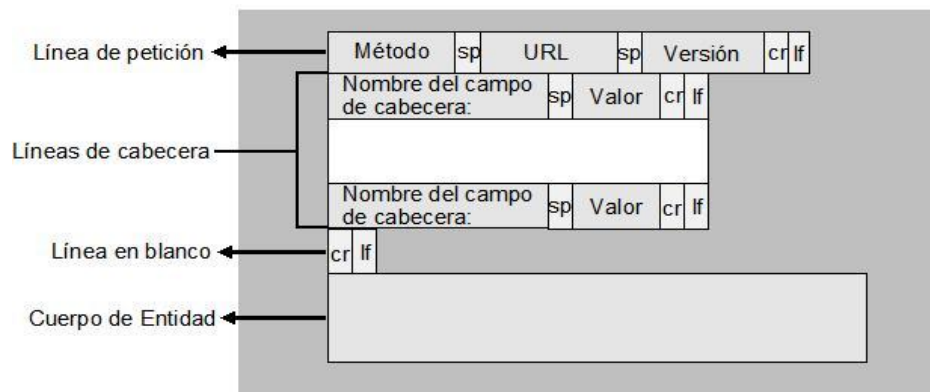


Figura 93. Formato general de un mensaje de petición

16.2.2.2 Mensaje HTTP de respuesta

A continuación, se proporciona un mensaje HTTP típico de respuesta. Este mensaje puede ser la respuesta al ejemplo de petición explicado antes.

```
HTTP/1.1 200 ok
Connection: close
Date: thu, 06 Aug 1998 12:00:15 GMT
Server: Apache/1.3.0 (Unix)
Last-Modified: Mon, 22 Jun 1988 09:23:24 GMT
Content-Length: 6821
Content-Type: text/html

(datos datos datos datos ...)
```

Revisemos cuidadosamente este mensaje de respuesta que tiene tres secciones: una **línea inicial de estatus**, seis **líneas de cabecera**, y luego el **cuerpo de entidad**. El cuerpo de entidad es la esencia del mensaje: contiene el propio

objeto pedido. La línea de estatus tienen tres campos: el campo de versión del protocolo, el código de estatus, y el correspondiente mensaje de estatus. En el ejemplo anterior la línea de estado indica que el servidor está utilizando HTTP/1.1 y que todo está OK. El servidor utiliza la línea de cabecera **Connection: close** para indicar al cliente que va a cerrar la conexión TCP después de enviar el mensaje. La línea de cabecera **Date:** indica la fecha y hora en que se creó y envió la respuesta HTTP por parte del servidor. La línea de cabecera **server** indica que el mensaje fue creado por un servidor web apache. La línea **Last-Modified:** indica la fecha y hora que el objeto fue creado o modificado por última vez. La línea de cabecera **Content-Length** indica el número de bytes del objeto enviado. La línea **Content-Type** indica que el objeto en el cuerpo de entidad es texto HTML.

Una vez que se ha descrito el ejemplo anterior, examinaremos el formato general de un mensaje de respuesta HTTP, a como se muestra en la Figura 94. El formato general del mensaje de respuesta se corresponde con el ejemplo anterior de un mensaje de respuesta.

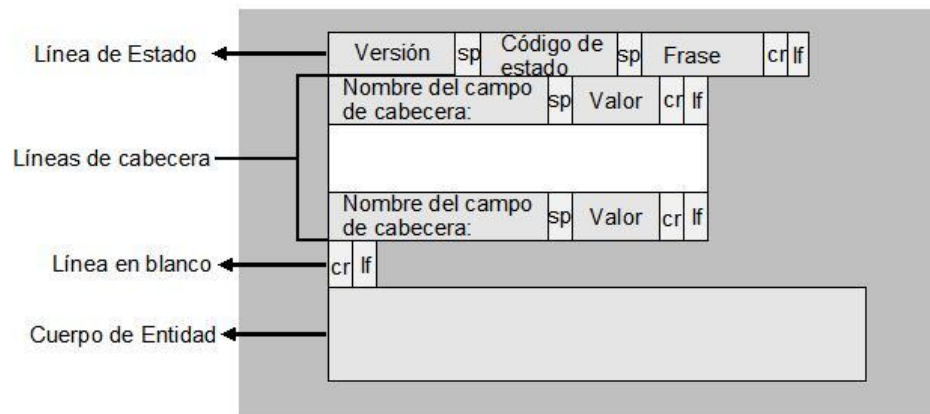


Figura 94. Formato general de un mensaje de respuesta

16.3 FTP

En una sesión FTP típica, el usuario está sentado enfrente de un host y quiere transferir archivos a/desde un host remoto. Para que el usuario pueda acceder a la cuenta remota debe de proporcionar una identificación de usuario y una palabra clave. Después de proporcionar esta información de autorización, el usuario puede transferir archivos desde el sistema local al sistema remoto y viceversa. El usuario interactúa con FTP a través de un agente de usuario FTP. Primeramente el usuario proporciona el nombre del host remoto, haciendo que el proceso FTP cliente establezca una conexión TCP con el proceso FTP servidor. El usuario proporciona entonces la identificación de usuario y la palabra clave, que son enviados sobre la conexión TCP como parte de los comandos FTP.

Tanto HTTP como FTP son protocolos de transferencia de archivos que tienen características comunes; por ejemplo, ambos funcionan sobre TCP. Sin embargo, los dos protocolos de la capa de aplicación muestran importantes diferencias. La diferencia más notable consiste en que FTP utiliza dos conexiones TCP paralelas para transferir un archivo: una **conexión de control** y una **conexión de datos**. La conexión de control se utiliza para enviar información de control entre las dos estaciones como: identificación del usuario, palabra clave, comandos para modificar el

directorio remoto y comandos para depositar y obtener archivos. La conexión de datos se utiliza para enviar un archivo. Dado que FTP utiliza una conexión de control separada, se dice que envía su información de control **fuera de banda**.

Como recordara, HTTP envía las líneas de cabecera de petición y respuesta sobre la misma conexión TCP que transporta el propio archivo transferido. Por este motivo, se dice que HTTP al igual que SMTP envía su información de control **en banda**.

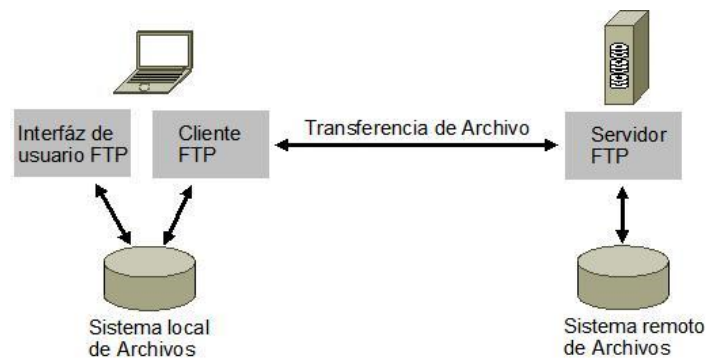


Figura 95. Transferencia de archivos entre sistemas locales y remotos

Cuando un usuario da comienzo a una sesión FTP con un host remoto, el lado del cliente de FTP inicia primeramente una conexión TCP de control con el lado del servidor sobre el puerto del servidor número 21. El lado del cliente de FTP envía una identificación de usuario a través de esta conexión de control. El lado del cliente FTP también envía sobre la conexión de control comandos que cambian el directorio remoto. Cuando el lado del servidor recibe sobre la conexión de control un comando de transferencia de archivo, el lado del servidor inicia una conexión TCP de datos con el cliente. FTP envía exactamente un archivo sobre la conexión de datos, y después cierra la conexión de datos. Por tanto, con FTP la conexión de control permanece abierta mientras dure la sesión del usuario, de manera que las conexiones de datos son no persistentes.

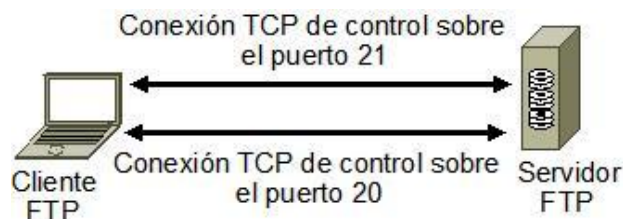


Figura 96. Conexiones de control y de datos

16.3.1 Comandos comunes FTP

Finalizamos esta sección con una breve descripción de algunos de los comandos FTP más comunes. A continuación se dan algunos de los comandos más habituales.

16.3.1.1 USER: Nombre de usuario

Utilizado para enviar la identificación de usuario al servidor.

16.3.1.2 PASS: Palabra clave

Utilizado para enviar la palabra clave al servidor.

16.3.1.3 LIST

Utilizado para pedir al servidor que liste todos los archivos del directorio remoto en curso.

16.3.1.4 GET

Utilizado para recuperar un archivo del directorio actual del servidor.

16.3.1.5 PUT

Utilizado para almacenar un archivo en el directorio actual del servidor.

16.4 Correo Electrónico

El correo electrónico fue la aplicación más popular cuando Internet estaba en sus comienzos, y es una de las aplicaciones más importantes de Internet hasta la fecha. En esta sección se examinarán los protocolos de la capa de aplicación, que constituyen el corazón del correo electrónico en Internet. Antes de pasar a una descripción a profundidad de estos protocolos, tomemos un punto de vista de alto nivel del sistema de correo de Internet y de sus componentes claves.

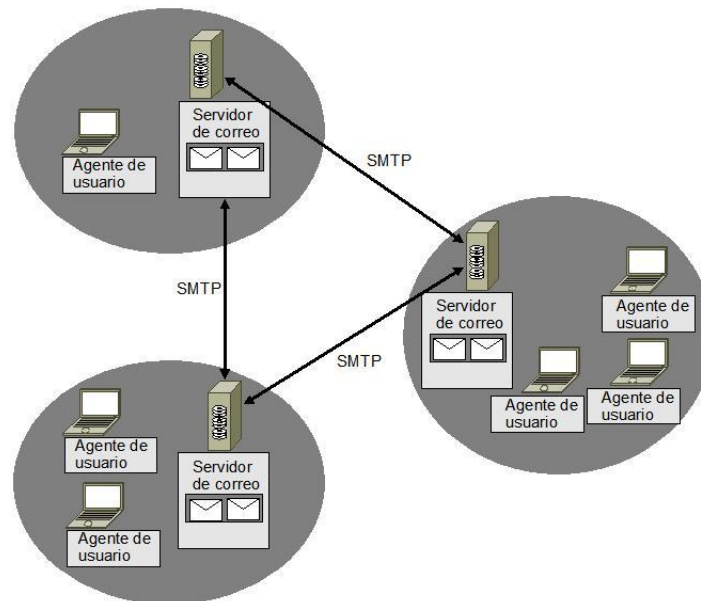


Figura 97. Perspectiva de alto nivel del sistema de correo electrónico.

La Figura 97 presenta una visión de alto nivel del sistema de correo de Internet y de sus componentes claves. Vemos en el diagrama que tiene tres componentes claves que son: agentes de usuario, servidores de correo, y el protocolo simple de transferencia de correo (SMTP). Los servidores de correo forman el núcleo de la infraestructura del correo electrónico. Cada destinatario tiene un buzón de correo ubicado en uno de los servidores de correo.

16.4.1 Protocolo de correo electrónico (SMTP)

El protocolo simple de transferencia de correo (SMTP) es el principal protocolo de la capa de aplicación para el correo electrónico de Internet. Utiliza el servicio fiable TCP de transferencia de datos para enviar el correo desde el servidor de correo del remitente hasta el del destinatario.

Como la mayoría de los protocolos de la capa de aplicación, SMTP tiene dos lados: el lado cliente, que ejecuta el servidor de correo del remitente, y el lado servidor, que ejecuta el servidor de correo del destinatario.

Ambos lados, cliente y servidor, de SMTP están presentes en todos los servidores de correo

SMTP es el núcleo del correo electrónico de Internet. SMTP transfiere mensajes desde el servidor de correo de un remitente al de un destinatario.

Para ilustrar la operativa básica de SMTP, revisemos un escenario común. Supongamos que Ruddy quiere enviar a Yoel un sencillo mensaje ASCII.

- Ruddy invoca a su agente de usuario para correo electrónico, proporciona la dirección de correo electrónico de Yoel (por ejemplo **yoel@telematica.edu**), compone un mensaje, y ordena al agente de usuario que envíe el mensaje.
- El agente de usuario de Ruddy envía el mensaje a su servidor de correo, donde es ubicado en una cola de mensaje.
- El lado cliente de SMTP, ejecutándose en el servidor de correo de Ruddy, ve el mensaje en la cola de mensajes, abre una conexión TCP con un servidor SMTP, que está ejecutándose en el servidor de correo de Yoel.
- Tras alguna sincronización SMTP inicial, el cliente SMTP envía el mensaje de Ruddy sobre la conexión TCP.
- En el servidor de correo de Yoel, el lado del servidor de SMTP recibe el mensaje. Entonces, el servidor de correo de Yoel lo deposita en su buzón de correo.
- Yoel invoca a su agente de usuario para leer convenientemente el mensaje.

El escenario de resume en la Figura 98.

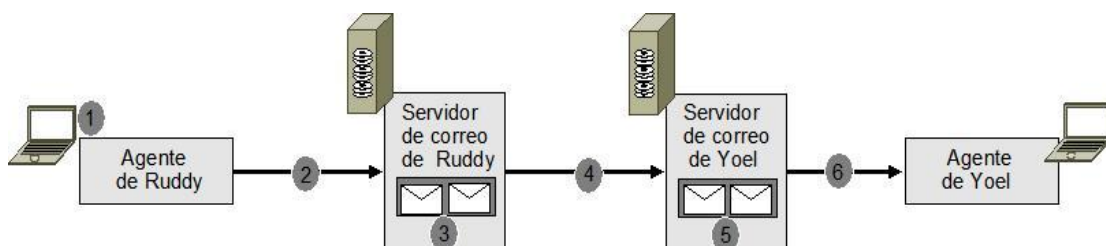


Figura 98. Ruddy envía un mensaje a Yoel

Veamos más cerca como transfiere SMTP un mensaje desde un servidor de correo emisor a otro receptor. En primer lugar, el cliente SMTP hace que TCP establezca una conexión sobre el **puerto 25** con el servidor SMTP receptor. Una vez que se ha establecido la conexión, el servidor y el cliente realizan alguna sincronización de la capa de aplicación.

Durante esta fase SMTP de sincronismo, el cliente SMTP indica la dirección de correo electrónico del remitente y del destinatario. Una vez que el cliente y el servidor SMTP se han presentado, el cliente envía el mensaje. El cliente repite este proceso sobre la misma conexión TCP si tiene que enviar otros mensajes al servidor, en caso contrario, ordena a TCP cerrar la conexión.

16.4.2 Formato de mensaje de correo

Cuando Ruddy escribe a Yoel una carta por correo ordinario, puede que incluya todo tipo de información periférica de cabecera al comienzo de la carta, como la dirección de Yoel, su propia dirección de repuesta y la fecha. De forma similar, cuando una persona envía un correo electrónico a otra, se incluye una cabecera que contiene información periférica y que precede al propio cuerpo del mensaje. Las líneas de cabecera y el cuerpo del mensaje se separan por una línea en blanco. Cada cabecera debe incluir una línea From: y To:, pudiendo incluir una línea de cabecera Subject: junto con otras líneas opcionales.

```
From: ruddy@telematica.edu
To: yoel@isi.edu
Subject: Buscando un nuevo horizonte
```

Figura 99. Cabecera típica de un mensaje SMTP

Después de la cabecera del mensaje sigue la línea en blanco, y después el cuerpo del mensaje

16.4.3 Extensión MIME para datos no ASCII

Las cabeceras definidas en el RFC 822 son adecuadas para enviar texto ASCII ordinario. Sin embargo, no son lo suficiente solventes para los mensajes multimedia, o para transportar formato de texto no ASCII. Para enviar otros contenidos distintos al texto ASCII, el agente de usuario emisor debe de incluir en el mensaje cabeceras adicionales. Estas cabeceras extras son las extensiones MIME (Multipurpose Internet Mail Extensions).

Las dos cabeceras MIME claves para soporte multimedia son Content-Type y Content-Transfer-Encoding. La cabecera Content-Type, permite que el agente de usuario receptor tome las acciones oportunas en relación con el mensaje, por ejemplo, indicando que el cuerpo del mensaje contiene una figura JPEG, el agente de usuario receptor puede dirigir el cuerpo del mensaje a una rutina del mensaje de descompresión de JPEG.

La cabecera Content-Transfer-Encoding, es necesaria porque los mensajes no ASCII deben ser codificados a un formato ASCII que no cause confusión a SMTP. Esta cabecera avisa al agente de usuario receptor de que el cuerpo del mensaje ha sido codificado a ASCII, e indica el tipo de codificación empleada.

16.4.4 Protocolos de acceso al correo

En esta sección se discutirán los protocolos más importantes de acceso o recuperación de correo desde un servidor de correo a un agente de usuario. Entre los cuales tenemos:

16.4.4.1 POP3

POP3 que se define en el RFC 1939, es un protocolo de acceso al correo extremadamente simple y limitado. POP3 comienza cuando el agente de usuario (cliente) abre una conexión TCP con el servidor de correo (el servidor) sobre el **puerto 110**. Cuando la conexión TCP está establecida, POP3 continúa con 3 fases: **Autorización, Transacción, Actualización**. En la primera fase, el agente de usuario envía el nombre de usuario y la palabra clave para autenticar al usuario que está descargando el correo. En la segunda fase el agente de usuario recupera los mensajes. La tercera fase ocurre cuando el agente de usuario haya emitido el comando **quit**, finalizando la sesión POP3.

16.4.4.2 IMAP

Con el acceso POP3, una vez descargado los mensajes en una maquina local, se pueden crear buzones de correos y mover los mensajes a los buzones. El usuario puede entonces realizar las operaciones de: borrar, mover, y buscar los mensajes en los buzones. El método POP3 tiene un problema para el usuario viajero que preferiría tener una jerarquía de correo en el servidor remoto y acceder desde cualquier dispositivo final sin importar geográficamente donde se encuentre en usuario final. El protocolo POP3 no proporciona ninguna forma de que el usuario pueda crear buzones remotos y asignar mensajes a los buzones.

Para dar solución a este problema surgió IMAP (Internet Mail Access Protocol), definido en el RFC 2060, es un protocolo de acceso de correo al igual que POP3, sin embargo este tiene muchas más funcionalidades que POP3, aunque es un poco más complejo.

El protocolo IMAP proporciona comandos que permiten a los usuarios crear buzones y mover los mensajes de un buzón a otro. IMAP también proporciona comandos para buscar en los buzones remotos mensajes que cumplan con criterios específicos. Otra característica importante es que IMAP tiene comandos que permiten al agente de usuario obtener componentes de mensajes. Por ejemplo, un agente puede obtener simplemente la cabecera de un mensaje, o simplemente una parte de un mensaje MIME.

Esta característica es importante cuando existe una conexión de bajo ancho de banda entre el agente de usuario y el servidor de correo. Con una conexión de bajo ancho de banda, el usuario podría no querer descargar todos los mensajes de buzón de correo, evitando especialmente los mensajes largos que puedan contener, por ejemplo, porciones de audio y video.

16.4.4.3 Correo electrónico basado en la web

Cada vez más los usuarios envían y acceden a su correo electrónico a través de navegadores web. Con este servicio, el agente de usuario es un navegador web habitual, y el usuario se comunica con su buzón de correo remoto vía HTTP. Cuando un destinatario quiere acceder a un mensaje en su buzón de correo, este es enviado desde el servidor de correo al navegador utilizando HTTP, en vez de POP3 o IMAP. Cuando un remitente desea enviar un mensaje de correo electrónico, este mensaje es enviado desde su navegador a su servidor de correo sobre HTTP, en vez de SMTP. Aun así, el servidor de correo del remitente sigue enviando y recibiendo mensajes de otros servidores de correo utilizando SMTP.

Esta solución de acceso al correo es enormemente adecuada para el usuario que va y viene. El usuario solo tiene que acceder a un navegador web para enviar y recibir mensajes. De hecho, muchas implementaciones de correo electrónico basadas en web utilizan un servidor IMAP para proporcionar la funcionalidad de buzones. En este caso, el acceso a los buzones y a los mensajes se realiza con scripts que se ejecutan en el servidor HTTP; los scripts utilizan el protocolo IMAP para comunicarse con el servidor IMAP.

16.5 Gestión de Red

En este tema se tratan solo los principios básicos de la gestión de red. La organización Internacional para la Estandarización (ISO; International Organization for Standardization) ha creado un modelo de gestión de red útil emplazando los escenarios anteriores en un entorno de trabajo más estructurado.

16.5.1 Áreas de gestión de red

Se definen cinco áreas en la gestión de red:

16.5.1.1 Gestión de rendimiento

El objetivo de la gestión de rendimiento es cuantificar, medir, informar, analizar y controlar el rendimiento de los distintos componentes de red. Estos componentes son, por ejemplo, los dispositivos individuales (enlaces, router y host).

16.5.1.2 Gestión de fallos

El objetivo de la gestión de fallos es registrar, detectar y responder a las condiciones de fallo en la red.

16.5.1.3 Gestión de configuración

La gestión de configuración permite al administrador de red hacer un seguimiento de los dispositivos de red y de las configuraciones hardware y software de estos dispositivos.

16.5.1.4 Gestión de cuentas

La gestión de cuentas permite al administrador de red especificar, registrar y controlar el acceso de los usuarios y dispositivos a los recursos de red.

16.5.1.5 Gestión de seguridad

El objetivo de la gestión de seguridad es controlar el acceso a los recursos de red de acuerdo con alguna política bien definida.

16.5.2 Infraestructura para la gestión de red

La gestión de red requiere la habilidad para supervisar, comprobar, sondear, configurar, y controlar los componentes hardware y software de una red. Dado que los dispositivos de red son distribuidos, como mínimo, el administrador de red debe ser capaz de recopilar datos de las entidades remotas, y de realizar cambios en la entidad remota.

Como se muestra en la Figura 100, existen tres componentes principales en una arquitectura de gestión de red: una entidad gestora, los dispositivos gestionados, y los protocolos de gestión de red.

La entidad gestora, es una aplicación con control humano que se ejecuta en una estación centralizada de gestión de red en el centro de operaciones de red. La entidad gestora es el lugar de actividad de la gestión de red; controla la recolección, procesamiento, análisis y/o visualización de la información de gestión

Un dispositivo gestionado es una parte del equipamiento de red que reside en la red gestionada. Un dispositivo gestionado puede ser un host, un router, un bridge, un hub, una impresora o un modem. En el dispositivo hay diversos objetos gestionados. Estos objetos gestionados son las partes reales de hardware de los dispositivos, y los conjuntos de parámetros de configuración de los dispositivos hardware y software. Estos objetos disponen de trozos de información que son recopilados en la base de información de gestión (MIB, Management Information Base), se verá que los valores de estos trozos de información están disponibles para la entidad gestora. Finalmente, existe un agente de gestión de red, que es un proceso residente que se ejecuta en cada dispositivo gestionado y que se comunica con la entidad gestora.

La tercera pieza de la arquitectura de gestión de red es el protocolo de gestión. El protocolo se ejecuta entre la entidad gestora y el dispositivo gestionado, permitiendo a la entidad gestora consultar el estado de los dispositivos, e indirectamente realizar acciones en dichos dispositivos a través de los agentes. Los agentes pueden utilizar el protocolo de gestión de red para informar a la entidad de eventos excepcionales.

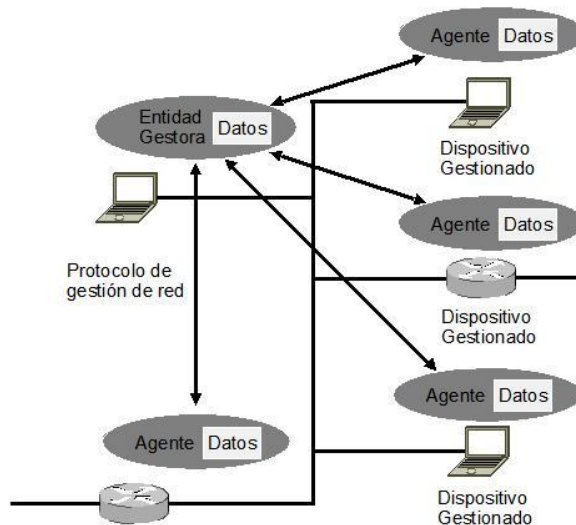


Figura 100. Componentes principales de una arquitectura de gestión de red

16.5.2.1 Base de información de gestión: MIB

Como se mencionó anteriormente, la base de información de gestión (Management Information Base, MIB) puede ser vista como un almacenamiento de información virtual, que contienen los objetos gestionados cuyos valores, en conjunto, reflejan el estado actual de la red. Estos valores pueden ser consultados y/o actualizados por una entidad de gestión enviando mensajes SNMP al agente que se está ejecutando en un dispositivo gestionado.

Los niveles inferiores del árbol de la MIB, muestran algunos de los módulos más importantes orientados al hardware (System e Interface), así como módulos que están asociados con algunos de los protocolos de internet más importantes. El RFC 3000 enumera todos los módulos MIB estandarizados.

Los objetos gestionados del sistema contienen información general sobre los dispositivos que están siendo gestionados.

16.5.3 Operaciones del protocolo SNMP y correspondencias de transporte

El protocolo SNMPv2 se utiliza para transportar la información MIB entre las entidades gestoras y los agentes que se ejecutan en representación de las entidades gestoras.

El uso más común de SNMP es un modo de petición-respuesta, en el que una entidad gestora SNMPv2 envía una petición a un agente SNMPv2, que recibe la petición, realiza alguna acción, y envía una respuesta a la petición. Típicamente, la petición es utilizada para consultar (obtener) o modificar (establecer) los valores de objetos MIB asociados con el dispositivo gestionado.

SNMPv2 define siete tipos de mensajes, conocidos genéricamente como Unidades de Datos de Protocolo (PDU; Protocol Data Unit).

Tipos de mensajes de PDU

- GetRequest.
- GetNextRequest.
- GetBulkRequest.
- InformRequest.
- SetRequest.
- Response.
- SNMPv2-Trap.

16.5.4 Protocolo de transporte utilizado por SNMP

Dada la naturaleza de petición-respuesta de SNMPv2, es conveniente destacar que, aunque las PDU pueden ser transportadas por medio de diferentes protocolos de transporte, típicamente son transportadas en un datagrama UDP.

El RFC 1906 establece que UDP es la correspondencia de transporte preferida. Dado que UDP es un protocolo de transporte no fiable, no existe garantía de que una petición, o su respuesta sean recibidas en un destino.

CAPÍTULO N° 3: DESARROLLO PRÁCTICO

PRÁCTICA N° 1: ASIGNACIÓN DE DIRECCIONES IPv4

Objetivo general

- Asignar direcciones IP correctamente a equipos finales.

Objetivos específicos

- Conocer cuáles son las direcciones IP habilitables de una dirección de Red para los equipos finales.
- Mostrar la importancia del Gateway en una Red.

Introducción

En la siguiente práctica estudiaremos el funcionamiento de asignación de direcciones IP. En esta práctica se asignara direcciones IP habilitables de una dirección de Red a los equipos finales pertenecientes a la Red LAN, de igual manera se estudiara la importancia de la dirección de Gateway.

Requerimientos

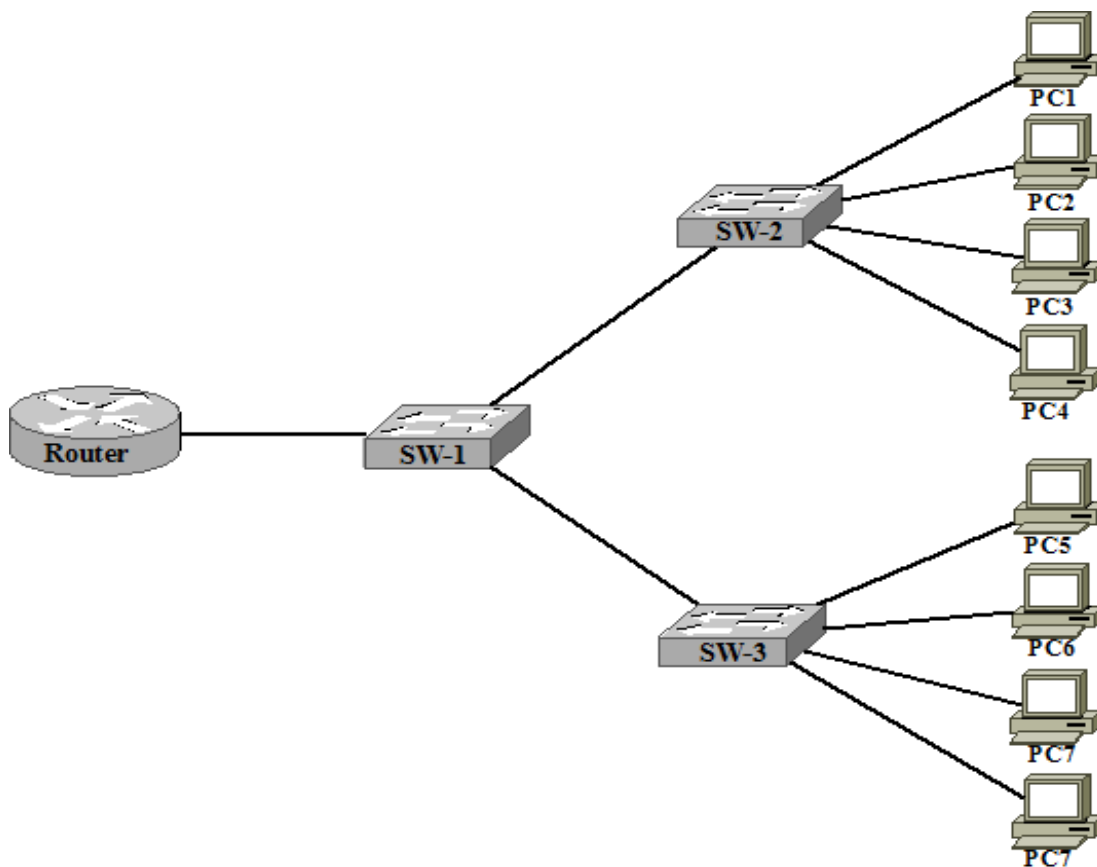
Para la realización de la práctica se necesitaron los siguientes requerimientos:

- Hardware:
 - Procesador con velocidad de 2.1 ghz
 - Memoria RAM de 1 GB.
 - Una PC con Windows.
- Software
 - Simulador de redes Packet Tracer 6.0 con los siguientes elementos:
 - 1 Router cisco de la serie 2811
 - 3 Switches 2960
 - 8 PCs

Conocimientos previos

Para la correcta realización de esta práctica el alumno deberá tener conocimientos básicos de direccionamiento IP, Puertas de enlace (Gateway), dirección de difusión y dirección de red.

Topología



Funcionalidad

En la figura de esta práctica se muestra la topología en la cual se estará asignando direcciones IP a los equipos finales y de igual manera se estará asignando la puerta de enlace para las PCs para comunicación entre los equipos finales pertenecientes a la red 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

ip address <i>dirección-ip .mascara de subred</i>	Asigna una dirección IP a una interfaz
interface <i>Tipo Número</i>	Cambios en el modo de configuración global al modo de configuración de interfaz
no shutdown	Habilita una interfaz
ping <i>ip-address</i>	sends an Internet Control Message Protocol (ICMP) echo request to the specified address

Datos del Router

Router			
Nombre	Interfaz	Dirección IP	Dirección de red
Router	Fa0/0	192.168.1.1/24	192.168.1.0/24

Enunciado**Asignación de direcciones IP**

Establecer las direcciones IP 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, rellendo el siguiente cuadro:

Nombre	Dirección IP	Máscara	Gateway
PC1			
PC2			
PC3			
PC4			
PC5			
PC6			
PC7			
PC8			

Pruebas de comunicación entre equipos finales

Habiendo asignado correctamente las direcciones IP y Gateway a las PCs y habiendo asignado la dirección al Router haremos la prueba de comunicación entre estos dispositivos mediante el comando Ping.

- Hacer ping desde PC1 a la dirección del Router.
- Hacer ping desde PC1 a la dirección del PC4.
- Hacer ping desde PC5 a la dirección del Router.
- Hacer ping desde PC5 a la dirección del PC8

1. 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 IP, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. De acuerdo a los conocimientos adquiridos ¿Qué rango de IP es assignable para la dirección de Red 192.168.10.0 mascara 255.255.255.0?
2. ¿Es correcto asignar la dirección 192.168.10.0 a la interfaz de un Router? Si ¿Por qué? No ¿Por qué?
3. ¿Cuáles son las clases de direcciones que existe en IPV4?
4. ¿Cuántos bits forman una dirección IPV4?
5. ¿Es correcto asignar la dirección 192.168.20.0 como puerta de enlace (Gateway)?
6. ¿Cuál es la dirección de difusión de la dirección de Red 192.168.1.0?

PRÁCTICA N° 2: **SEGMENTACIÓN DE REDES A TRAVÉS DE VLANs ESTÁTICAS**

OBJETIVOS

General

- Crear VLANs estáticas analizando el funcionamiento y forma de aplicarlas en un escenario de red cualquiera.

Específicos

- Poner en práctica los conocimientos del alumno, evaluando los temas de VLANs estáticas, subnetting, enrutamiento estático y DHCP utilizando la tecnología de cisco.
- Configurar subinterfaces y enrutamiento estático para la comunicación entre diferentes VLANs.
- Activar puertos troncales en los correspondientes Switches.

INTRODUCCIÓN

En la siguiente práctica estudiaremos el funcionamiento de las VLANs en modo Estático. El estudiante tendrá la posibilidad de simular una red física creando las VLANs en este caso modo estático. Las VLANs son importantes en una topología de red ofreciendo mayor seguridad a la red ya que brindan la facilidad de agrupar equipos en diferentes segmentos, esto con el fin de restringir el acceso a información a usuarios que no pertenecen a la misma VLANs.

REQUERIMIENTOS

Para la realización de la práctica se necesitaron los siguientes requerimientos:

- **Hardware**
 - Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.
- **Software**
 - Simulador Cisco PacketTracer 6.0 con los siguientes elementos:
 - 2 Routers serie 2811
 - 3 Switches serie 2960
 - 8 PCs
 - 7 Laptops
 - 1 Servidor
 - 2 Access Point

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica se deberá tener conocimientos en implementación de VLANs (Leer capítulo 2, Tema de VLANs, subnetting y enrutamiento estático, ya que será de mucha importancia para realizar la misma.

TOPOLOGÍA

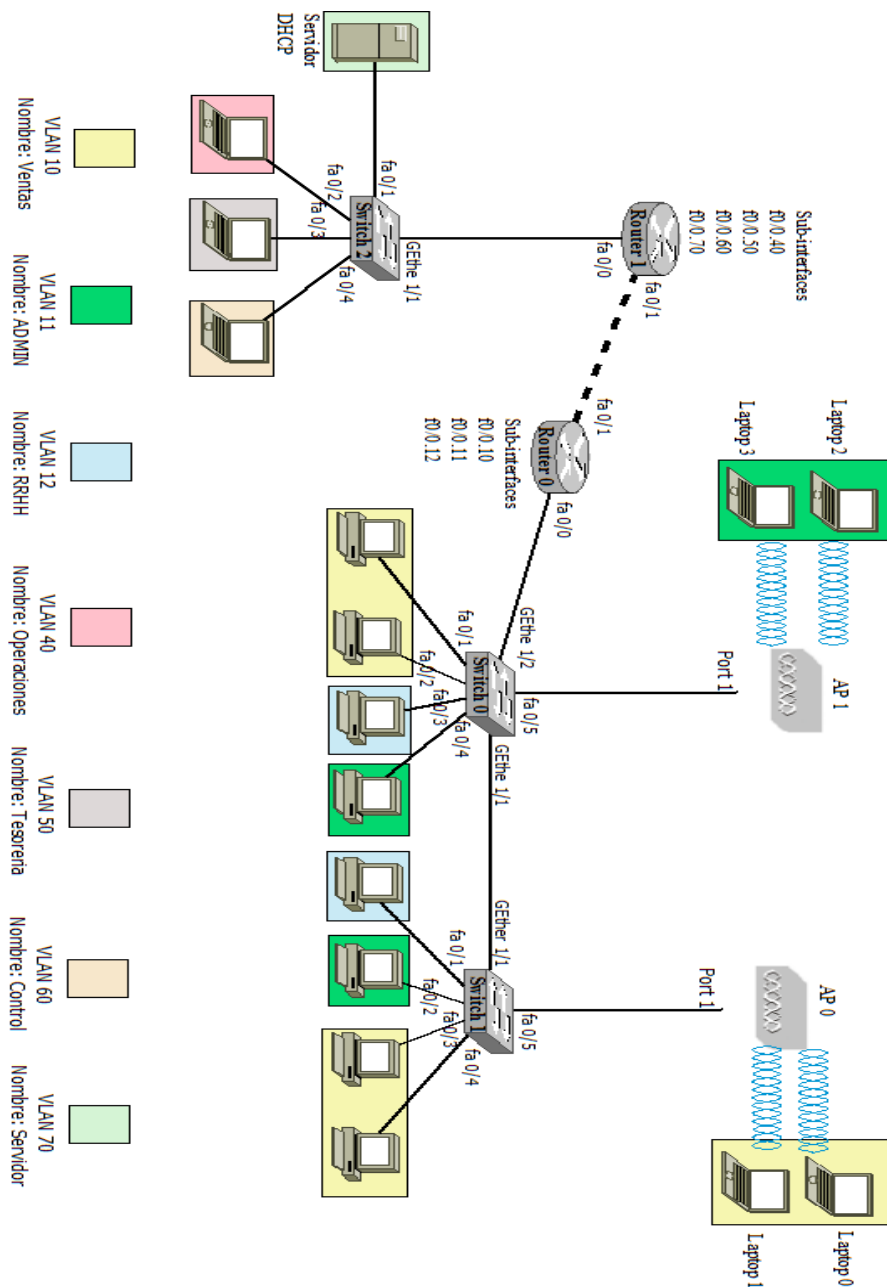


Figura 101. Topología de práctica de VLAN estática

FUNCIONALIDAD

La Figura 101, muestra la topología en la cual se estarán aplicando las configuraciones de VLANs estáticas, al igual que se estarán implementando de manera conjunta diferentes tecnologías y protocolos a como se mencionan a continuación:

- Se deberán crear siete VLANs de manera estática, con sus datos especificados en su respectivo cuadro que se muestra posteriormente llamado Switch 0 y Switch 1.
- Habrá al menos un protocolo de enrutamiento, en nuestro caso enrutamiento estático para la comunicación entre las diferentes subredes pertenecientes a la red LAN.
- 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 tanto entre las diferentes VLANs.
- Se deberá configurar un servidor DHCP para la asignación de direcciones de manera dinámica a los equipos finales pertenecientes a las diferentes VLANs.
- Se deberá configurar dos Access Point, para brindar cobertura inalámbrica a los usuarios móviles que tengan acceso a la red.

COMANDOS DE AYUDA

Comando	Descripción
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
ip address <i>dirección-ip .mascara de subred</i>	Asigna una dirección IP a una interfaz
interface range fastethernet <i>Inicio puerto-finalización del puerto</i>	Configura un rango de interfaces
interface <i>Tipo Número</i>	Cambios en el modo de configuración global al modo de configuración de interfaz
ip helper-address	Especifica una dirección de destino de UDP broadcasts para la asignación de direcciones DHCP
ip route <i>destination-pre x destination-pre x-mask {ip-address interface-type [ip-address]}</i>	Establece una ruta estática
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 <i>vlan-id</i>	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
vlan <i>vlan-id</i>	Crea una vlan

DATOS DE LOS DISPOSITIVOS

Switch 0 Switch 1			
Nombre del Switch	Número de VLANs	Nombre	Puertos Access
Switch 0	10	Ventas	Fa 0/1, Fa 0/2
	12	RRHH	Fa 0/3
	11	ADMIN	Fa 0/4, Fa0/5
Switch 1	10	Ventas	Fa0/3, Fa0/4, Fa0/5
	12	RRHH	Fa0/1
	11	ADMIN	Fa0/2

Switch 2		
Número y nombre de VLANs	Dirección subred	Puertos Access
40 Operaciones	192.168.0.0/24	Fa0/2
50 Tesorería	192.168.1.0/24	Fa0/3
60 Control	192.168.2.0/24	Fa0/4
70 Servidor	192.168.3.0/24	Fa0/1

Routers				
Equipo	Interfaz	Subinterfaces	Dirección IP	Etiquetado de VLAN
Router 1	Fa0/1 10.0.0.1/24	Fa0/0.40	192.168.0.1/24	40
		Fa0/0.50	192.168.1.1/24	50
		Fa0/0.60	192.168.2.1/24	60
		Fa0/0.70	192.168.3.2/24	70
Router 0	Fa0/1 10.0.0.2/24	Fa0/0.10	172.168.10.1/24	10
		Fa0/0.11	172.168.11.1/24	20
		Fa0/0.12	172.168.12.1/24	30

Pools			
Pool Name	Gateway	Inicio de dirección	Máximo usuarios
Operación	192.168.0.1	192.168.0.2/24	20
Tesorería	192.168.1.1	192.168.1.2/24	20
Control	192.168.2.1	192.168.2.2/24	20

Ventas	172.168.10.1	172.168.10.2/24	20
ADMIN	172.168.11.1	172.168.11.2/24	20
RRHH	172.168.12.1	172.168.12.2/24	20

Servidor DHCP		
Equipo	Dirección IP	Gateway
Servidor DHCP	192.168.3.1/24	192.168.3.2/24

ENUNCIADO

En esta práctica se pretende configurar una topología de red como la mostrada anteriormente en la Figura 101. Se deberá seguir paso a paso la configuración de las diferentes tecnologías que se implementara en este escenario.

Creación de VLANs

Como punto de inicial, se deberán crear las VLANs en los dispositivos (Switch 0, Switch 1 y Switch 2) cuyos datos se encuentran en el cuadro llamado Switch 0 Switch 1, y el cuadro llamado Switch 2.

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, y Switch 2.

Método de encaminamiento del tráfico

En los dispositivos Router 0 y Router 1 crear subinterfaces y establecer el etiquetado de VLAN para cada subinterfaz creada en estos dispositivos, para permitir comunicación entre las distintas VLANs mediante una misma interfaz física. Recuerde introducir correctamente la dirección IP para a cada subinterfaz creada en estos dispositivos. Ver cuadro Routers.

1. Router 0

Para comunicar el tráfico de las redes pertenecientes a las VLANs 10,11,12 con las redes pertenecientes a las VLANs 40,50,60,70, se deberá crear una ruta por defecto por cada red perteneciente a las VLANs 40,50,60,70 esto por el hecho del que el Router 0 no conoce dichas redes.

Nota: La dirección destino de las rutas por defecto será la ip de la interfaz FastEthernet 0/1 del Router 1.

2. Router 1

Para comunicar el tráfico de las redes pertenecientes a las VLANs 40, 50, 60,70 con las redes pertenecientes a las VLANs 10, 11,12, se deberá crear una ruta por defecto por cada red perteneciente a las VLANs 10, 11,12 esto por el hecho del que el Router 1 no conoce dichas redes.

Nota: La dirección destino de las rutas por defecto será la IP de la interfaz FastEthernet 0/1 del Router 0

Configuración de los Access Point

En esta sección es donde se deberá configurar los equipos que permitirán el acceso inalámbrico para los usuarios que tengan equipos con soporte a conexión WIFI, y tengan autorización para el acceso al servicio. Para cumplir con esta tarea, se debe configurar 2 Access Point (AP0 y AP1) que se ilustran en la figura de la topología en esta práctica

1. AP0

En el dispositivo AP0 en el Puerto 1 se configurara el SSID con el nombre de “oficina 1”.

Para los dispositivos Laptop0 y Laptop1 se deberá agregar físicamente un módulo Linksys-WPC300N para conexión wifi, tendrá también que configurar en la interfaz Wireless0 el mismo SSID que se configuro en el dispositivo AP0 “oficina 1” para que se conecte de manera inalámbrica con este dispositivo.

2. AP1

En el dispositivo AP1 en el Puerto 1 se configurara el SSID con el nombre de “oficina 2”.

Para los dispositivos Laptop2 y Laptop3 se deberá agregar físicamente un módulo Linksys-WPC300N para conexión wifi, tendrá también que configurar en la interfaz Wireless0 el mismo SSID que se configuro en el dispositivo AP1 “oficina 2” para que se conecte de manera inalámbrica con este dispositivo

Servicio DHCP

En el escenario de red de esta práctica, se dará servicio DHCP para las asignaciones de direcciones IP de manera dinámica, esto con la finalidad de que todos los usuarios pertenecientes a cada una de las VLANs puedan obtener su IP sin necesidad de introducirla de manera manual.

En el Servidor DHCP establecer su respectiva dirección IP y default Gateway a como se indica en el cuadro Servidor DHCP.

En el dispositivo Servidor DHCP se le deberá agregar los pools con el rango de direcciones que se les estarán asignando a los usuarios de cada VLANs y el número de usuarios máximos que pueden estar conectados para cada red virtual (VLANs).

Configuración de dispositivos finales (PC)

En las PC y en las Laptops pertenecientes a cada VLANs configurarlas de manera que estas efectúen una petición DHCP para obtener su correspondiente dirección IP.

Tiempo estimado de solución

➤ 2 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 VLANs, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Qué afirmaciones son verdaderas con respecto al uso de subinterfaces en el enrutamiento entre VLANs?
 - a) Las subinterfaces no tienen contención para el ancho de banda.
 - b) Se requieren más puertos de Switch que en el enrutamiento entre VLAN tradicional.
 - c) Se requieren menos puertos de Switch que en el enrutamiento entre VLAN tradicional.
 - d) Resolución de problemas más sencillo que con el enrutamiento entre VLAN tradicional.
 - e) Conexión física menos compleja que en el enrutamiento entre VLAN tradicional.

2. ¿Qué afirmaciones son verdaderas acerca del comando interface Fa0/0.10?
 - a) El comando aplica la VLAN 10 a la interfaz fa0/0 del Router.
 - b) El comando se utiliza en la configuración del enrutamiento entre VLAN del Router 0.
 - c) El comando configura una subinterfaz.
 - d) El comando configura la interfaz fa0/0 como un enlace troncal.
 - e) El comando no incluye una dirección IP, ya que ésta se aplica a la interfaz física.

3. ¿Cuáles son los tres elementos que deben utilizarse cuando se configura una interfaz de Router para el enlace de la VLAN?
 - a) Una subinterfaz por VLAN.
 - b) Una interfaz física para cada subinterfaz.
 - c) Una red o subred IP para cada subinterfaz.
 - d) Un enlace troncal por VLAN.
 - e) Un dominio de administración para cada subinterfaz.
 - f) Una encapsulación de protocolo de enlace compatible para cada interfaz.

4. ¿Qué es importante tener en cuenta al configurar las subinterfaces de un Router cuando se implementa el enrutamiento entre VLAN?
 - a) La interfaz física debe tener una dirección IP configurada.
 - b) Los números de las subinterfaces deben coincidir con el número de identificación de la VLAN.
 - c) El comando no shutdown se debe ejecutar en cada Subinterfaz.
 - d) La dirección IP de cada subinterfaz debe ser la dirección del gateway predeterminado para cada subred de la VLAN.

PRÁCTICA N° 3: **SEGMENTACIÓN DE REDES A TRAVÉS DE VLANs DINÁMICAS**

OBJETIVOS

General

- Configurar VLANs dinámicas en un ambiente de red con múltiples Switches.

Específicos

- Analizar el funcionamiento e interoperabilidad de las VLANs modo dinámico.
- Comprender el funcionamiento del protocolo VTP de propagación de las VLANs.

INTRODUCCIÓN

En esta práctica estudiaremos el funcionamiento de las VLANs en modo dinámico, la cual es uno de los tipos de VLANs que existe y uno de los más utilizados en la actualidad. El estudiante tendrá la posibilidad de simular una red física creando las VLANs en este caso modo dinámico. Las VLANs son importantes en una topología de red ofreciendo mayor seguridad a la red ya que brindan la facilidad de agrupar equipos en diferentes segmentos, esto con el fin de restringir el acceso a información a usuarios que no pertenecen a la misma VLANs.

REQUERIMIENTOS

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.

➤ Software

- Simulador Cisco Packet Tracer 6.0 con los siguientes elementos:
 - 3 Routers serie 2811.
 - 1 Switch serie 3560.
 - 9 Switches serie 2960.
 - 1 Router Linksys-WRT300N-Wifi.
 - 2 Tablets PC.
 - 2 Servidores.
 - 12 Ordenadores.

CONOCIMIENTOS PREVIOS

Para llevar a cabo la realización con éxito de la siguiente práctica, se necesita que el alumno tenga conocimientos bases acerca del establecimiento y funcionamiento de una VLANs estáticas e implementarlos en esta práctica para la correcta comprensión y funcionalidad de la misma.

TOPOLOGÍA

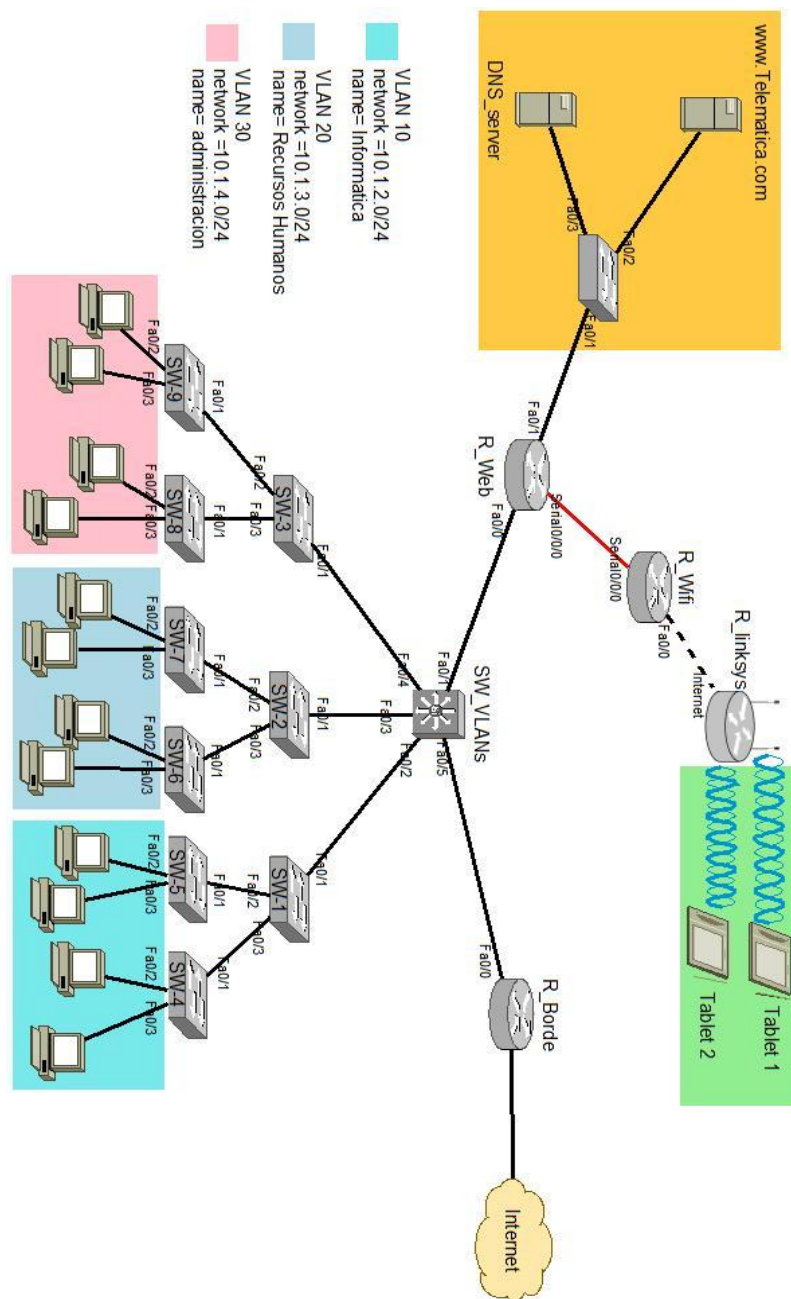


Figura 102. Topología de práctica de VLAN dinámica

FUNCIONALIDAD

En la topología planteada anteriormente lo que se pretende es crear una red LAN privada, la cual estará compuesta por tres segmentos de VLANs las cuales serán nombradas como: informática, recursos humanos y administración, al igual que también constara de conectividad wifi para los clientes con dispositivos inalámbricos, constara con una área

respectiva de servidores las cuales ofrecerán los servicios de DNS y HTTP para que los clientes visiten los sitios alojados en el servidor web.

COMANDOS DE AYUDA

Comando	Descripción
vtp domain {nombre dominio}	Establece un dominio VTP
ip routing	Establece encaminamiento de tráfico en un switch de capa 3
no switchport	Indica a una interfaz que no sea de Switch
show vlan brief	Muestra las VLANs existentes
show vtp status	Muestra el estado de VTP
interface vlan {num vlan}	Habilita una interfaz de VLAN
vlan database	Entra a la base de datos de las VLANs

DATOS DE LOS DISPOSITIVOS

Routers			
Nombre	Fa0/0	Fa0/1	Serial0/0/0
R_borde	10.1.5.65/28		
R_web	10.1.1.2/24	10.1.5.17/28	10.1.5.33/28
R_wifi	10.1.5.49/28		10.1.5.34/28

R_linksys			
Interfaz	Dirección IP	Gateway	DNS
Internet	10.1.5.50/28	10.1.5.49/28	10.1.5.19/28
LAN	192.168.0.1/24	-----	-----

Servidores		
Nombre	Dirección IP	Gateway
DNS_Server	10.1.5.19/28	10.1.5.17/28
Web_Server	10.1.5.18/28	10.1.5.17/28

Dominio	
Nombre dominio	Dirección IP
www.telematica.com	10.1.5.18/28

Switch Multicapa			
Nombre de VLANs	Número de VLANs	Interfaz	Dirección IP
Informática	10	VLAN10	10.1.2.0/24
Recursos Humanos	20	VLAN20	10.1.3.0/24
Administración	30	VLAN30	10.1.4.0/24
-----	-----	Fa0/1	10.1.1.1/24
-----	-----	Fa0/5	10.1.5.66/28

Switches de capa 2		
Nombre	Número de VLANs	Puertos Access
Switch4	10	Fa0/2 - Fa0/3
Switch5	10	Fa0/2 - Fa0/3
Switch6	20	Fa0/2 - Fa0/3
Switch7	20	Fa0/2 - Fa0/3
Switch8	30	Fa0/2 - Fa0/3
Switch9	30	Fa0/2 - Fa0/3

ENUNCIADO

Asignación de direcciones IP

Asignar a IP a las interfaces de los equipos que se mencionan a continuación:

- R_borde.
- R_web.
- R_wifi.
- R_linksys.
- Switch Multicapa (Fa0/1, Fa0/8).

Dicha direcciones se facilitan en el cuadro Routers y en el cuadro Switch Multicapa. Recuerde de asignar de manera correcta estas direcciones en a las interfaces que se indican en los dos cuadros antes mencionados.

Creación y propagación de VLANs

En el equipo Switch Multicapa se crearan las VLANs con sus nombres y direcciones IP a como se muestra en el cuadro llamado Switch Multicapa.

En el Switch Multicapa emplearemos el uso del protocolo VTP para la propagación de las VLANs, por lo que se deberá crear un dominio VTP llamado TELEMATICA y su funcionalidad será modo server.

Nota: Recuerde que el Switch multicapa tiene la facultad de actuar como un encaminador de datagramas, de manera que debemos de configurar esta opción, no olvide utilizar el comando `ip routing`. Por otra parte, no olvide que en este equipo se les asignaran IP a algunas de sus interfaces físicas, por lo tanto recuerde utilizar el comando `no switchport`. En este equipo en donde se crearan las VLANs, por lo tanto debe de crear un dominio VTP, de manera que este actuará como servidor de las VLANs para los otros equipos. Utilice el comando `vtp domain nombre_dominio`.

Configuración de Switches de capa 2

En los Switch de capa 2 de la serie 2960 configurarlos a modo cliente VTP, de manera que estos también pertenezcan al dominio VTP que fue creado en el servidor VTP (Switch Multicapa).

En los Switches 4, 5, 6, 7, 8, 9, configurar los puertos en modo access a como se especifica en el cuadro Switches de capa 2, de modo que en estos equipos es donde se segmentan los puertos que permitirán el tráfico de la VLAN especificada. Recuerde configurar estos equipos a modo cliente VTP, y que estos pertenezcan al dominio VTP creado en el Switch Multicapa.

Configuración de encaminamiento estático

En el Switch Multicapa establecer la siguiente ruta por defecto:

- Una ruta estática para que encamine el tráfico perteneciente a las diferentes VLANs hacia la zona de servidores.

En el router R_web establecer las siguientes rutas estáticas:

- Una ruta estática que permita el paso del tráfico proveniente de los servidores hacia los clientes de las diferentes VLANs (1 ruta por VLAN).
- Una ruta estática que permita el paso del tráfico proveniente de los servidores hacia los usuarios que estén conectados de manera inalámbrica.
- Una ruta por defecto que permita el paso del tráfico de los servidores al router de borde.

En el dispositivo R_wifi establecer la siguiente ruta estática:

- Una ruta estática que permita el paso del tráfico de los usuarios conectado al inalámbrico hacia la red LAN.

En el router R_borde establecer las siguientes rutas estáticas:

- Una ruta por defecto para encaminar el tráfico con destino hacia las diferentes VLANs (1 ruta por VLAN).
- Una ruta por defecto que permita el tráfico con destino hacia el área de los servidores.
- Una ruta por defecto que permita el paso del tráfico con destino a los usuarios inalámbricos.

Configuración del Router Linksys

En este equipo establecer las direcciones IP a las interfaces a como se muestra en el cuadro R_linksys.

En la pestaña GUI de este equipo configurar el servicio DHCP para los usuarios que estén conectados de manera inalámbrica, asignándole un rango de direcciones IP (192.168.0.100 – 192.168.0.149), con un máximo de 50 usuarios conectados.

Servicio DHCP

En el Switch Multicapa también se ofrecerá el servicio DHCP para la asignación de IP dinámicas a los clientes. Por lo tanto se harán las siguientes configuraciones:

- Deberá crear tres pool con el nombre correspondiente de cada VLANs para cada uno.
- Asignar un default router para cada pool que en este caso será la dirección IP de cada una de las VLANs.
- Asignar la dirección del servidor DNS para la resolución de nombres.

Indicar a los equipos finales (PC) para que realicen una petición DHCP para su respectiva asignación de IP.

Indicar que los usuarios (Tablets) conectados vía wifi obtengan sus direcciones mediante DHCP.

Configuración de los servicios DNS y WEB

El en servidor llamado web_server establecer su dirección IP, dirección de gateway y dirección DNS a como se muestra en el cuadro Servidores. En éste equipo editar el archivo index.html y establecer un rotulo que diga Ingeniería en Telemática.

En el servidor DNS establecer su dirección IP y default gateway a como se muestra en el cuadro Servidores. En este equipo agregar el dominio y dirección IP del servicio web al cual se le estará dando servicio de resolución de nombre a como se muestra en la cuadro llamado dominio.

Tiempo estimado de solución

- 2 horas

Prueba de análisis

1. ¿Por qué la VLAN 1 no fue creada?
2. ¿Cuál es la razón por la cual se crea un dominio VTP en el Switch Multicapa?
3. Ejecute el comando show vlan brief y analice el resultado. Haga un comentario del resultado
4. ¿Cuál es el objetivo de configurar el servidor VTP en el Switch Multicapa y en los demás no?
5. ¿Ejecute el comando show vtp status y analice el resultado? ¿Qué logra observar?

6. ¿Cuál es la diferencia de asignar un puerto de un Switch modo trunk a configurarlo en modo access?
7. En cualquiera de los Switch antes mencionados (4, 5, 6, 7, 8, 9) ejecute el comando `show vlan brief` y analice el resultado. Es el mismo resultado que el obtenido al ejecutarlo en el Switch multicapa. ¿Cuál es la diferencia?
8. ¿Con respecto al tipo de VLANs configurado en la práctica anterior, cual tipo de VLANs puede ser el más factible implementar en una topología de red por ejemplo en una empresa? ¿Por qué?
9. ¿Cuáles son las ventajas de utilizar VTP en una topología de red anterior?

PRÁCTICA N° 4: ALTA DISPONIBILIDAD A NIVEL DE ENLACE CON SPANNING TREE Y BALANCEO DE CARGA CON ETHERCHANNEL

OBJETIVOS

General

- Analizar los mecanismos de encaminamiento de tráfico y balanceo de carga a nivel de enlace.

Específicos

- Comprender la causa de la formación de bucles en una red conmutada.
- Conocer los diferentes mecanismos que ofrece STP para la prevención y recuperación ante bucles en la red.
- Utilizar las tecnologías que nos ofrece STP para asegurar y optimizar el funcionamiento de la red a nivel de enlace.
- Comprender el funcionamiento de la tecnología EtherChannel.
- Estudiar los mecanismos que ofrece EtherChannel para el balanceo de carga y optimización de ancho de banda.

INTRODUCCIÓN

En la presente práctica se pretende analizar y comprender el funcionamiento del protocolo STP Spanning Tree, siendo uno de los protocolos de suma importancia en la capa de enlace, ya que previene muchas anomalías como lo son los bucles en una red conmutada, siendo estos problemas los más comunes en las redes de empresas e instituciones, dándole un menor rendimiento y calidad a las respuestas ante un usuario. En esta práctica abordaremos algunos puntos importantes desde la evolución del protocolo STP tradicional basado en el estándar 802.1D hasta los más actuales tales como RSTP, MSTP.

Se estudiara y pondrá en uso otra de las tecnologías importantes a como lo es EtherChannel, para garantizar una mayor calidad en lo que respecta al tiempo de respuesta de un usuario a un servicio que preste una determinada empresa o institución, al igual que la optimización del ancho de banda de la misma. Se abarcará los diferentes métodos de balanceo de carga que esta tecnología ofrece.

REQUERIMIENTOS

Para la realización de esta práctica se requiere como mínimo lo siguiente:

- **Hardware**
 - Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.
- **Software**
 - Simulador Cisco Packet Tracer versión 6.0 con los siguientes elementos:
 - 8 Switches serie 2960.
 - 1 Switch serie 3560.
 - 24 PCs
 - 5 Servidores.

- 2 Routers serie 2811.
- Un Cloud.

CONOCIMIENTOS PREVIOS

Para dar marcha a la realización correcta de esta práctica es necesario que el estudiante tenga como mínimo conocimientos bases acerca de la teoría y operación del protocolo Spanning Tree (STP), al igual que también de la teoría y operación de la tecnología de EtherChannel. En el capítulo correspondiente a estas dos tecnologías abarcan lo que es el funcionamiento e importancia de la presencia de estas dos tecnologías en una red conmutada. Por lo que se le sugiere al alumno leer, analizar y comprender los puntos que se abarcan en los capítulos previos.

TOPOLOGÍA

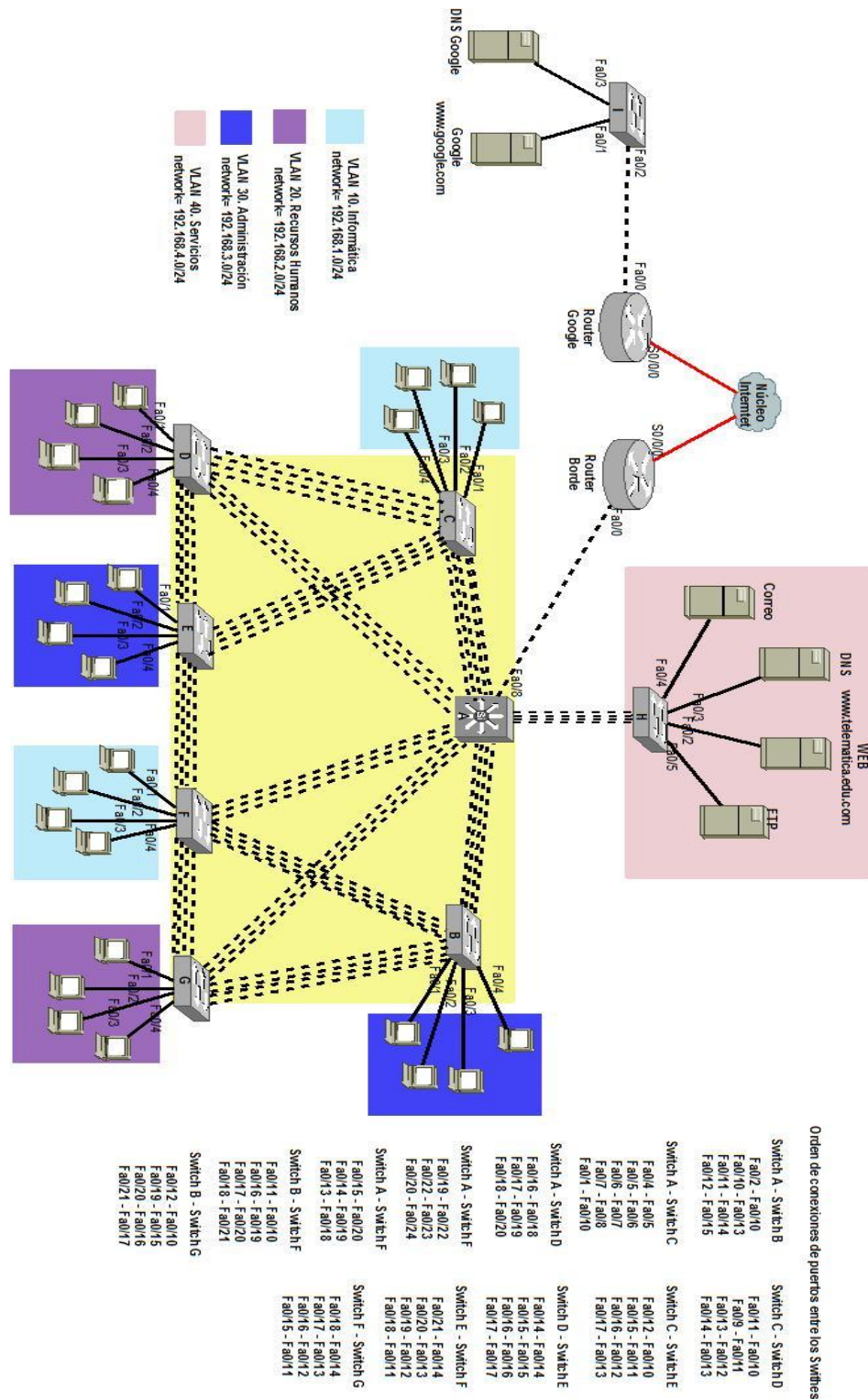


Figura 103. Topología de práctica de Spnning Tree & EtherChannel

FUNCIONALIDAD

En la topología planteada anteriormente se pretende crear una red LAN la cual estará segmentada en diferentes redes virtuales (VLANs) que representaran diferentes departamentos. Se estará implementando VLANs tipo dinámicas propagadas por el protocolo VTP, para el encaminamiento del tráfico se estará implementando enrutamiento dinámico (RIP), y la asignación de direcciones para los usuarios de las diferentes VLANs será mediante DHCP.

Uno de los puntos más importantes a tratar a como es uno de los objetivos de esta práctica es implementar el protocolo STP para la convergencia inmediata de la red y evitar tormentas de tráfico mediante broadcast. Se estará configurando ciertas tecnologías soportadas por el protocolo antes mencionado como lo es la parte de seguridad para mantener la integridad y el buen funcionamiento de la red.

En la topología también se muestra un área de servidores, en la que se prestan servicios WEB, DNS, FTP y Correo. Lo que se pretende en esta área es implementar otra importante tecnología planteada en los objetivos de esta práctica a como es EtherChannel, se estarán implementando diferentes funcionalidades que esta tecnología presta a como lo son los diferentes métodos de balanceo de carga que ofrece.

Por último la topología consta de una salida hacia el exterior en donde los usuarios de la LAN tendrán acceso a un servicio Web alojado en la nube.

COMANDOS DE AYUDA

Comando	Descripción
vlan {num-vlan}	Agregar las nuevas VLANs a la BD
vtp domain {nombre del dominio}	Crear un dominio VTP
vtp mode {server client}	Establece el modo de operación del dominio.
switchport access vlan{ num-vlan}	Configurar puertos de acceso
ip dhcp pool nombre	Crear un pool.
default-router ip	Asignar un default-router
dns-server ip	Asignar un DNS para cada VLANs
network ip mascara	Asignar una red al pool
router rip	Crear enrutamiento dinámico con rip
access-list num {permit deny} direccion-red mascara-Wilcard	Anadir una lista de acceso
ip nat {inside outside}	Establece una interfaz operando NAT.
ip nat inside source list {número de lista de acceso}	Establece la operación de traducción de las redes a la interfaz de salida.
interface {interfaz de salida} overload	
spanning-tree mode tipo-STP	Configurar el tipo de STP

spanning-tree vlan {rango-vlan} priority {prioridad}	Asignar un puente primario o secundario basándose en prioridad del dispositivo
spanning-tree vlan {rango-vlan} port-priority {prioridad}	Cambiar la prioridad de Puerto a un switch
spanning-tree portfast	Configurar PortFast en una interfaz
spanning-tree guard root	Configurar Root Guard en una interfaz
spanning-tree bpduguard {enable disable}	Configurar BPDU Guard en una interfaz
channel-protocol {pagp lacp}	Configurar el protocolo de negociación etherchannel para una interfaz o un grupo
channel-group num-grupo mode {active pasive desirable auto}	Asignar un número de grupo a las interfaces y el modo de negociación

DATOS DE LOS DISPOSITIVOS

Switch Multicapa			
Nombre y Número de VLAN	Dirección de subred	Dirección IP	Interfaz
VLAN10 Informática	192.168.1.0/24	192.168.1.1/24	VLAN10
VLAN20 Recursos Humanos	192.168.2.0/24	192.168.2.1/24	VLAN20
VLAN30 Administración	192.168.3.0/24	192.168.3.1/24	VLAN30
VLAN40 Servicios	192.168.4.0/24	192.168.4.1/24	VLAN40
-----	192.168.5.0/30	192.168.5.1/30	Fa0/8

Switches de Capa 2		
Nombre	Rango de Interfaces	Modo de operación
B	Fa0/1 - Fa0/4	Access
C	Fa0/1 - Fa0/4	Access
D	Fa0/1 - Fa0/4	Access
E	Fa0/1 - Fa0/4	Access
F	Fa0/1 - Fa0/4	Access
G	Fa0/1 - Fa0/4	Access

Servidores Internos			
Nombre	Dirección IP	Dirección de red	Default Gateway
DNS	192.168.4.2/24	192.168.4.0/24	192.168.4.1/24
Web	192.168.4.4/24		
FTP	192.168.4.5/24		
Correo	192.168.4.3/24		

Dispositivos Externos				
Equipo	Nombre	Dirección IP		Dirección de Red
Cloud	Núcleo de internet	Serial0	Serial1	-----
		-----	-----	
Router	Router Google	Fa0/0	Serial0/0/0	201.127.233.0/30
		100.1.1.1/24	201.127.233.2/30	100.1.1.0/24
Server	Google	Fa0/0		Default Gateway
		100.1.1.3/24		100.1.1.1/24
Server	DNS	Fa0/0		
		100.1.1.2/24		

Router de Borde			
Nombre	Dirección de Red	Dirección IP	Interfaz
router borde	192.68.5.0/24	192.168.5.1/24	Fa 0/0
	201.127.233.0/30	201.127.233.1/30	Serial0/0/0

Dominios		
Equipo	Dominio	Dirección IP
DNS Interno	www.telematica.edu.com	192.168.4.4
	www.google.com	100.1.1.3
DNS externo	www.google.com	100.1.1.3

ENUNCIADO

Asignación de direcciones IP

Asignar IP a los dispositivos de enrutamiento y a los que prestan servicios, tales como:

- Switch Multicapa (ver cuadro Switch Multicapa).
- Servidor de Correo, DNS, Web, FTP. (ver cuadro Servidores internos).
- Router Borde (cuadro Router de borde).
- Router Google (cuadro Dispositivos externos).
- Servidor de Google (cuadro Dispositivos externos).

Creación y propagación de VLANs

Crear las VLANs en el dispositivo llamado “A” que es el Switch Multicapa, con las especificaciones indicadas en el cuadro Switch Multicapa.

En el Switch Multicapa configurar un dominio VTP llamado TELEMATICA y ponerlo a modo Server, de manera que este servirá para la propagación de las VLANs a los dispositivos que estarán a modo cliente.

En los dispositivos B, C, D, E, F, G que son dispositivos de capa 2 se estarán configurando en el mismo dominio VTP que el servidor, de manera que en ellos se habilitaran y se propagaran las VLANs que fueron creadas en el servidor VTP (Switch Multicapa).

Posteriormente en estos mismos dispositivos (B, C, D, E, F, G) configurar los puertos en modo **access** indicados en el cuadro Switches de capa 2, la cual les permite a los usuarios finales pertenecer a la VLAN que se le asigne a cada puerto.

Asignación de direcciones mediante DHCP a los equipos finales.

Crear 3 pool en el Switch A llamados cada uno con el nombre de la VLANs que han sido creados anteriormente, indicando a cada pool la dirección de red correspondiente a cada VLANs, asignar un default router con su respectiva dirección IP (dirección de interfaz de la VLAN), y por ultimo asignar la dirección del servidor DNS para la resolución de nombres a los servicios que consumirán los miembros de las diferentes VLANs.

Configurar los dispositivos finales (PC) de manera que estos realicen una petición DHCP para obtener su respectiva dirección IP conforme a la VLAN que pertenezca cada equipo.

Configuración de Protocolo de enrutamiento

Para dar enrutamiento al tráfico entre las redes existentes, se estará usando RIP como protocolo de enrutamiento y debe ser configurado en los siguientes equipos:

En el dispositivo Switch Multicapa llamado A configurar enrutamiento mediante RIP, en donde se anunciaran las redes de las VLANs existentes hacia el router de Borde mediante la interfaz que está directamente conectada con el Switch Multicapa.

➤ Switch A.

```
Switch#configure terminal
Switch(config)#router rip
Switch(config-router)#network 192.168.1.0
Switch(config-router)#network 192.168.2.0
Switch(config-router)#network 192.168.3.0
Switch(config-router)#network 192.168.4.0
Switch(config-router)#network 192.168.5.0
Switch(config-router)#end
```

En el equipo Router de borde configurar enrutamiento dinámico al igual que en el Switch Multicapa, de manera que éste anunciara las redes tanto de las VLANs como la red punto a punto que tiene con el Switch multicapa mediante la interfaz de salida Serial0/0/0.

➤ **Router de borde.**

```
Router#configure terminal  
Router(config)#router rip  
Router(config-router)#network 192.168.5.0  
Router(config-router)#network 201.127.233.0  
Router(config-router)#end
```

Del lado externo en el equipo Router Google configurar RIP, de manera que este anunciara la red de los servidores hacia el exterior mediante la interfaz de salida Serial0/0/0.

➤ **Router Google.**

```
Router#configure terminal  
Router(config)#router rip  
Router(config-router)#network 100.1.1.0  
Router(config-router)#network 201.127.233.0  
Router(config-router)#end
```

Configuración de NAT

En el dispositivo Router de Borde configurar el protocolo NAT, de modo que los equipos que conforman la red interna podrán salir al exterior (Internet) mediante una única dirección pública. En este caso estaremos implementando NAT sobrecargado. Recuerde que para aplicar NAT correctamente debemos:

- Establecer de manera correcta las listas de acceso correspondiente a las direcciones que se estarán traduciendo.
- Indicar las interfaces en este equipo que estarán implementando NAT.
- Asocie la lista de acceso, especificando la interfaz de salida para la traducción.

En este momento si se han hecho bien las configuraciones anteriores debe de haber comunicación de los dispositivos de la LAN hacia el exterior (Internet).

Configuración de Spanning Tree (STP)

Para enmarcar la práctica en uno de sus objetivos, se iniciará a configurar el protocolo STP con algunas de sus características y funcionalidades planteadas en los siguientes puntos.

Como primer tarea a configurar el funcionamiento de esta tecnología en esta práctica, se deberá indicar en el dispositivo Switch Multicapa (A) el tipo de STP que queramos que se ejecute, cada uno de estos tipos abordados en el capítulo correspondiente a esta tecnología.

En el dispositivo Switch multicapa (A) configurarlo como puente raíz de manera predeterminada basándose en la prioridad del equipo, de modo que este será el punto de referencia central para los demás dispositivos que conformaran el árbol de expansión en la red y será por él en donde pase todo el tráfico de la VLANs existentes. Elegimos este dispositivo para cumplir un punto importante al momento de asignar un puente raíz por defecto, que es tomar en cuenta la posición física de la raíz con respecto al centro de la red de capa 2.

Configurar en el Switch de capa de enlace (capa 2) llamado B como puente raíz secundario basándose en la prioridad del equipo, esto con el objetivo de que si ocurre un fallo con respecto al puente raíz primario, el Switch B secundario pase a tomar el papel de puente raíz primario.

Nota: Para lo antes mencionado la prioridad del puente raíz primario puede ser de 4096, y para el puente raíz secundaria puede ser de 28672.

Después de tener elegido tanto el puente raíz primario como el secundario pasaremos a configurar otras opciones para influir en la decisión del STP con respecto al puerto raíz en los Switches que forman parte del árbol de expansión. Una de estas opciones que se configurará en esta práctica es poner una prioridad considerable (32) en el puerto raíz que STP elige por defecto en los switch B, C, D, E, F, G, esto a como dijimos anteriormente para influir a esta tecnología a que elija siempre a estos puertos como raíces al momento de iniciar la convergencia por el protocolo STP.

Nota: Recordar revisar en cada Switch antes mencionado cual fue el puerto raíz que se eligió por defecto con el comando `Switch # show spanning-tree` y a ese puerto se le cambiara el parámetro de la prioridad.

También se estarán configurando uno de los métodos existentes de suma importancia a como lo es PortFast para permitir una convergencia STP aún más rápida. Enfocándonos en lo abordado en la teoría en donde se hablan acerca de este método adicional configuraremos este método en los dispositivos que están en la capa de acceso (B, C, D, E, F, G) en los puertos en donde se encuentren conectadas estaciones finales de trabajo (PC).

Nota: En esta práctica los puertos en los switch que tienen conectados PC van en el rango de 1 a 4.

Otro de los puntos que estaremos configurando en esta práctica es la protección del correcto funcionamiento y la eficiencia del STP, esta protección se basa en la implementación de otras de las características más importantes existentes que integra esta tecnología. Estas características son Root Guard y BPDU Guard que en su momento se abarcaron en este documento en el capítulo de STP.

En esta práctica aplicaremos Root Guard en los Switch de capa 2 (B, C, D, E, G) cuyas interfaces estén propensas a que se puedan conectar nuevos Switches al momento de una alteración física en la estructura de la red, en este caso el rango de puertos que se le estará asignando esta tecnología en estos dispositivos es de la Fa0/22 a la Fa0/24.

De igual manera en estos mismos dispositivos configurar BPDU Guard en las interfaces que estén conectadas a estaciones de trabajo finales y estén ejecutando PortFast.

Configuración de EtherChannel

Para cumplir con otro objetivo propuesto planteados al inicio de ésta práctica, se deberá configurar EtherChannel con las funcionalidades que esta tecnología provee.

Como primer punto lo que se tiene que tener definido es con que protocolo de negociación queremos que nuestro EtherChannel opere. En esta práctica se estará implementando esta tecnología en la área de servidores, de modo que en el Switch Multicapa A elegir la interfaces físicas (fa0/3, 0/21, 0/23, 0/24) que están conectadas al Switch H que comunica el área de servidores, indicarle con que protocolo de negociación trabajaran estas interfaces, en este caso utilizamos el protocolo pagp y por ultimo a estas interfaces asignarles un número de grupo (1) al cual pertenecerán e indicarle el modo que estas participaran al momento de iniciar la negociación en este caso modo desirable.

Teniendo configurado en el Switch Multicapa ejecutando la tecnología EtherChannel lo que nos resta es que en el Switch H habilitaremos esta tecnología para una posterior negociación con el equipo A. En el equipo H elegir las interfaces físicas (fa0/6, 0/7, 0/8, 0/9) que están conectadas con el equipo A y de igual manera que se hizo en A elegiremos el protocolo de negociación que implementaran (pagp), asignar un número (1) de grupo a estas interfaces

e indicarle el modo en que actuaran al iniciar la negociación con el equipo A, a diferencia que A se configuró en modo desirable, en H configurarlo en modo auto, de modo que quien iniciara la negociación será el Switch Multicapa A.

De este modo tenemos configurado en esta red la tecnología EtherChannel, tanto con un dispositivo activo (A) que se encargara de generar la negociación con los puertos y un dispositivo pasivo (B) que se encargara de responder a esa negociación.

Configuración de Servicios externos e internos

En esta sección se deberá configurar servicios WEB y DNS para el uso de los clientes de la LAN.

Nota: Solamente se estará configurando los servidores WEB y DNS internos y externos, los demás no se abarcarán en ésta práctica.

En el servidor web interno llamado Web asignar su respectiva dirección IP dirección IP del DNS y default Gateway a como se indica en el cuadro Servidores Internos. En este equipo se editará el archivo index.html y se deberá establecer un rotulo que diga TELEMATICA.

De la misma manera se estará haciendo en el Servidor externo llamado Google, asignando su respectiva dirección IP dirección IP del DNS y default Gateway a como se indica en el cuadro Dispositivos Externos en la sección Server. En este equipo también se editará el archivo index.html y se deberá establecer un rotulo que diga Google.

En el Servidor interno llamado DNS asignar su respectiva dirección IP y default gateway a como se muestra en el cuadro Servidores Internos y agregar los nombres de dominio del sitio web con su respectiva dirección IP para la correspondiente resolución de nombres a como se muestra en el cuadro dominios.

De la misma manera que en paso anterior se hará en el servidor externo llamado google.

Tiempo estimado de solución

➤ 6 horas

Prueba de análisis

1. ¿Qué comando puede utilizar para configurar el Switch A como puente raíz con la VLAN 10, asumiendo que los otros switches están usando los valores STP por defecto?
2. Asumiendo que se eligió el switch A como raíz primario. ¿Qué comando se utiliza para configurar el switch B como puente raíz secundario?
- 3.Cuál de los Switches siguientes se establecerá como puente raíz, utilizando la información de la tabla siguiente?
¿Qué Switch se establecerá como puente raíz secundario si el primario falla?

Nombre del Switch	Prioridad de puente	Dirección MAC	Costo puerto
A	32,768	00-d0-10-34-26-a0	Todos 19
B	32,768	00-d0-10-34-24-a0	Todos 4
C	32,767	00-d0-10-34-27-a0	Todos 19
D	32,769	00-d0-10-34-24-a1	Todos 19

En las preguntas de la 4 a la 7 nos basaremos en una red que contiene 2 switches A y B, la prioridad de puente y las direcciones MAC para cada uno son las siguientes: 32,768:0000.aaaa.aaaa y 32,768 0000.bbbb.bbbb respectivamente.

4. ¿Cuál de los dos Switches se establecerá como puente raíz?
5. Si la prioridad del switch B cambiara a 10,000, ¿Cuál de ellos será el puente raíz?
6. Si la prioridad del switch B cambia a 32,769, ¿Cuál de ellos será el puente raíz?
7. Si se introduce un switch C con una prioridad y una MAC: 40,000:0000.0000.cccc, ¿cuál de ellos se establecerá como puente raíz secundario?
8. ¿Qué comando de configuración se necesita para seleccionar LACP en un EtherChannel como protocolo de negociación?
9. ¿Qué comando se usa para ver el estado de todos los puertos que forman un EtherChannel?
10. ¿Cuál comando puedes usar para verificar que algoritmo hash o método de balanceo de carga se está utilizando en un EtherChannel?

**PRÁCTICA N° 5: ALTA DISPONIBILIDAD A NIVEL DE
ENLACE CON SPANNING TREE Y SEGURIDAD CON
EL PROTOCOLO PPP CON AUTHENTICACIÓN CHAP**

OBJETIVOS

General

- Implementar el protocolo punto a punto (PPP) y funciones que proporciona, incorporando también la funcionalidad del protocolo de árbol de expansión STP.

Específicos

- Configurar el método de autenticación CHAP.
- Designar un puente raíz primario y un secundario para una mayor estabilidad en el funcionamiento de STP.

INTRODUCCIÓN

En esta práctica se estará configurando el protocolo PPP con modo de operación, en este caso con autenticación CHAP. Se podrá simular una red física configurando el protocolo punto a punto (PPP) y la funcionalidad del protocolo Spanning Tree para evitar bucles en el nivel de enlace. Con la implementación y realización de esta práctica, se comprenderá el funcionamiento de PPP incorporando Spanning Tree, para darle una mejor seguridad a la red en el nivel de enlace.

REQUERIMIENTOS

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.

➤ Software

- Simulador Cisco Packet Tracer 6.0 con los siguientes elementos:
- - 3 Router 2811
 - 12 Switches 2960
 - 1 Switch 3560
 - 11 PCs

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica el alumno deberá tener conocimientos de implementación de VLANs, enrutamiento estático funcionamiento de Spanning Tree y PPP, ya que será de mucha importancia para la realización y correcta funcionalidad de la misma.

TOPOLOGÍA

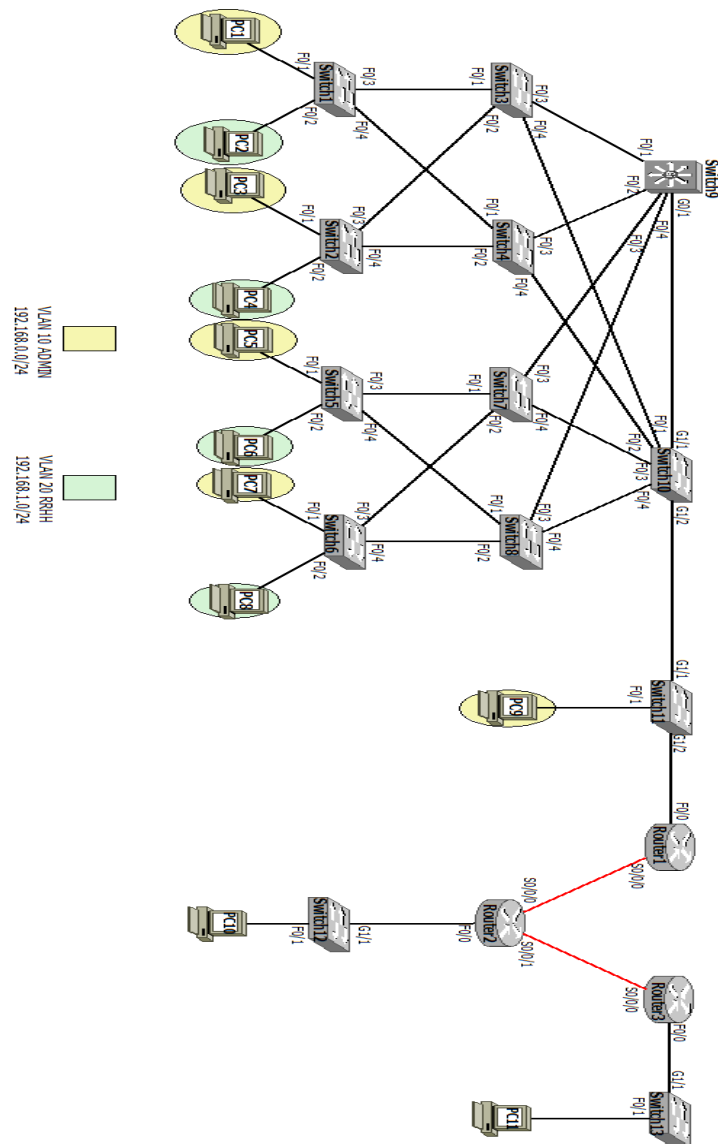


Figura 104. Topología de práctica de Spanning Tree & PPP

FUNCIONALIDAD

La [Figura 104](#), muestra la topología en la cual se estarán aplicando conjuntamente diferentes tecnologías que operan en la capa de enlace, en esta práctica se estará implementando la funcionalidad del protocolo de árbol de expansión para la prevención de bucles en el nivel de enlace (Spanning Tree) y el protocolo de conexión punto a punto (PPP) para dar seguridad en el nivel de enlace. Además de otras tecnologías como VTP para la propagación de las VLANs de manera dinámica entre otras.

COMANDOS DE AYUDA

Comando	Descripción
hostname <i>host-name</i>	Establece el nombre del dispositivo
ip address <i>dirección-ip .mascara de subred</i>	Asigna una dirección IP a una interfaz
interface range fastethernet <i>Inicio puerto-finalización del puerto</i>	Configura un rango de interfaces
interface <i>Tipo Número</i>	Cambios en el modo de configuración global al modo de configuración de interfaz
ip route <i>destination-pre x destination-pre x-mask {ip-address interface-type [ip-address]}</i>	Establece una ruta estática
ip routing	Habilita ruteo de IP en el Switch Multicapa
no shutdown	Habilita una interfaz
no switchport	Deshabilita las características de un Puerto de switch
show spanning-tree vlan <i>vlan-id</i>	Muestra si Spanning Tree está ejecutando una VLAN
spanning-tree vlan <i>vlan-id priority Prioridad</i>	Establece la prioridad de árbol de expansión para su uso en el ID de puente
switchport access vlan <i>vlan-id</i>	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
switchport trunk encapsulation dot1q	Configura troncal para la encapsulación 802.1Q
vlan <i>vlan-id</i>	Crea una vlan
username <i>router-nombre password password</i>	Crea una entrada de usuario para la autenticación Router remoto
encapsulation ppp	Permite la encapsulación PPP
ppp authentication chap	Permite la autenticación CHAP

DATOS DE LOS DISPOSITIVOS

Routers			
Equipo	Interfaz	Dirección IP	Dirección de subred

Router 1	Fa0/0	192.168.2.2/24	192.168.2.0/24
	Serial0/0/0	10.0.0.1/24	10.0.0.0/24
Router 2	Fa0/0	192.168.3.1/24	192.168.3.0/24
	Serial0/0/0	10.0.0.2/24	10.1.1.0/24
	Serial0/0/1	10.0.1.1/24	10.0.1.0/24
Router 3	Fa0/0	192.168.4.1/24	192.168.4.0/24
	Serial0/0/0	10.0.1.2/24	10.0.1.0/24

Switch Multicapa				
Número VLAN	Nombre VLAN	Interfaz	Dirección IP	Dirección de subred
10	ADMIN	VLAN10	192.168.0.1/24	192.168.0.0/24
20	RRHH	VLAN20	192.168.1.1/24	192.168.1.0/24
-----	-----	GigabitEthernet0/1	192.168.2.1/24	192.168.2.0/24

Switches de capa 2		
Nombre	Número VLANs	Puertos Access
Switch13	-----	-----
Switch 12	-----	-----
Switch 11	10	Fa0/1
Switch 10	-----	-----
Switch 8	-----	-----
Switch 7	-----	-----
Switch 6	10, 20	Fa0/1, Fa0/2
Switch 5	10, 20	Fa0/1, Fa0/2
Switch 4	-----	-----
Switch 3	-----	-----
Switch 2	10, 20	Fa0/1, Fa0/2
Switch 1	10, 20	Fa0/1, Fa0/2

ENUNCIADO**Asignación de direcciones IP**

Establecer las direcciones IP a las interfaces de los equipos que se mencionan a continuación:

- Router 1.
- Router 2.
- Router 3.
- Switch Multicapa (ver cuadro Switch Multicapa)

Ver cuadro Routers para la asignación de IP a las interfaces de los Routers.

Creación y propagación de VLANs

Crear las VLANs en el dispositivo Switch Multicapa llamado Switch 9 con los datos facilitados en el cuadro Switch Multicapa, se le deberán asignar el nombre y las direcciones IP a cada una de las interfaces virtuales (VLANs) creadas. En este equipo se deberá configurar el rango de interfaces desde la FastEthernet 0/1 hasta la Fa0/4.

Recuerde indicar que la interfaz GigabitEthernet0/1 del Switch Multicapa opere como interfaz de ruteo de tráfico, es decir que esta no actuara como puerto de Switch, sino que funcionará como puerto de Router con el comando no switchport. Ver cuadro comandos de ayuda para más información.

Establecer un dominio VTP llamado cisco e indicarle a este equipo que opere en modo server VTP, esto con el objetivo de propagar las VLANs creadas en este equipo hacia los demás Switches de capa 2.

En los Switches de capa 2 configurarlos para que operen bajo el domino VTP creado anteriormente en el Switch Multicapa, con la diferencia que estos funcionaran en modo cliente VTP. En estos equipos es donde también seleccionaremos los puertos Access que pertenecerán a cada una de las VLANs (ver cuadro Switches capa 2).

Protocolo de encaminamiento de tráfico

Establecer rutas estáticas en los dispositivos Router 1, Router 2, Router 3, y Switch Multicapa, de manera que estos equipos permitan enviar tráfico desde un punto de la red a otro, es decir que pueda haber comunicación entre todas las subredes (VLANs) pertenecientes a la LAN.

1. Switch Multicapa

En el Switch Multicapa establecer una ruta por defecto que permita el paso del tráfico de las VLANs con destino al Router 1.

2. Router 1

Una ruta por defecto que permita el paso del tráfico de las VLANs proveniente del Switch Multicapa hacia el Router 2.

En el Router 1 establecer una ruta por defecto que permita el paso de tráfico con destino hacia las VLANs (1 ruta por VLAN).

3. Router 2

Una ruta por defecto que permita el paso del tráfico con destino hacia las VLANs.

Una ruta por defecto que permita el paso del tráfico con destino hacia la red del Router 3.

4. Router 3

Una ruta por defecto que permita el paso del tráfico con destino a la red del Router 2 y a la red de las VLANs.

Servicio DHCP

En el equipo Switch Multicapa, se deberá establecer dos pool llamados cada uno con el nombre de las VLANs que han sido creadas anteriormente.

A cada pool indicar la dirección de subred de cada VLAN, además para cada pool establecer un default Router en este caso la dirección IP de cada interfaz virtual (VLANs).

Indicar a los equipos finales que deberán realizar una petición DHCP, con el propósito de que estos obtengan una dirección IP ofrecidas por los pool creados anteriormente, de manera que cada dispositivo final obtendrá su correspondiente dirección IP conforme a la configuración del puerto del equipo que envió la solicitud DHCP. Recuerde de revisar si al equipo final se le asignó su respectiva IP.

Configuración de Spanning Tree

Con los conocimientos adquiridos tanto en la teoría y en la práctica anterior con respecto al funcionamiento e importancia del protocolo Spanning Tree en una red en el nivel de enlace, se deberá implementar dicha tecnología en esta práctica para evitar los bucles o inundación de tráfico en la red, puesto que en este escenario de red en lo que corresponde al nivel de enlace se encuentran muchas rutas redundantes entre los Switches de capa 2.

Se deberá configurar el equipo Switch Multicapa como puente raíz primario, y al Switch 10 como puente raíz secundario, este último con el objetivo de que si ocurre alguna anomalía con el Switch Multicapa este tome el papel de puente raíz primario.

Para cumplir con la tarea antes mencionada, se debe configurar estos papeles en cada uno de los equipos (Switch Multicapa y Switch 10) indicando en el Switch Multicapa que sea el puente raíz primario de todas las VLANs creadas y el Switch 10 indicaremos que actúe como puente raíz secundario igual para todas las VLANs. Recuerde consultar la tabla de comandos de ayuda para que pueda cumplir con la tarea antes mencionada.

Configuración de PPP

De manera global en cada Router indicar un nombre de usuario y una clave. En este caso el nombre de usuario será el nombre del equipo, y la clave será cisco para cada equipo.

Para cumplir con uno de los objetivos de la práctica, en este apartado se deberá configurar en protocolo PPP con autenticación CHAP para cada interfaz serial de los Routers.

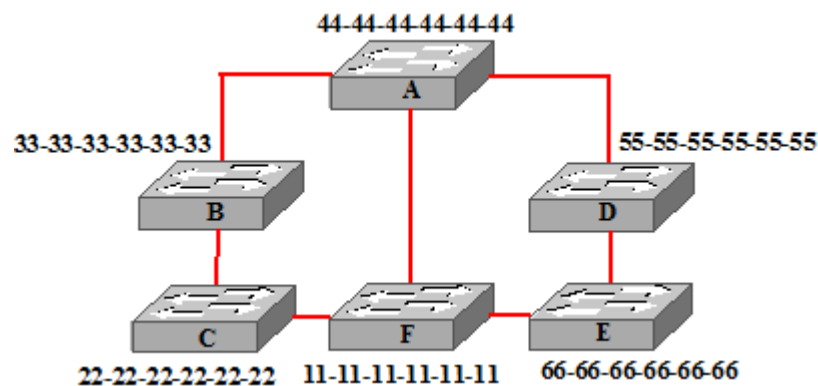
Tiempo estimado de solución

➤ 4 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 Spanning Tree y PPP, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

- De acuerdo a la siguiente ilustración. Cada Switch se muestra con su dirección Mac ¿Qué Switch será elegido como puente raíz Spanning Tree si los Switches se configuran con sus valores de prioridad?



- De acuerdo a la siguiente ilustración ¿Qué afirmaciones son verdaderas con respecto a que representa el valor de costo 23 para el switch4? (elijá dos opciones)

```
Switch4# show spanning-tree
VLAN0001
Spanning tree enabled protocol ieee
Root: ID Priority 24577
      Address    0019.2f8d.d200
      Cost       23
      Port       16 (FastEthernet0/14)
      Hello Time 3 sec Max Age 30 sec Forward Delay 15 sec
<resultado omitido>
```

- El costo representa la ruta de menor costo para el switch4 hasta el Switch raíz. Un costo de 23 es el valor publicado por el puerto 16 en el Switch (más cercano) al Switch raíz.
- Switch4 agrega el costo de un enlace FastEthernet a 23 para determinar el costo total para alcanzar el Switch raíz.
- El switch4 está conectado mediante un enlace FastEthernet a un Switch que a su vez está conectado directamente con el Switch raíz a través del enlace gigabit Ethernet.
- Un costo de 23 es el valor publicado por el puerto 16 en el Switch "upstream" (más cercano) al Switch raíz.

- e) El Switch raíz publica un costo de 23 que es menor a cualquier otro Switch en el dominio Spanning-Tree de la VLAN001
3. ¿Por qué es importante que un administrador de red considere el diámetro de la red spanning-tree cuando elige un puente raíz?
- a) La distancia de cableado entre los Switches es de 100 metros.
 - b) El límite del diámetro de la red es de 9.
 - c) La convergencia es más ya que el BPDU viaja lejos De la raíz.
 - d) Los BPDU pueden descartarse debido a los temporizadores de expiración.
4. ¿Cuáles son la dos afirmaciones verdaderas acerca de la operación predeterminada de STP en un entorno conmutado capa 2 que tiene conexiones redundantes entre switches? (elija dos opciones).
- a) El Switch raíz es el Switch con los puertos de velocidad más.
 - b) Las decisiones sobre qué puerto bloquear cuando dos puerto tienen igual costo depende de la prioridad e identidad del puerto.
 - c) Todos los puertos de enlace están designados y no Bloqueados.
 - d) Los puentes raíz tienen todos sus puertos establecidos como puerto raíz.
 - e) Cada uno de los Switches que no son raíz tienen un solo puerto raíz.
5. ¿Cómo puede el administrador de la red influir para que un Switch con STP se vuelva el puente raíz primario?
- a) Configura todas las interfaces del Switch como puerto raíz estática.
 - b) Cambiar la BPDU a un valor más bajo que el de otros Switches de la red.
 - c) Asigna al Switch una dirección IP más baja que los otros Switches de la red.
 - d) Establece la prioridad del Switch a un valor más pequeño que el de los Switches de la red.
6. ¿Cuáles de las siguientes son tres características del protocolo CHAP? (Elija tres opciones).
- a) Intercambia un número de desafío aleatorio durante la sesión para verificar la identidad.
 - b) Envía una contraseña de autenticación para verificar la identidad.
 - c) Impide la transmisión de información de conexión en texto sin cifrar.
 - d) Desconecta la sesión PPP si la autenticación falla.
 - e) Inicia un protocolo de enlace de tres vías.
 - f) Es vulnerable a los ataques de reproducción.
7. ¿Cuáles de las siguientes son tres afirmaciones que describen correctamente la autenticación de PPP? (elije tres opciones).
- a) PAP envía la contraseña en texto sin cifrar.
 - b) PAP usa un protocolo de enlace de tres vías para establecer un enlace.
 - c) PAP proporciona protección contra ataques reiterados de ensayo y error.
 - d) CHAP usa un protocolo de enlace de tres vías para establecer un enlace.

- e) CHAP utiliza un desafío respuesta que está basado en el algoritmo de hash MD5.
 - f) CHAP realiza la verificación mediante desafíos repetitivos.
8. Consulte la ilustración ¿Que operaciones es verdadera acerca de las operaciones PPP?

```
Serial is up, line protocol is up
Hardware is HD64570
Internet address is 200.200.200.1/24
MTU 1500 bytes, BW 1544 Kbit, DLY 20000 usec,
reliability 255/255, txload 1/255, rxload 1/255
Encapsulation PPP, Loopback not set
Keepalive set (10 sec)
LCP Open
Open: IPCP, CDPCP
Last input 00:00:04, output 00:00:04, output hang never
Last clearing of "show interface" counters 00:08:59
Input queue: 0/75/0/0 (size/max/drops/flushes); Total output drops : 0
queueing strategy: weighted fair
Output queue: 0/1/256 (active/max total/threshold/drops)
Conversations 0/1/256 (active/max active/max total)
Reserved Conversations 0/0 (allocated/max allocated)
Available Bandwidth 1158 kilobits/sec
```

- a) La capa dos esta desactivada.
- b) Las negociaciones LCP, IPCP y CDPCP están en curso.
- c) Solo la fase de establecimiento de enlace se completó con éxito.
- d) Tanto la fase de establecimiento de enlace como la fase de capa de red se completaron con éxito.

PRÁCTICA N° 6: **BALANCEO DE CARGA HSRP**

OBJETIVOS

General

- Implementar el protocolo HSRP.

Específicos

- Designar Router maestro y Router de respaldo en Routers redundantes para evitar puntos de fallos únicos en la red.
- Establecer grupos HSRP para cada interfaz de VLANs.

INTRODUCCIÓN

En la siguiente práctica estudiaremos el funcionamiento del protocolo HSRP. Se podrá simular una red física configurando HSRP. Con la implementación y realización de esta práctica el alumno Será capaz de entender el funcionamiento de HSRP y poder aplicar este tipo de protocolo.

REQUERIMIENTOS

- **Hardware**
 - Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.
- **Software**
 - Simulador Bosson Netsim 8.0
 - 2 Switches 3550
 - 2 Switches 2960
 - 8 PCs

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica el alumno deberá tener conocimientos de implementación de HSRP (Leer capítulo 2, Tema HSRP, ya que será de mucha importación para realizar la misma.

TOPOLOGÍA

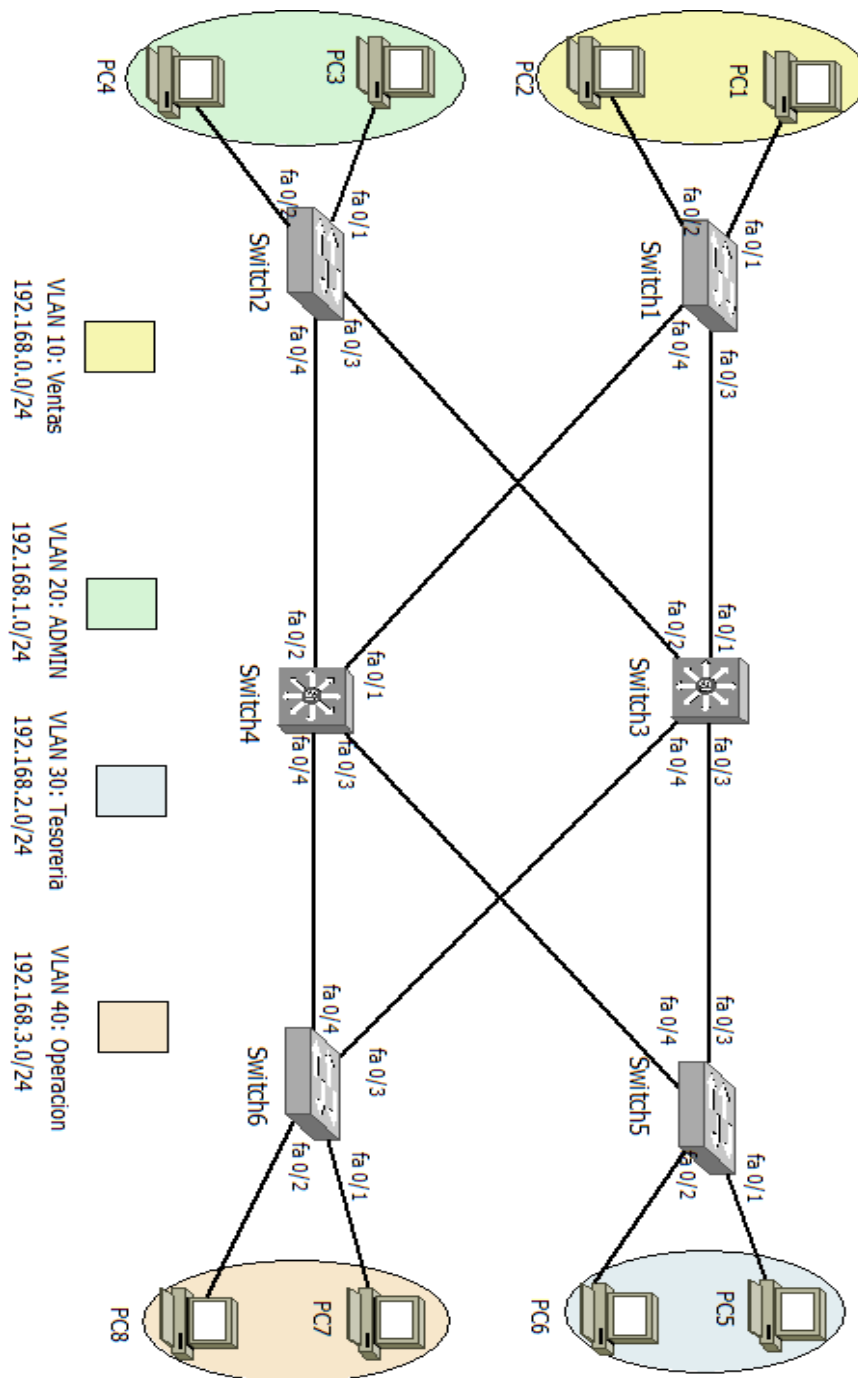


Figura 105. Topología de práctica de HSRP

FUNCIONALIDAD

La Figura 105 muestra la topología en la cual se estarán aplicando la tecnología HSRP. Además de que se estará usando de manera conjunta con diferentes tecnologías y protocolos para que cumplan las siguientes funcionalidades:

- Se deberán crear cuatro VLANs, de manera que estas se propagarán dinámicamente mediante VTP.

- Se deberán asignar dirección IP correspondiente en las interfaces de las VLANs en los dos Switches Multicapa.
- Se habilitará el protocolo HSRP en cada interfaz de las VLANs asignándole un grupo por VLAN con su correspondiente prioridad.
- Los equipos finales se deberán comunicar haciendo ping entre ellos.

COMANDOS DE AYUDA

Comando	Descripción
ip address <i>dirección-ip mascara de subred</i>	Asigna una dirección IP a una interfaz
interface range fastethernet slot/starting-port - ending-port	Configura un rango de interfaces
interface <i>tipo número</i>	Cambios en el modo de configuración global al modo de configuración de interfaz
no shutdown	Habilita una interfaz
show standby	Muestra información HSRP
standby [Número de grupo] ip [Dirección IP]	Crea o habilita el grupo HSRP con un número de grupo y la dirección IP virtual
standby [número-grupo] priority <i>prioridad</i>	Establece el valor de prioridad se utiliza para elegir el Router activo, el valor por defecto es 100, y el rango de valor de prioridad configurable es de 1 a 255, donde 255 es la más alta prioridad
standby [número-grupo] preempt	Permite preferencia de HSRP en un dispositivo
standby [group-number] track [interfaces]	Permite restar prioridad si la interfaz indicada cae

DATOS DE LOS DISPOSITIVOS

Switches Multicapas				
Equipo	Nombre y número de VLAN	Dirección IP	Dirección de subred	Dirección virtual y Grupo HSRP
Switch 3	Ventas (VLAN 10)	192.168.0.253/24	192.168.0.0/24	192.168.0.254/24 Grupo 1
	ADMIN (VLAN 20)	192.168.1.253/24	192.168.1.0/24	192.168.1.254/24 Grupo 2
	Tesorería (VLAN 30)	192.168.2.253/24	192.168.2.0/24	192.168.2.254/24 Grupo 3
	Operación (VLAN 40)	192.168.3.253/24	192.168.3.0/24	192.168.3.254/24 Grupo 4
Switch 4	Ventas (VLAN 10)	192.168.0.252/24	192.168.0.0/24	192.168.0.254/24 Grupo 1
	ADMIN (VLAN 20)	192.168.1.252/24	192.168.1.0/24	192.168.1.254/24

				Grupo 2
	Tesorería (VLAN 30)	192.168.2.252/24	192.168.2.0/24	192.168.2.254/24
				Grupo 3
	Operación (VLAN 40)	192.168.3.252/24	192.168.3.0/24	192.168.3.254/24
				Grupo 4

Rango de puertos Access		
Equipo	Puertos Access	VLAN ID
Switch 1	Fa0/1 – 2	10
Switch 2	Fa0/1 - 2	20
Switch 5	Fa0/1 - 2	30
Switch 6	Fa0/1 - 2	40

ENUNCIADO

Creación y propagación de VLANs

Crear las VLANs en el dispositivo llamado Switch3 (ver cuadro Switches-Multicapa), habilitando y asignándoles las direcciones IP a cada una de las interfaces virtuales en el Switch3 y Switch4. Los Switches deberán configurarles las interfaces FastEthernet 0/1-4 como troncales.

Configurar un dominio VTP llamado cisco para la propagación de cada una de las VLANs creadas en el Switch3 y ponerlo modo servidor VTP.

En el Switch1, Switch2, Switch4, Switch5, y Switch6 configurar VTP en modo cliente, y las interfaces a como se muestra en el cuadro llamada rango de puertos Access.

Configuración de HSRP

Configurar HSRP para cada interfaz de la VLAN ya antes creada en El Switch1 y Switch2. (Ver cuadro Switches Multicapa).

En el Switch3 la prioridad del grupo HSRP será de 120 (VLAN 10 Y 20), prioridad 100 (VLAN 30 Y 40). Para el Switch4 la prioridad HSRP será de 100 (VLAN 10 Y 20), prioridad 120 (VLAN 30 Y 40). Recuerde indicar a los Switches Multicapa que hagan el papel de ruteo con el comando IP routing.

Tiempo estimado de solución

➤ 4 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 HSRP, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. R2 es el Router activo en su implementación HSRP, y tiene la prioridad HSRP por defecto. Para que el Router R3 que espera para convertirse en el Router activo, lo que tiene que suceder es:
 - a) Prioridad de R3 debe establecerse en 105.
 - b) Prioridad de R3 se debe establecer en 10.
 - c) R3 deben configurarse con el comando preempt.
 - d) R3 tiene que reiniciarse.
2. ¿Cuántos mensajes timer hello por defecto se envía en HSRP?
3. ¿Cuáles son la dirección IP de origen y destino de los paquetes de saludo HSRP?
4. ¿Puedo configurar más de un grupo de reserva con el mismo número de grupo?
5. ¿Cuál será la dirección MAC es un Router virtual en el grupo HSRP 20?
6. ¿Son mensajes HSRP TCP o UDP?
7. ¿Es posible ejecutar HSRP entre dos Routers en dos interfaces diferentes?

PRÁCTICA N° 7: BALANCEO DE CARGA VRRP

OBJETIVOS

General

- Implementar el protocolo VRRP en un escenario de red en donde se encuentren salidas o gateways redundantes con un destino en común.

Específicos

- Establecer grupos VRRP.
- Asignar una IP virtual a cada grupo VRRP.
- Reducir la sobrecarga de tráfico en un Gateway o puerta de enlace.

INTRODUCCIÓN

En esta práctica, se deberá implementar el protocolo VRRP con algunas de sus funcionalidades que este ofrece. Se analizará el funcionamiento de este protocolo y la interacción de los dispositivos que se encuentren ejecutando esta tecnología.

Además del protocolo VRRP en esta práctica, se deberán implementar tecnologías básicas como Spanning Tree, NAT, VLANs entre otras.

REQUERIMIENTOS

- **Hardware**
 - Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.
- **Software**
 - Simulador Boson Netsim 8.0
 - 2 Routers 2811.
 - 2 Switches 3550.
 - 2 Switches 2960.
 - 4 PC.

CONOCIMIENTOS PREVIOS

Para dar marcha a la realización de esta práctica, se debe de dominar el tema que corresponde al funcionamiento del protocolo VRRP, esta información facilitada en el capítulo correspondiente a esta tecnología en este documento, en donde se abordan puntos importantes que en esta práctica se estarán configurando.

TOPOLOGÍA

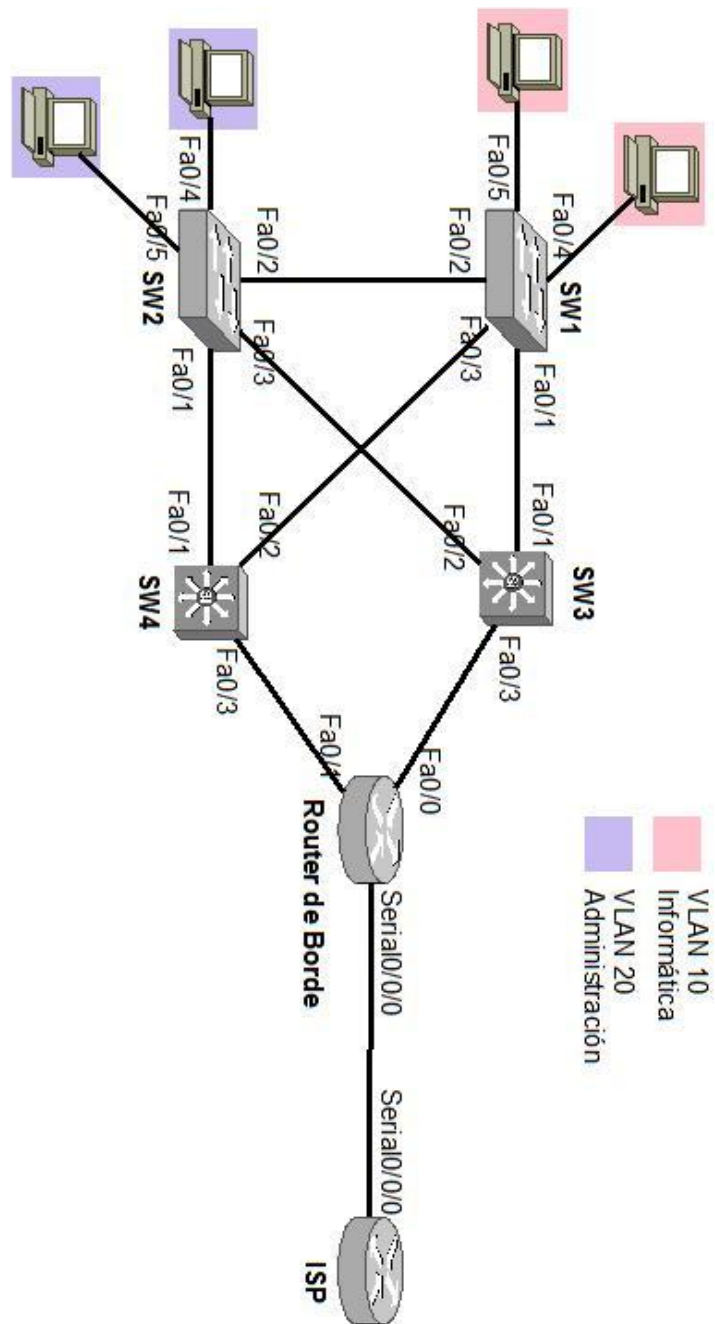


Figura 106. Topología de práctica de VRRP

FUNCIONALIDAD

La Figura 106 muestra una red en un ambiente con redundancia de puertas de enlaces para la salida a internet.

En este escenario, se encuentra segmentado en dos áreas que son: Informática y administración, por lo tanto existirán dos VLANs con los nombres antes dichos. También se configurará la tecnología Spanning Tree para evitar los bucles o inundaciones en la red.

Para la salida de los usuarios de cada VLANs a internet, se implementará NAT sobrecargado para las traducciones de direcciones.

COMANDOS DE AYUDA

Comandos	Descripción
show vrrp	Muestra la configuración que se ha establecido en un equipo miembro de un grupo VRRP.
vrrp {número de grupo} ip {dirección ip virtual}	Establece una ip virtual para un grupo VRRP
vrrp {número de grupo} priority {prioridad}	Establece la prioridad a un miembro de un grupo VRRP.
vrrp {número de grupo} shutdown	Deshabilita la configuración VRRP establecida.

DATOS DE LOS DISPOSITIVOS

Routers			
Equipo	Interfaz	Dirección IP	Dirección de Subred
Router de Borde	Serial0/0/0	215.15.15.1/24	215.15.15.0/24
	Fa0/0	192.168.3.2/30	192.168.3.0/30
	Fa0/1	192.168.4.2/30	192.168.4.0/30
ISP	Serial0/0/0	215.15.15.2/24	215.15.15.0/24

Switches Multicapa					
Equipo	Nombre y Número de VLAN	Interfaz	Dirección IP	Dirección Subred	Grupo VRRP IP virtual
SW3	Informática 10	VLAN 10	192.168.1.1/24	192.168.1.0/24	Grupo 1 192.168.1.100/24
	Administración	VLAN 20	192.168.2.1/24	192.168.2.0/24	Grupo 2 192.168.2.100/24
	-----	Fa0/3	192.168.3.1/24	192.168.3.0/30	-----
SW4	Informática 10	VLAN 10	192.168.1.2/24	192.168.1.0/24	Grupo 1

					192.168.1.100/24
	Administración	VLAN 20	192.168.2.2/24	192.168.2.0/24	Grupo 2 192.168.2.100/24
	-----	Fa0/3	192.168.4.1/30	192.168.4.0/30	-----

Switches de capa 2		
Equipo	Interfaces Access	ID de VLAN
Switch 1	Fa0/4 - 5	10
Switch 2	Fa0/4 - 5	20

ENUNCIADO

Creación y propagación de VLANs

Para dar inicio a esta práctica, en el Switch Multicapa llamado SW3 se deberán crear 2 VLANs, de modo que la red estará seccionada en 2 departamentos. También en el Switch Multicapa llamado SW4 se deberán habilitar interfaces virtuales correspondiente a las VLANs creadas en SW3. Se deberá establecer los parámetros para cada una de las interfaces virtuales (VLANs) activadas en ambos Switches. Estos datos se proveen en la tabla Switches Multicapa.

Para seguir con la tarea de esta práctica, usaremos VTP para la propagación de las VLANs creadas en SW3 a los demás Switches. En el Switch Multicapa SW3 se creará un dominio llamado TELEMATICA y lo pondremos a operar a modo server.

Para los Switches SW1, SW2, y SW4 se deben de configurar a modo cliente VTP para que estos operen bajo el mismo dominio que el servidor VTP. En este punto también es donde se define los puertos access en los Switches SW1 y SW2 (ver cuadro Switches de capa 2).

Método de encaminamiento de tráfico

En este apartado es donde configuraremos el método de encaminamiento del tráfico generados por todas las VLANs, para que haya comunicación entre VLANs diferentes, y para la salida a internet.

En el Router de borde agregar rutas por defecto para que el tráfico generado por las VLANs sea encaminado con destino al exterior satisfactoriamente.

Una ruta por defecto para que el tráfico que venga del exterior con destino a cualquiera de las VLANs también sea encaminado de manera correcta. Recuerde que por cada subred (VLAN) se establece una ruta por defecto.

Configuración de NAT

A como hemos hecho en algunas de las prácticas anteriores a esta, implementaremos la tecnología NAT que servirá para las traducciones de direcciones IP, de modo que los usuarios pertenecientes a la LAN privada puedan navegar

al exterior mediante una misma dirección pública. Para llevar a cabo lo antes mencionado se estará haciendo lo siguiente:

- Establecer de manera correcta las listas de acceso correspondiente a las direcciones que se estarán traduciendo.
- Indicar las interfaces en este equipo que estarán implementando NAT.
- Asocie la lista de acceso, especificando la interfaz de salida para la traducción.

Configuración de Spanning Tree

En esta sección debido a los enlaces redundantes en el nivel de enlace que presenta este escenario de red, se deberá configurar el protocolo de árbol de expansión (STP) para evitar anomalías no deseadas.

En el Switch multicapa llamado SW3 establecer Spanning Tree en modo PVSTP, también a este equipo configurarlo para que opere como puente raíz primario. De igual manera en el Switch Multicapa llamado SW4 se hará las mismas configuraciones, con la diferencia que este equipo se deberá configurar para que opere como puente raíz secundario.

Para cumplir correctamente lo antes descrito nos basaremos en la prioridad de cada uno de los equipos, de modo que al Switch SW3 se le establecerá una prioridad por debajo de todos los Switches de la red para que de manera predeterminada se elija como puente raíz primario. De igual manera se hará lo mismo en el Switch SW4, con la diferencia que a este se le establecerá una prioridad por debajo de los demás Switches pero por encima de la prioridad del Switch SW3.

Configuración de VRRP

Para cumplir con el objetivo principal de esta práctica, se deberá configurar el protocolo VRRP que es un protocolo para aprovechar la redundancia de puertas de enlaces (gateway) para usuarios con un destino en común (internet). Para una mayor comprensión de la funcionalidad de este protocolo consulte el capítulo correspondiente a esta tecnología.

En el Switch llamado SW3, se establecerán 2 grupos VRRP uno para cada interfaz de VLAN creada en este equipo. El grupo 1 se le establecerá una prioridad de 110 y al grupo 2 una prioridad de 100, esto con el fin de que la interfaz de la VLAN 10 en este equipo sea el maestro en este grupo, y la interfaz VLAN 20 sea un respaldo del grupo 2.

De igual manera se hará la misma configuración en el Switch SW4, se crearan los grupos en este caso grupo 1 y grupo 2, el grupo 1 para la interfaz VLAN 10 y grupo 2 para la interfaz VLAN 20, con la diferencia que la interfaz VLAN 10 en el grupo 1 de este equipo será un respaldo y su prioridad será de 100, mientras que la interfaz VLAN 20 será el maestro del grupo 2 y su prioridad será de 110.

Recuerde que para cada grupo se tiene que establecer una IP para el router virtual y esta será el gateway para los usuarios finales. Utilice los datos proporcionados en el cuadro Switches Multicapa para la realización de esta tarea y también apoyarse con el cuadro comandos de ayuda.

Tiempo estimado de solución

➤ 4 horas

Preguntas de Análisis

1. ¿Cuál es el objetivo de crear 2 grupos VRRP en cada Switch?
2. ¿Por qué para los usuarios finales se establece como puerta de enlace la dirección del router virtual y no la dirección de la interfaz lógica de la VLAN?
3. Ejecute el comando show vrrp en el Switch SW3. Haga un resumen de la información generada al ejecutar este comando.
4. ¿Cuál es la prioridad que se establece por defecto a un grupo VRRP?

PRÁCTICA N° 8: FRAME RELAY E INTERVLANS

OBJETIVOS

General

- Implementar la tecnología Frame Relay poniendo en práctica los conocimientos adquiridos en base a la teoría.

Específicos

- Establecer subinterfaces en los Routers que formen parte en este escenario de red para comunicación punto a punto Frame Relay.

INTRODUCCIÓN

Con la implementación y realización de esta práctica, se pretende que los estudiantes sean capaces de entender y poder aplicar Frame Relay Punto a Punto, además de comprender el funcionamiento de la misma.

REQUERIMIENTOS

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB

➤ Software

- Simulador Cisco Packet Tracer 6.0. con los siguientes elementos:
 - 9 Switches 2960.
 - 12 PCs
 - 3 Routers 2811.
 - 1 Nube PT.

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica el alumno deberá tener conocimientos de implementación de Frame Relay, ya que será de mucha importancia para realizar la misma.

TOPOLOGÍA

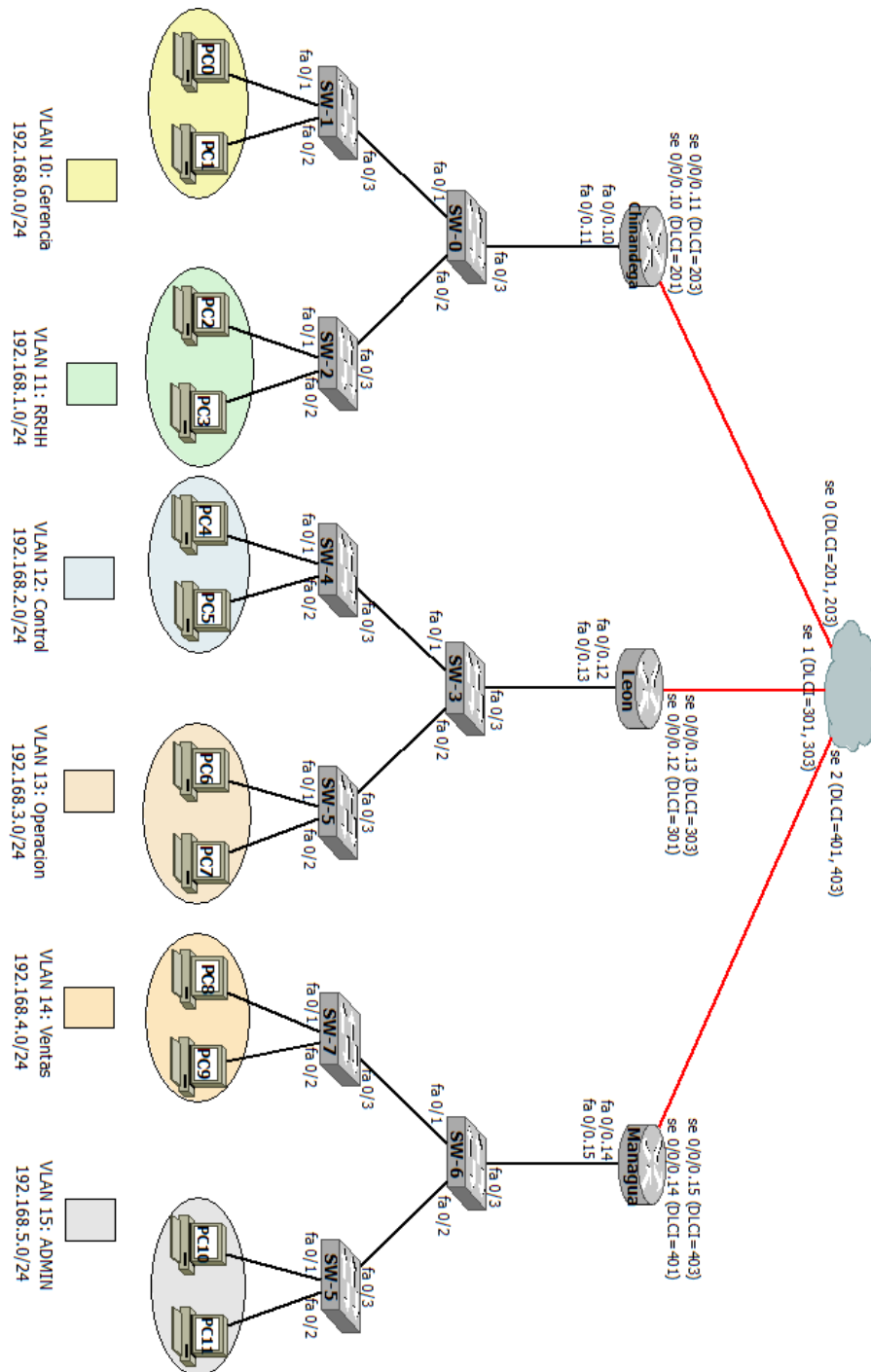


Figura 107. Topología de práctica de Frame Relay & InterVLANs

FUNCIONALIDAD

La Figura 107 muestra la topología en la cual se estarán aplicando Frame Relay. Además de que se estarán usando de manera conjunta con diferentes tecnologías y protocolos a como se menciona a continuación:

- Se deberán crear Seis VLANs, asignándoles direcciones IP dinámicamente a los equipos finales perteneciente a cada una de ellas mediante el protocolo DHCP.
- Habrá al menos un protocolo de enrutamiento (enrutamiento estático), para la comunicación entre las diferentes subredes.
- Se crearán en la interfaz FastEthernet de los Routers, sub-interfaces para comunicación entre VLANs distintas.
- Se crearán en la interfaz Serial de los Routers, sub-interfaces asignándole DLCI para cada sub-interfaz.
- Habrá que configurar la nube Frame Relay los DLCI que se configuró en las sub-interfaces de los Routers.

COMANDOS DE AYUDA

Comando	Descripción
show frame-relay lmi [<i>type number</i>]	Muestra información sobre la interfaz de gestión local (LMI)
show frame-relay map	Muestra las entradas del mapa Frame Relay actuales 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 interface-dlci <i>dlci</i>	Asigna el identificador de conexión de enlace de datos (DLCI) a una interfaz o sub-interfaz que se conectará a una red Frame Relay
[no] frame-relay inverse-arp [<i>protocolo</i>] [<i>dlci</i>]	Activa o desactiva Dirección inversa Resolución Protocol (ARP) en una interfaz
frame-relay lmi-type {ansi cisco q933a}	Selecciona el tipo de interfaz de gestión local (LMI)
frame-relay map <i>protocol protocol-address dlci</i> [broadcast] [ietf cisco]	Define la correspondencia entre una dirección de protocolo de destino y el DLCI utiliza para conectarse a la dirección de destino

DATOS DE LOS DISPOSITIVOS

Switch 0, Switch 3 y Switch 6				
Equipo	Nombre y Número de VLAN	Dominio VTP	Modo de operación VTP	Dirección de subred
SW0	Gerencia 10	Chinandega	Server	192.168.0.0/24
	RRHH 11	Chinandega	Server	192.168.1.0/24
SW3	Control 12	León	Server	192.168.2.0/24
	Operación 13	León	Server	192.168.3.0/24
SW6	Ventas 14	Managua	Server	192.168.4.0/24

	ADMIN 15	Managua	Server	192.168.5.0/24
--	----------	---------	--------	----------------

Puertos Access en los Switches		
Equipo	ID de VLAN	Puertos Access
Switch 1	10	Fa0/1 - 2
Switch 2	11	Fa0/1 – 2
Switch 4	12	Fa0/1 – 2
Switch 5	13	Fa0/1 – 2
Switch 7	14	Fa0/1 – 2
Switch 8	14	Fa0/1 - 2

Routers				
Equipo	Subinterfaz	Dirección IP	Dirección de subred	DLCI
Chinandega	Fa0/0.10	192.168.0.254/24	192.168.0.0/24	-----
	Fa0/0.11	192.168.1.254/24	192.168.1.0/24	-----
	Serial0/0/0.10	172.168.0.1/24	172.168.0.0/24	201
	Serial0/0/0.11	172.168.3.1/24	172.168.3.0/24	203
León	Fa0/0.12	192.168.2.254/24	192.168.2.0/24	-----
	Fa0/0.13	192.168.3.254/24	192.168.3.0/24	-----
	Serial0/0/0.12	172.168.3.2/24	172.168.3.0/24	301
	Serial0/0/0.13	172.168.2.1/24	172.168.2.0/24	303
Managua	Fa0/0.14	192.168.4.254/24	192.168.4.0/24	-----
	Fa0/0.15	192.168.5.254/24	192.168.5.0/24	-----
	Serial0/0/0.14	172.168.2.2/24	172.168.2.0/24	401
	Serial0/0/0.15	172.168.0.2/24	172.168.0.0/24	403

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 subinterfaces Serial, al igual que su DLCI correspondiente a cada una. (Ver tabla Routers).

Indicar que en las subinterfaces seriales de los Routers operen en modo Frame Relay punto a punto, e indicar la encapsulación Frame Relay para cada subinterfaz serial en estos equipos.

Creación y propagación de VLANs

Crear las VLANs en los dispositivo llamado SW-0, SW-3, SW-6 (ver cuadro SW-0, SW-3, SW-6). A los Switches antes mencionados se les deberán configurar las interfaces FastEthernet 0/1-3 como troncales.

Configurar un dominio VTP en los Switches 0, 3, y 6 con su respectivo nombre a como se indica en el cuadro llamado Switches 0, 3, y 6. Recuerde establecer estos en modo server VTP.

En los Switches SW1, SW2, SW4, SW5, SW7, SW8 configurarlos VTP modo cliente y las interfaces a como se muestra en la tabla (ver cuadro puertos Access en los Switches).

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 12,13(una por cada VLANs).

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 14,15(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 10,11 (una por cada VLANs).

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 14,15 (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 10,11 (una por cada VLANs).

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 12,13 (una por cada VLANs).

Configuración DHCP

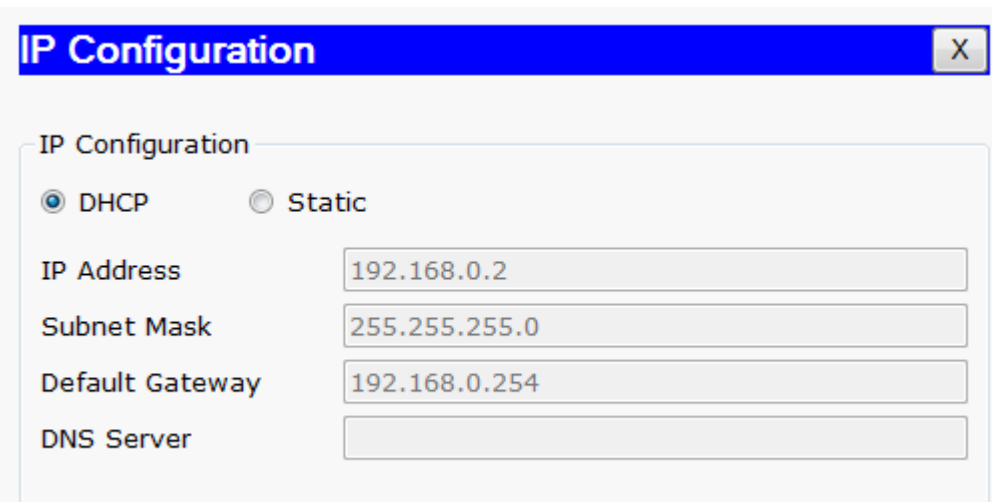
En el Router llamado Chinandega establecer 2 pools llamados Gerencia y RRHH indicando la dirección de red para cada uno, en este caso la dirección de red de la VLAN Gerencia y la VLAN RRHH, asignar un default Router para cada uno en este caso la dirección IP de la Subinterfaz Fa0/0.10 y Fa0/0.11.

En el Router llamado León establecer 2 pools llamados Control y Operación indicando la dirección de red para cada uno, en este caso la dirección de red de la VLAN Control y la VLAN Operación, asignar un default Router para cada uno en este caso la dirección IP de la Subinterfaz Fa0/0.12 y Fa0/0.13.

En el Router llamado Managua establecer 2 pools llamados Ventas y ADMIN indicando la dirección de red para cada uno, en este caso la dirección de red de la VLAN Ventas y la VLAN ADMIN, asignar un default Router para cada uno en este caso la dirección IP de la Subinterfaz Fa0/0.14 y Fa0/0.15.

Indicar a los equipos finales que deberán realizar una petición de DHCP, para que estos obtengan su dirección IP, la cual deberá ser perteneciente a la VLAN a la que cada uno de estos pertenezca.

En todos los equipos se deberán configurar que los equipos realicen una petición de DHCP para que el pool que hemos creado asigne dirección de manera dinámica a cada uno de los equipos, lo que se tendrá que hacer en todos los equipos se muestra en la siguiente figura:



Verificar que las direcciones IP han sido asignadas a cada uno de los equipos finales mediante DHCP.

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.

Nota: será de mucha importancia al asociar un nombre lógico con cada DLCI en la configuración de la nube Frame Relay fijarse bien las IP de las subinterfaces seriales configuradas en cada Routers que están asociadas con un DLCI para correcta comunicación esto es muy importante para configurar el siguiente punto.

Ejemplo: La interfaz serial 0 de la nube Frame Relay para el DLCI = 201 DLCI que es un DLCI de una subinterfaz serial del Router Chinandega fijarnos qué dirección IP tiene la subinterfaz serial de este Router con éste DLCI y que otro

Router tiene una IP en la misma Red en su subinterfaz serial en este caso Managua así que el nombre asociado a este DLCI en la nube Frame Relay en el puerto Serial 0 podría ser Chinandega-Managua.

Para el caso de la interfaz serial2 de la nube Frame Relay que está conectado directamente al Router2 el nombre sería Managua-Chinandega así asociamos comunicación del Router Chinandega con Managua y viceversa.

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.

Recuerde La nota agregada en para la interfaz Serial 0 al igual del Ejemplo también expuesto anteriormente

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.

Recuerde La nota agregada en para la interfaz Serial 0 al igual del Ejemplo también expuesto para ese caso.

Ordenando puertos seriales de la nube Frame Relay para correcta comunicación

En la pestaña config de la nube Frame Relay ir al menú llamado Frame Relay le aparecerá cuatro pequeños menús desplegables el primero y tercer menú muestra nombre de las seriales y el segundo y cuarto menú los nombres que agregamos como asociación de los DLCI configurados en el paso anterior.

seleccionemos el primer menú con un nombre de serial y por defecto el segundo menú mostrara un nombre el cual asociamos a esa interfaz serial, seleccionemos luego en el tercer menú una interfaz serial correcta de acuerdo al seleccionado en el primer menú y por defecto el cuarto menú mostrara un nombre asociado a ese puerto también.

Ejemplo:

1. El primer menú seleccionamos el puerto serial0.
2. El segundo menú muestra el nombre Chinandega-Managua
3. El tercer menú seleccionamos un puerto serial que por defecto haga que el cuarto menú muestre un nombre relacionado con el nombre del segundo menú, este nombre podría ser Managua-Chinandega
4. Luego agregara todo lo seleccionado.

Realizar estos pasos hasta que tenga asociado comunicación de Chinandega a Managua y viceversa, Chinandega a León y viceversa, León a Managua y viceversa.

Tiempo estimado de solución

➤ 4 horas

Preguntas de análisis

En este apartado se pretende que los estudiantes sean capaces de analizar y responder las siguientes cuestiones en base al tema de Frame Relay, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Qué se requiere para configurar múltiples DLCI en una interfaz de Router individual?
 - a) Una sola dirección de red
 - b) Múltiples LMI
 - c) Sub-interfaces
 - d) Múltiples direcciones IP en la misma subred
2. Al utilizar Frame Relay, ¿qué comando muestra tanto la información DLCI como la información LMI en la interfaz Serial 0?
 - a) show interface serial 0
 - b) show interface frame-relay
 - c) show frame-relay serial 0
 - d) show frame-relay information
3. En una conexión Frame Relay punto a punto, ¿qué requiere cada conexión punto a punto?
 - a) Un tipo de encapsulamiento único
 - b) Una sub-interfaz separada
 - c) Una dirección IPX
 - d) Una conexión física directa
4. ¿Cuál es la ventaja de un servicio público de Frame Relay?
 - a) El servicio es gratuito
 - b) Un proveedor de servicio administra el equipo de la red
 - c) Elimina la necesidad de comprar equipo
 - d) Elimina la necesidad de DLCI
5. ¿Qué comando de Router muestra las tablas Frame Relay?
 - a) show interface
 - b) show frame-relay route
 - c) show ip route
 - d) show frame-relay map

PRÁCTICA N° 9: LISTAS DE CONTROL DE ACCESO

OBJETIVOS

General

- Aplicar Listas de Accesos fusionándolas con diferentes tecnologías y protocolos.

Específicos

- Establecer Listas de Accesos para dar seguridad a una red LAN.
- Mostrar la importancia de las listas de acceso para permitir y denegar tráfico.
- Conocer la importancia de la máscara WildCard a la hora de hacer uso de Listas de Accesos.

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) según su tipo y finalidad de uso, además de comprender el funcionamiento de las mismas para dar seguridad a los equipos de una red LAN ante una red externa. De igual manera se pretende que conozcan la importancia que tiene la máscara WildCard al momento de la implementación de las Listas de Accesos.

REQUERIMIENTOS

Para la realización de esta se requerirá de lo siguiente:

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.

➤ Software

- Simulador Cisco Packet Tracer 6.0 con los siguientes elementos:
 - 8 Servidores
 - 3 Routers 1841
 - 8 Switches 2950
 - 1 Switch Multicapa 3560
 - 9 PCs

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica es necesario que el estudiante tenga conocimientos básicos acerca del uso y funcionamiento de las diferentes tipos de Listas de Control de Acceso proporcionado previamente en los aspectos teóricos, además conocimientos básicos de los siguientes temas:

- Direccionamiento IP (Subnetting).
- Creación y propagación de VLANs.
- Asignación dinámica de direcciones IP mediante DHCP.

- Rutas estáticas.
- Implementación de protocolos de enrutamiento (RIP).
- NAT.

TOPOLOGÍA

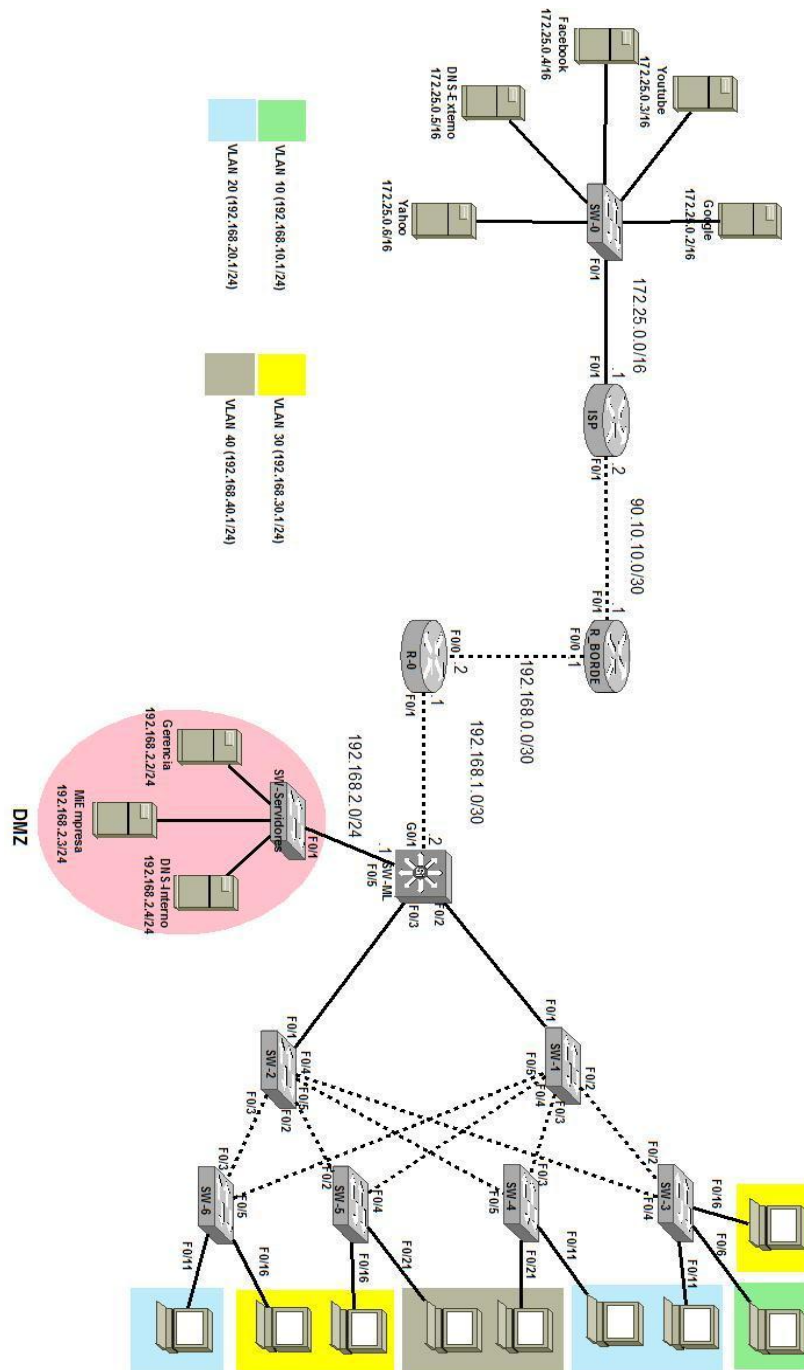


Figura 108. Topología de práctica de Listas de Control de Acceso

FUNCIONALIDAD

La Figura 108 muestra la topología en la cual se estarán aplicando diferentes tipos de Listas de Control de Acceso (ACL), según su finalidad de uso. Además de que se estarán usando de manera conjunta con diferentes tecnologías y protocolos:

- Se deberán crear cuatro VLANs, asignándoles direcciones IP dinámicamente a los equipos finales perteneciente a cada una de ellas mediante el protocolo DHCP, agregando la dirección del Servidor DNS interno que poseerán los equipos de la red LAN.
- Habrá al menos un protocolo de enrutamiento (RIP), para la comunicación entre las diferentes redes pertenecientes a la LAN, así como también entre las VLANs que serán creadas.
- Se crearán rutas estáticas para dar encaminamiento de tráfico a los equipos pertenecientes a la LAN.
- Aplicar NAT para la traducción de direcciones IP provenientes de la LAN al exterior.
- Los equipos finales deberán acceder a cada uno de los servidores pertenecientes a la red LAN y WAN mediante DNS.
- Por último y como objetivo principal de la práctica, se deberán crear Listas de Control de Acceso para dar seguridad a cada uno de los equipos perteneciente a la red interna, manteniendo integridad y seguridad en el control de flujo de tráfico en la red LAN así como también hacia la red externa o WAN.

COMANDOS DE AYUDA

Comando	Descripción
access-list access-list-number {deny permit} protocol source source-wildcard [operator [port]] destination destination-wildcard [operator [port]]	Define una lista Extendida IP de control de acceso (ACL) para el tipo de tráfico especificado por el parámetro de protocolo
dns-server [ip servidor dns]	configura la dirección de un servidor DNS
encapsulation frame-relay	
interface fastEthernet [número de interfaz]	Accede a la interfaz especificada
interface range fastEthernet [número de interfaces]	Accede al rango de interfaz especificados
ip access-group {access-list-number access-list-name} {in out}	Controla el acceso a una Interfaz
ip address [ip-address subnet-mask]	Asigna dirección IP a una interfaz
ip dhcp pool [nombre del pool]	Crea un POOL para asignar direcciones de manera dinámica

ip route [red/mascara siguiente salto]	
ip routing	Activa ruteo en un switch multicapa
network [network-address]	activa el protocolo de enrutamiento especificado en la red especificada
no shutdown	Habilita una Interfaz
no switchport	Habilita la interfaz de un Switch multicapa a modo de ruteo
router rip	Habilita el enrutamiento RIP en un dispositivo
show access-lists [access-list-number access-list-name]	Muestra el contenido de las ACL actuales
show ip route	Muestra la tabla de enrutamiento IP
show vlan	muestra información de las VLANs
show vlan brief	muestra los parámetros para todas las VLAN, contiene el nombre, el estado y los puertos que le asigna la VLAN
switchport access	Establece una interfaz modo acceso
switchport mode trunk	Establece una interfaz modo troncal
vlan [número-vlan name nombre-vlan]	crea una VLAN especificando el nombre y el número
vlan database	Accede a la base de datos de VLANs
vtp domain [nombre del dominio]	Crea un dominio VTP
vtp mode client	Establece el dominio VTP como modo Cliente
vtp mode server	Establece el dominio VTP como modo Servidor

DATOS DE LOS DISPOSITIVOS

VLANs						
Nombre y número	Dirección de red	Dirección IP	Interfaces de Acceso			
			SW-3	SW-4	SW-5	SW-6
VLAN 10 (RRHH)	192.168.10.0/24	192.168.10.1/24	F0/6-10	F 0/6-10	F 0/6-10	F 0/6-10
VLAN 20 (Contabilidad)	192.168.20.0/24	192.168.20.1/24	F0/11-15	F0/11-15	F0/11-15	F0/11-15
VLAN 30 (Tesorería)	192.168.30.0/24	192.168.30.1/24	F0/16- 20	F0/16- 20	F0/16- 20	F0/16- 20
VLAN 40 (Gerencia)	192.168.40.0/24	192.168.40.1/24	F0/21-24	F0/21-24	F0/21-24	F0/21-24

Switch Multicapa (SW_ML)			
Nombre	Interfaz	Dirección IP	Dirección de subred
SW_ML	Fa0/5	192.168.2.1/24	192.168.2.0/24
	GigabitEthernet0/1	192.168.1.2/30	192.168.1.0/30

Routers		
Nombre	Direcciones IP	
	FastEthernet 0/0	FastEthernet 0/1
R-0	192.168.0.2/30	192.168.1.1/30
R_Borde	192.168.0.1/30	90.10.10.1/30
ISP	172.25.0.1/30	90.10.10.2/30

Switches de capa 2		
Nombre	Puertos Troncales	Puerto Access
SW0	--	--
SW-Servidores	--	--
SW-1	0/1-5	--
SW-2	0/1-5	--
SW-3	0/1-5	0/6-24
SW-4	0/1-5	0/6-24
SW-5	0/1-5	0/6-24
SW-6	0/1-5	0/6-24

Servidores			
Nombre	Dirección IP	Gateway	Nombre de Dominio
Externos			
Google	172.25.0.2/16	172.25.0.1	www.google.com
YouTube	172.25.0.3/16	172.25.0.1	www.youtube.com
Facebook	172.25.0.4/16	172.25.0.1	www.facebook.com
DNS-Externo	172.25.0.5/16	172.25.0.1	--
Yahoo	172.25.0.6/16	172.25.0.1	www.yahoo.com
Internos			
Gerencia	192.168.2.2/24	192.168.2.1	www.gerencia.com
MiEmpresa	192.168.2.3/24	192.168.2.2	www.miempresa.com
DNS-Interno	192.168.2.4/24	192.168.2.2	--

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:

- ISP.
- R_Borde.
- R0.
- SW-ML.

Asignar direcciones IP de estática a cada uno de los servidores interno y externos a como se aprecia en la tabla llamada servidores.

Recuerde establecer correctamente las direcciones IP para un funcionamiento satisfactorio.

Creación y propagación de VLAN

Crear las VLANs en el dispositivo llamado SW-ML (ver cuadro VLANs), asignándoles las direcciones IP a cada una. Al Switch se le deberá configurar las interfaces FastEthernet 0/2 y 0/3 como troncales.

Configurar en el dispositivo SW-ML un dominio VTP llamado miempresa para la propagación de cada una de las VLANs dentro la red LAN y ponerlo modo servidor.

En los Switch SW-1 y SW-2 configurar el dominio VTP modo cliente y las interfaces a como se muestra en el cuadro llamado Switches de capa 2.

En los Switches SW-3, SW-4, SW-5 y SW-6 configurar el dominio VTP modo cliente y las interfaces a como se muestra en el cuadro llamada Switches de capa 2.

En el Switch SW-ML visualizar si las VLANs fueron creadas correctamente.

En cualquiera de los switch SW-3, SW-4, SW-5 o SW-6, visualice si las VLANs fueron propagadas correctamente y verifique las interfaces que cada una de estas tienen asignadas.

Asignación de direcciones IP mediante DHCP a los equipos finales

Configurar cuatro pools llamados cada uno con el nombre de las VLANs que han sido creadas anteriormente, indicando el rango de direcciones IP la cual se estarán asignando de manera dinámica a cada uno de los equipos finales que realice una petición DHCP, de igual manera se deberá indicar el servidor DNS que tendrá cada uno de los equipos finales, con el fin de que estos puedan acceder a los servidores externos e internos por medio de los nombres de dominios.

Indicar a los equipos finales que deberán realizar una petición DHCP, para que estos obtengan su dirección IP, la cual deberá ser perteneciente a la VLAN a la que cada uno de estos pertenezcan.

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.

Verificar que las direcciones IP han sido asignadas a cada uno de los equipos finales mediante DHCP

Configuración de protocolo de enrutamiento (RIP)

Para encaminar el tráfico entre cada una de las subredes existente de la LAN, se estará usando RIP como protocolo de enrutamiento, y debe ser configurado en los equipos que se mencionan a continuación:

En el router R0 configurar RIP para que éste anuncie las subredes que conozca o estén en su tabla de enrutamiento a los demás routers.

De igual manera se hará en el router R_Borde para anunciar las subredes que este conozca incluyendo la subred DMZ.

Verifique en las tablas de enrutamiento de cada router antes configurado si existen las subredes anunciada por cada uno.

Configuración de protocolo de enrutamiento (Estático)

Se debe definir rutas para obligar a los paquetes a tomar una ruta determinada entre su origen y su destino, estas rutas serán definidas en los siguientes equipos:

En el Switch SW-ML, crear una ruta estática dándole un camino determinada a todas las subredes internas que dependen de este equipo.

En el Router R0 indicar la ruta que seguirán los paquetes provenientes de SW-ML, hacia el R_Borde que posee la LAN.

En el Router R0 crear una ruta estática para cada una de las subredes pertenecientes en la LAN (VLANs, DMZ), indicándole una ruta de retorno a estas.

En el R_Borde, crear dos rutas estáticas, dándole así una ruta alternativa a todos los paquetes que provengan de la red WAN hacia el interior, así como los provenientes de la red LAN al exterior.

Traducción de direcciones IP (NAT Sobrecargado)

Definir listas de acceso IP estándar que permita a las direcciones locales internas que se deben traducir.

Asocie la lista de acceso, especificando la interfaz de salida.

Especificar la interfaz interna y marcarla como conectada al interior.

Especificar la interfaz externa y marcarla como conectada al exterior.

Visualizar si las listas de acceso han sido creadas correctamente.

De manera gráfica verificar si las direcciones IP de nuestra red LAN están siendo traducidas, enviando un mensaje de cualquiera de los equipos finales de nuestra red LAN hacia los servidores, analizando cada uno de los paquetes para hacer la verificación.

Creación de las Listas de Control de Acceso (ACL)

Crear las Listas de Control de Acceso que se desean implementar detallando el tipo que se estará usando y la interfaz del dispositivo a la cual se le estará aplicando:

➤ Para los Servidores externos:

- Gerencia Acceso a todos los servidores Externos.
- Tesorería, Contabilidad tendrán acceso al servidor de Facebook.
- RRHH, y una máquina de contabilidad acceso a Youtube.
- Todos los equipos tendrán acceso a Google, excepto los de la VLAN de tesorería.
- Permitir que todos los equipos de la DMZ puedan acceder a todos los servidores externos.
- Crear una lista de acceso denegando cualquier otro tipo de tráfico.

➤ Para los servidores Internos

- Acceso de todos los equipos al servidor mi Empresa.
- Solo gerencia y tesorería accederán al servidor gerencia.
- Solo gerencia podrá hacer transferencia de ficheros por medio de FTP.
- Crear una lista de acceso denegando cualquier otro tipo de tráfico.

Implementación de las Listas de Control de Acceso (ACL)

Implemente las Listas de Control de Acceso creadas en el apartado anterior en los interfaces de los dispositivos que crea conveniente y verifique su correcto funcionamiento realizando las debidas pruebas.

Tiempo estimado de solución

➤ 4 horas

Prueba de análisis

1. ¿Cuáles son las principales ventajas de crear e implementar Listas de Control de Acceso?
2. Según la figura que representa la topología de esta práctica ¿En cuál o cuáles de los equipos tubo que crear e implementar las Listas de Control de Acceso?
3. ¿Cuantos tipos de Listas de Control de Acceso uso para dar solución a la práctica e ilustre donde las implementa?
4. ¿Qué nuevas listas de las Listas de Control de Acceso cree usted que sea necesario implementar y por qué?

PRÁCTICA N° 10: TELEFONÍA SOBRE IP (VoIP)

OBJETIVOS

General

- Crear una central telefónica VoIP.

Específicos

- Configurar el Call Manager Express en un router Cisco serie 2811.
- Conectar y comunicar teléfonos IP de redes diferentes.

INTRODUCCIÓN

La implementación y realización de esta práctica, está orientada al desarrollo de una plataforma de VoIP, con base en el diseño e implementación de un sistema de comunicación de voz sobre IP a través de varias redes conectadas mediante los protocolos TCP/IP.

Se pretende que los estudiantes sean capaces de entender la sopa de protocolos que se ven involucrados en la transmisión de voz sobre IP y puedan realizar el ejercicio detallado a continuación.

REQUERIMIENTOS

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.

➤ Software

- Simulador Cisco Packet Tracer 6.0 con los siguientes elementos:
 - 4 Routers serie 2811.
 - 5 Switches 2950-24.
 - 9 Teléfonos IP.
 - 6 PCs.
 - 2 Servidores.
 - Cable directo.
 - Cable serial DTE.

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica es necesario que el estudiante tenga conocimientos básicos del funcionamiento de los protocolos que usa VoIP a la hora que este entra en funcionamiento (ver Capítulo 2, Telefonía sobre VoIP), además conocimientos básicos de los siguientes temas:

- Creación de VLANs.
- Asignación dinámica de direcciones IP mediante DHCP.
- Implementación de protocolos de enrutamiento (RIP).

- Hacer uso de Sub-Interfaces.

TOPOLOGÍA

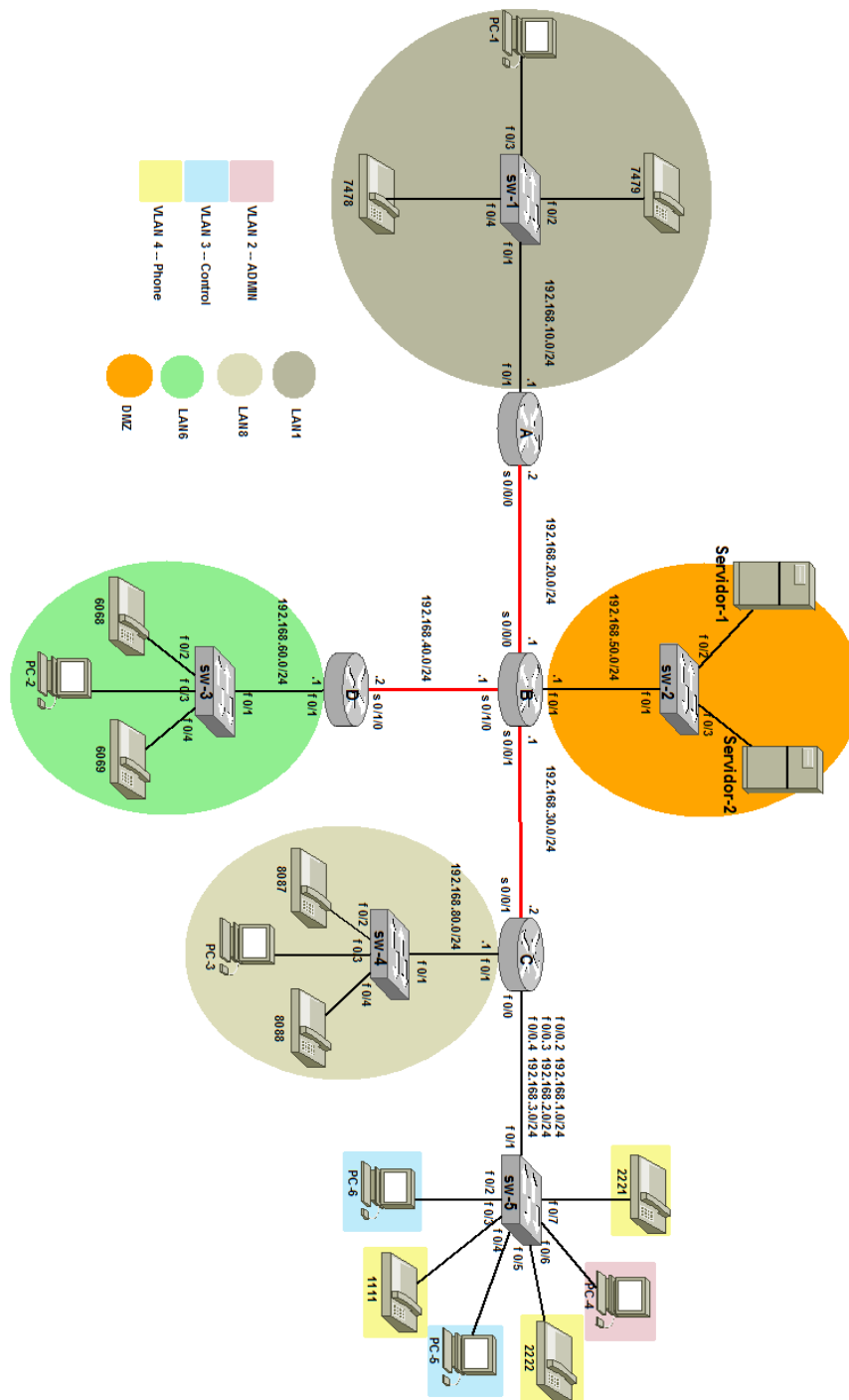


Figura 109. Topología de práctica de Telefonía sobre VoIP

FUNCIONALIDAD

La Figura 109 muestra la topología en la cual se estarán creando diferentes centrales telefónicas y comunicándolas entre sí. Se deberá configurarse un protocolo de enrutamiento (RIP), para la comunicación entre las diferentes redes existentes. Todas las centrales telefónicas deben comunicarse entre sí, por lo que se deberá de configurar el grupo de número de destinatarios en cada uno de los routers en los que fueron creados cada uno de los servicios telefónicos.

Existe una DMZ en la cual todos los equipos (PCs y Teléfonos IP) deberán tener conexión con ambos servidores.

COMANDOS DE AYUDA

Comandos	Descripción
auto assign [número] to [número]	La asignación automática de números de extensión a los botones
default router [dirección IP]	Dirección IP del router por defecto
ephone-dn [número]	Definición de la entrada del directorio
interface fastEthernet [número de interfaz]	Accede a la interfaz especificada
ip address [dirección-IP Máscara-de-Red]	Asigna dirección IP a una interfaz
ip dhcp pool [nombre del pool]	Crea un POOL para asignar direcciones de manera dinámica
ip source-address [dirección-ip] port [número-de-puerto]	
max-dn [número]	Definir el número máximo de números de la guía
max-ephones [número]	Define el máximo número de teléfonos
network [dirección-de-Red]	Activa el protocolo de enrutamiento especificado en la red especificada
no shutdown	Habilita una Interfaz
number [número]	Asigne el número de teléfono a la entrada especificada
option [código-dhcp] ip [dirección-IP]	Opciones de configuración de DHCP, obligatorio para la configuración de VoIP
switchport access	Establece una interfaz modo acceso
switchport voice vlan [número-VLAN]	Definir la VLAN en la que los paquetes de voz serán manejados
telephony-service	Configuración del router para servicios de telefonía
vlan [número-vlan name nombre-vlan]	Crea una VLAN especificando el nombre y el número

DATOS DE LOS DISPOSITIVOS

Routers				
Nombre	Dirección IP			
	F 0/1	S 0/0/0	S 0/0/1	S 0/1/0
A	192.168.10.1/24	192.168.20.2/24	--	--
B	192.168.50.1/24	192.168.20.1/24	192.168.30.1/24	192.168.40.1/24
C	192.168.80.1/24	--	192.168.30.2/24	--
D	192.168.60.1/24	--	--	192.168.40.2/24

Subinterfaces			
Nombre	Dirección de Sub-interfaces		
	F 0/0.2	F 0/0.3	F 0/0.4
C	192.168.1.1/24	192.168.2.1/24	192.168.3.1/24

Servidores		
Nombre	Dirección IP	Gateway
Servidor1	192.168.50.2/24	192.168.50.1
Servidor2	192.168.50.3/24	

VLANs			
Número	Nombre	Dirección IP	Interfaces
2	Control	192.168.1.1/24	F 0/2, F 0/4
3	ADMIN	192.168.2.1/24	F 0/6
4	Phones	192.168.3.1/24	F 0/3, F 0/5, F 0/7

Switches						
Interfaces FastEthernet	Configuraciones	Nombres				
		SW-1	SW-2	SW-3	SW-4	SW-5
0/1	Modo	access	access	access	access	access
	N° VLAN	1	1	1	1	1
0/2	Modo	voice		voice	voice	access
	N° VLAN	1		1	1	3
0/3	Modo					voice
	N° VLAN					4
0/4	Modo	voice		voice	voice	access

	N° VLAN	1		1	1	3
0/5	Modo					voice
	N° VLAN					4
0/6	Modo					access
	N° VLAN					2
0/7	Modo					voice
	N° VLAN					4

ENUNCIADO**Asignación de direcciones IP**

Asignar dirección IP a las interfaces de los siguientes equipos:

- Router A.
- Router B.
- Router C.
- Router D.

Asignar direcciones IP estática a los servidores a como se detalla en el cuadro servidores.

Creación de VLANs

En el SW-5 crear las VLANs a como se detallan en el cuadro nombrado VLANs

Asignar direcciones IP a cada una de las sub-interfaces en el router C para que cada una de las VLANs pueda tomar sus respectivas direcciones IP.

Configuración de puertos para los Switches

En los Switches sw-1, sw-3, sw-4 y sw-5 configurar las interfaces a como se especifica en el cuadro llamado Switches.

- Sw-1
- Sw-3
- Sw-4
- Sw-5

Configurar las interfaces para las PCs pertenecientes a sus diferentes VLANs.

Configurar las interfaces para los Teléfonos IP que pertenecerán a la VLAN Phone.

Asignación de direcciones IP mediante DHCP a los teléfonos IP y a las PC

Configurar un pool en los router A, C y D con sus respectivos nombres, indicando el rango de direcciones IP la cual se estarán asignando de manera dinámica a cada uno de los equipos finales (Teléfonos IP y PCs) que realice una petición DHCP.

Nota:

- Se debe tomar en cuenta que se necesitara asignar direcciones IP de manera dinámica a los Teléfonos IP, se necesitara crear un pool para cada VLANs en el switch Sw-5, asignándoles las direcciones IP establecidas en la tabla VLANs
 - Router A
 - Router C
 - Router D

Indicar a las PCs que deberán realizar una petición de DHCP, para que obtengan su dirección IP de manera dinámica.

Poniendo el cursor sobre cada uno de los teléfonos IP, verificar si estos han obtenido su dirección IP de acuerdo al pool que fue configurado en cada uno de sus redes.

Configuración de protocolo de enrutamiento (RIP)

Para encaminar el tráfico entre cada una de las redes, se estará usando RIP como protocolo de enrutamiento, y debe ser configurado en los siguientes equipos:

- Router A
- Router B
- Router C
- Router D

Una vez establecido RIP como protocolo de enrutamiento en los Routers, debe existir conexión entre los equipos de las diferentes redes, verificar si dicha conexión existe.

Configurar el servicio de Telefonía IP

Configurar el Administrador de llamadas de servicio expreso de telefonía IP en los routers A, C y D; lo que permitirá habilitar VoIP.

- Router A
- Router C
- Router D

Configuración de la guía telefónica de los teléfonos IP

Aunque los teléfonos IP ya están conectados, se necesita configuración complementaria antes que sea capaz de comunicarse. Es necesario configurar en los routers, que se le estableció el servicio de telefonía con CME para asignar un número a cada uno de los teléfonos IP.

- Router A
- Router C
- Router D

Ahora es el turno de verificar si los teléfonos tomaron su número y dirección IP de manera correcta.

Nota: Si tiene apagando y/o desconectado el teléfono IP recordar conectarlo y/o encenderlo.

Verificar si se comunican teléfonos de la misma central telefónica. Por ejemplo realizar una llamada del teléfono IP que tiene al número 6068 al que posee el número 6069.

Ya existe la comunicación entre los teléfonos IP entre las mismas centrales telefónicas, es el turno de configurar para que todos los teléfonos puedan comunicarse entre sí.

Configuración para que todos los teléfonos se puedan comunicar entre si

Ya existe la comunicación entre los teléfonos IP entre las mismas centrales telefónicas, es el turno de configurar para que todos los teléfonos de diferentes centrales telefónicas puedan comunicarse entre sí.

- Router A
- Router C
- Router D

Tiempo estimado de solución

- 4 horas

Preguntas de análisis

1. Explique cuál es la ventaja principal que en una empresa se use telefonía sobre VoIP que la telefonía que es usada convencionalmente.
2. A la hora de crear las centrales telefónicas, es de gran relevancia crear una por cada VLAN existente o es posible crear una sola central para toda la empresa. Explique las ventajas y desventajas de ambos escenarios.

PRÁCTICA N° 11: PROTOCOLOS DE ENRUTAMIENTO CON DIRECCIONAMIENTO IPv4

OBJETIVOS

General

- Configurar protocolos de enrutamiento dinámico y entender su interrelación con otros protocolos.

Específicos

- Escribir la configuración de protocolos de enrutamiento internos dentro de diferentes sistemas autónomos.
- Establecer y configurar el protocolo BGP en los enrutadores de borde de cada uno de los sistemas autónomos.
- Estudiar e implementar la redistribución de rutas en los enrutadores de borde, para poder comunicar los diferentes sistemas autónomos.

INTRODUCCIÓN

Con la implementación y realización de esta práctica, se pretende que los estudiantes sean capaces de entender y poder configurar los diferentes tipos de protocolos de enrutamiento dinámico internos (EIGRP, RIP, IS-IS, OSPF), y a su vez comunicar los diferentes sistemas autónomos a través del protocolo de enrutamiento externo BGP. Este protocolo se configurará en el/los Routers de borde de cada Sistema Autónomo siendo este/estos, los Routers que funcionarán con dos tipos de protocolo de enrutamiento: BGP y el protocolo de enrutamiento interno que corresponda con el sistema autónomo.

REQUERIMIENTOS

Para la realización de esta se requerirá de lo siguiente:

- **Hardware**
 - Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.
- **Software**
 - Emulador de redes GNS3 0.8.4.
 - 10 Routers de la serie 3640.
 - 4 PCs.

CONOCIMIENTOS PREVIOS

Para la correcta realización 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 a su vez del protocolo de comunicación entre diferentes sistemas autónomos (BGP) de Acceso proporcionado previamente en los aspectos teóricos.

TOPOLOGÍA

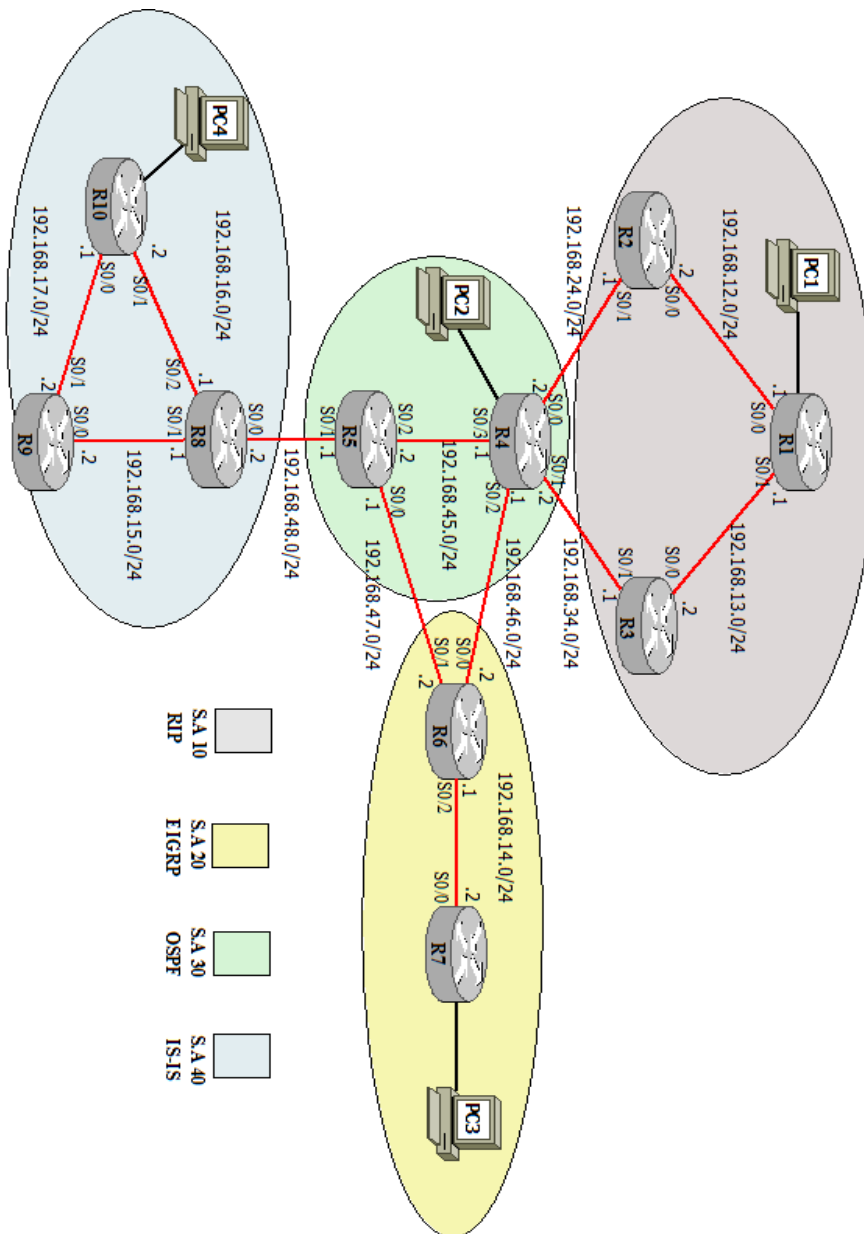


Figura 110. Topología de práctica de protocolos de enrutamiento con direccionamiento IPv4

FUNCIONALIDAD

La Figura 110 muestra la topología en la cual se estarán aplicando diferentes tipos de protocolos de enrutamiento dinámico fusionándolo con BGP

- Se deberá establecer un protocolo de enrutamiento dinámico para cada Sistema Autónomo como lo muestra la topología.

- En los Routers de Borde de cada Sistema Autónomo se configurará BGP para la comunicación entre diferentes Sistemas Autónomos.
- En los Routers de Borde de cada Sistema Autónomo se redistribuirá rutas entre BGP y el protocolo de enrutamiento que se estableció en ese Sistema Autónomo.

COMANDOS DE AYUDA

Comando	Descripción
Comandos de RIP	
router rip	Permite enrutamiento de Routing Information Protocol (RIP)
version 2	Permite RIP versión 2
Comandos de OSPF	
Router ospf proceso-id	Entra en el modo de configuración del Router para un proceso de OSPF
network direccion mascara-wilcard area area-id	Define las interfaces que ejecutan OSPF y que las zonas que operan
log-adjacency-changes	Traza los cambios en estado de adyacencia
Comandos de EIGRP	
router eigrp número-sistema-autónomo	Entra en el modo de configuración del Router para EIGRP
redistribute protocolo metric banda ancha- delay reliability carga MTU	Configura la redistribución en EIGRP
Comandos de IS-IS	
router isis	Entra en el modo de configuración del Router para IS-IS
ip router isis	Permite IS-IS en una interfaz
net direccion-net	Configura Título Entidad de Red (NET) para IS-IS
Comandos de BGP	
router bgp número-sistema-autónomo	Entra en el modo de configuración del Router para BGP
bgp log-neighbor-changes	Habilita el registro de los cambios de adyacencia de vecinos para monitorear la estabilidad del sistema de enrutamiento BGP y para ayudar a detectar problemas.

neighbor <i>direccion-ip next-hop-self</i>	Define al vecino como siguiente salto del vecino
neighbor <i>direccion-ip remote-as número-sistema-autónomo</i>	Establece una relación de vecino BGP
no synchronization	Deshabilita la publicación exclusivamente de redes con las cuales está en condiciones de comunicarse utilizando una ruta aprendida por un protocolo de enrutamiento interior
Comandos Generales	
clock rate <i>clock-rate</i>	Establece la velocidad de reloj para un equipo de comunicaciones de datos (DCE)
ip address <i>dirección-ip mascara de subred</i>	Asigna una dirección IP a una interfaz
interface <i>tipo número</i>	Cambios en el modo de configuración global al modo de configuración de interfaz
network <i>dirección -red</i>	Activa el protocolo de enrutamiento especificado en la red especificada
no auto-summary	No restaura la conducta por default de sumarización automática de rutas de subredes en rutas a nivel de red.
no shutdown	Habilita una interfaz
redistribute <i>protocolo [proceso-id] [metric {valor-metrica transparent}] [subnets]</i>	Configura la redistribución en el protocolo especificado
redistribute static	Permite la publicación de una ruta estática cuando una interfaz no está definida
redistribute connected	Permite anunciar por el tipo de protocolo las interfaces directamente conectados que forman parte de la red del tipo de protocolo
show ip route	Muestra la tabla de enrutamiento IP
show running-config	Muestra el archivo de configuración activo

DATOS DE LOS DISPOSITIVOS

Routers					
Nombre	Direcciones IP				Protocolo y S.A
	Se 0/0	Se0/1	Se0/2	Se0/3	
R1	192.168.12.1/24	192.168.13.1/24	----	----	RIP (S.A 10)
R2	192.168.12.2/24	192.168.24.2/24	----	----	RIP y BGP (S.A 10)
R3	192.168.13.2/24	192.168.34.2/24	----	----	RIP y BGP (S.A 10)
R4	192.168.24.1/24	192.168.34.1/24	192.168.46.1/24	192.168.45.1/24	OSPF y BGP (S.A 30)
R5	192.168.47.1/24	192.168.48.1/24	192.168.45.2/24	----	OSPF y BGP (S.A 30)
R6	192.168.46.2/24	192.168.47.2/24	192.168.14.1/24	----	EIGRP y BGP (S.A 20)
R7	192.168.14.2/24	----	----	----	EIGRP (S.A 20)
R8	192.168.48.2/24	192.168.15.1/24	192.168.16.1/24	----	ISIS Y BGP (S.A 40)
R9	192.168.15.2/24	192.168.17.2/24	----	----	ISIS (S.A 40)
R10	192.168.17.1/24	192.168.16.2/24	----	----	ISIS (S.A 40)

Interfaz FastEthernet				
Nombre	Direcciones IP	Puerto local	Puerto remoto	Dirección de Host remoto
R1	192.168.0.1/24			
R4	192.168.1.1/24			
R7	192.168.2.1/24			
R10	192.168.3.1/24			
PC1	192.168.0.2/24	30000	20000	127.0.0.1
PC2	192.168.1.2/24	30001	20001	127.0.0.1
PC3	192.168.2.2/24	30002	20002	127.0.0.1
PC4	192.168.3.2/24	30003	20003	127.0.0.1

ENUNCIADO**Asignación de direcciones IP**

Se deberá asignar las direcciones IP a las interfaces FastEthernet de las PC y Routers tal y como aparecen en el cuadro llamado Interfaz FastEthernet.

Nota: Se recomienda ejecutar primero El Software Virtual PC y Luego GNS3, 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 cuadro llamado Interfaz FastEthernet para configuración de puerto local y puerto remoto de las PCs en GNS3

Ejemplo:

Si la maquina PC1 en Virtual PC tiene el puerto local 20001 y puerto remoto 30001 nosotros configuraremos una de las PC en GNS3 con puerto local 30001 y puerto remoto 20001 y solo esa máquina podrá usar ese puerto local y remoto y así de las misma manera haremos para las demás PCs en GNS3 fijándonos que puerto local y remoto tienen las PCs en Virtual PC. de esta manera conectamos Virtual PC con GNS3

Cabe mencionar que las direcciones de las PC se configuran en la consola de Virtual PC no en GNS3 a diferencia de los Routers.

- R1.
- R4.
- R7.
- R10.
- VPCS 1.
- VPCS 2.
- VPCS 3.
- VPCS 4.

De igual manera a los Routers se le estarán asignando las direcciones IP a las interfaces Seriales a como se muestra en el cuadro llamado Routers.

- R1.
- R2.
- R3.
- R4.
- R5.
- R6.
- R7.
- R8.
- R9.
- R10.

Configuración de protocolo de enrutamiento (RIP)

Para enrutar el tráfico entre cada una de las redes existente del Sistema Autónomo 10, se estará usando RIP como protocolo de enrutamiento, en cada uno de los Routers (Router 1, Router 2, Router 3) pertenecientes a este Sistema Autónomo ver cuadro de Routers para ver cuáles son los datos para los Routers pertenecientes a este Sistema Autónomo.

- R1.
- R2.
- R3.

Configuración de protocolo de enrutamiento (OSPF)

El protocolo público conocido como "Primero la ruta más corta" (OSPF) es un protocolo de enrutamiento de estado del enlace no patentado.

Se deberá de configurar este protocolo para enrutar el tráfico entre cada una de las redes existente del Sistema Autónomo 30 ver cuadro de Routers para ver cuáles son los datos para los Routers pertenecientes a este Sistema Autónomo y configurar el Protocolo antes mencionado.

- R4.
- R5.

Configuración de protocolo de enrutamiento (EIGRP)

EIGRP es un protocolo de enrutamiento propietario de Cisco basado en IGRP.

Con frecuencia, se describe EIGRP como un protocolo de enrutamiento híbrido que ofrece lo mejor de los algoritmos de vector-distancia y del estado de enlace.

Se deberá de configurar este protocolo para enrutar el tráfico entre cada una de las redes existente del Sistema Autónomo 20 ver cuadro de Routers para ver cuáles son los Routers pertenecientes a este Sistema Autónomo y configurar el Protocolo antes mencionado.

- R6.
- R7

Configuración de protocolo de enrutamiento (IS-IS)

IS-IS es un protocolo de IGP desarrollado por DEC, Y suscrito por la ISO en los años 80 como protocolo de Routing para OSI, IS-IS Y OSPF son similares Ambos son protocolos de Routing de estado de enlace. Ambos están basados en el algoritmo SPF basado a su vez en el algoritmo de Dijkstra.

Se configurara el protocolo ISIS en el Sistema Autónomo 40 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.

- R8.
- R9.
- R10.

Configuración de Redistribución de rutas en los sistemas autónomos

La redistribución de rutas para que dos dispositivos (Routers o Switches capa 3) intercambien información de enrutamiento es preciso, en principio, que ambos dispositivos 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.

Sin embargo, cuando un dispositivo aprende información de enrutamiento a partir de diferentes fuentes (por ejemplo. rutas estáticas o a través de diferentes protocolos) Cisco IOS permite que la información aprendida por una fuente sea publicada hacia otros dispositivos utilizando un protocolo diferente.

Por ejemplo, que una ruta aprendida a través de RIP sea publicada hacia otros dispositivos utilizando OSPF. 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 (otro protocolo, rutas estáticas o directamente conectadas).

1. Redistribución de rutas en el sistema autónomo 10

En el sistema autónomo 10 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.

Nota: recuerde que los Routers de bordes de cada sistema autónomo son los únicos que se establecerán dos tipos de protocolo de enrutamiento IP (BGP para comunicación con otros sistemas autónomos y el protocolo de enrutamiento interno del sistema autónomo).

Ver cuadro de Routers y topología para observar cuales son los Routers de bordes pertenecientes a este sistema autónomo (los Routers que tienen dos protocolos de enrutamiento) y ver cuál es nuestro Router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, luego configurar la redistribución de rutas de BGP en el protocolo RIP protocolo interno del sistema autónomo y de RIP en BGP.

- R2.
- R3.

2. Redistribución de rutas en el sistema autónomo 30

En el sistema autónomo 30 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.

Ver cuadro de 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 BGP y se deberá configurar como tal. tendrá igual que ver cuáles son nuestros Routers vecinos conectados directamente a nuestros dos Routers de los otros sistemas autónomos conectado directamente al sistema autónomo 30, luego configurar la redistribución de rutas de BGP en el protocolo OSPF protocolo interno del sistema autónomo y de OSPF en BGP.

- R4.
- R5.

3. Redistribución de rutas en el sistema autónomo 20

En el sistema autónomo 20 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

Ver cuadro de Routers y topología para observar cuales son los Routers de bordes pertenecientes a este sistema autónomo (los Routers que tienen dos protocolos de enrutamiento) y ver cuál es nuestro Router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, luego configurar la redistribución de rutas de BGP en el protocolo EIGRP protocolo interno del sistema autónomo y de EIGRP en BGP.

- R6.

4. Redistribución de rutas en el sistema autónomo 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

Ver cuadro de Routers y topología para observar cuales son los Routers de bordes pertenecientes a este sistema autónomo (los Routers que tienen dos protocolos de enrutamiento) y ver cuál es nuestro Router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, luego configurar la redistribución de rutas de BGP en el protocolo IS-IS protocolo interno del sistema autónomo y de IS-IS en BGP.

- R8.

Tiempo estimado de solución

- 6 horas

Preguntas de análisis

La siguiente línea de código está presente en la tabla de enrutamiento:

```
O 10.16.1.0/27 [110/129] vía 192.168.1.5, 00:00:05, Serial0/0/1
```

1. ¿Qué indica el número 129 en este resultado?
 - a) El costo de este enlace tiene un valor de 129.
 - b) La frecuencia de reloj en esta interfaz serial está establecida en 129,000.
 - c) El Router de siguiente salto está a 129 saltos de distancia de este Router.
 - d) Esta ruta ha sido actualizada 129 veces en esta tabla de enrutamiento.

2. ¿Cuáles son las dos afirmaciones que describen EIGRP? (Elija dos opciones).
 - a) EIGRP se puede utilizar con Routers Cisco y con Routers que no son Cisco.
 - b) EIGRP envía updates disparados cada vez que hay un cambio en la topología que influye en la información de enrutamiento.
 - c) EIGRP tiene una métrica infinita de 16.
 - d) EIGRP envía una actualización parcial de la tabla de enrutamiento, que incluye solamente rutas que se han cambiado.
 - e) EIGRP transmite sus actualizaciones a todos los Routers en la red.

3. ¿Cuáles son las afirmaciones verdaderas con respecto a las métricas? (Elija dos opciones).
 - a) RIP utiliza el ancho de banda como métrica.
 - b) OSPF utiliza el retardo como métrica.
 - c) EIGRP utiliza el ancho de banda como métrica.
 - d) OSPF utiliza el costo basado en el ancho de banda como métrica.
 - e) RIP utiliza el retardo como métrica.
 - f) EIGRP utiliza el conteo de saltos solamente como métrica.

4. ¿Cuál de las siguientes opciones describe mejor la operación de los protocolos de enrutamiento por vector de distancia?
 - a) La única métrica que utilizan es el conteo de saltos.
 - b) Sólo envían actualizaciones cuando se agrega una nueva red.
 - c) Envían sus tablas de enrutamiento a los vecinos directamente conectados.
 - d) Inundan toda la red con actualizaciones de enrutamiento.

5. ¿Cuál es el propósito de un protocolo de enrutamiento?
 - a) Se utiliza para desarrollar y mantener tablas ARP.
 - b) Proporciona un método para segmentar y re ensamblar los paquetes de datos.

- c) Permite que un administrador cree un esquema de direccionamiento para la red.
 - d) Permite que un Router comparta información acerca de redes conocidas con otros Routers.
 - e) Ofrece un procedimiento para codificar y decodificar datos en bits para el reenvío de paquetes.
6. Un administrador de red está evaluando RIP versus EIGRP para una red nueva. La red será sensible a la congestión y debe responder rápidamente a cambios de topología. ¿Cuáles son dos buenas razones para elegir EIGRP en lugar de RIP en este caso? (Elija dos opciones).
- a) EIGRP utiliza actualizaciones periódicas.
 - b) EIGRP sólo actualiza vecinos afectados.
 - c) EIGRP utiliza actualizaciones de broadcast.
 - d) Las actualizaciones de EIGRP son parciales.
 - e) EIGRP utiliza el eficiente algoritmo Bellman-Ford.

PRÁCTICA N° 12: PROTOCOLOS DE ENRUTAMIENTO CON DIRECCIONAMIENTO IPv6

OBJETIVOS

General

- Configurar protocolos de enrutamiento dinámico, entender su interrelación con otros protocolos y poner en funcionamiento el protocolo IPV6.

Específicos

- Escribir la configuración de protocolos de enrutamiento internos dentro de diferentes sistemas autónomos implementando el protocolo IPV6.
- Establecer y configurar el protocolo BGP en los enrutadores de borde de cada uno de los sistemas autónomos.
- Estudiar e implementar la redistribución de rutas en los enrutadores de borde, para poder comunicar los diferentes sistemas autónomos.

INTRODUCCIÓN

Con la implementación y realización de esta práctica, se pretende que los estudiantes sean capaces de entender y poder configurar los diferentes tipos de protocolos de enrutamiento dinámico internos (RIP, IS-IS, OSPF), implementando el protocolo IPV6 y a su vez comunicar los diferentes sistemas autónomos a través del protocolo de enrutamiento externo BGP. 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: BGP y el protocolo de enrutamiento interno que corresponda con el sistema autónomo.

REQUERIMIENTOS

- **Hardware**
 - Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.
- **Software**
 - Emulador de redes GNS3 0.8.4
 - 10 Routers de la serie 3640
 - 4 PCs.

Nota: Se deberá de descargar el software Virtual PC para hacer uso de los 4 PC

CONOCIMIENTOS PREVIOS

Para la correcta realización 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 (BGP) de Acceso proporcionado previamente en los aspectos teóricos.

TOPOLOGÍA

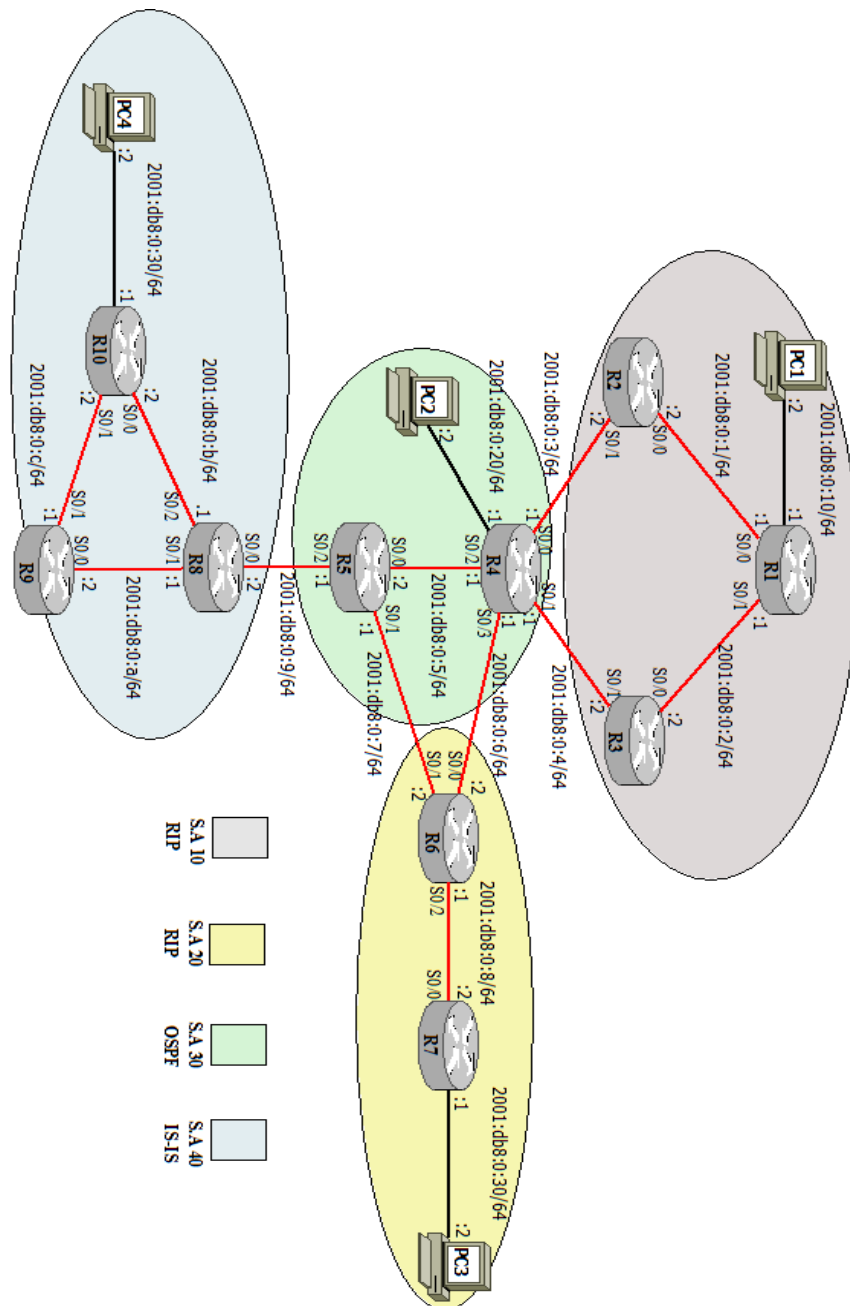


Figura 111. Topología de práctica de protocolos de enrutamiento con direccionamiento IPv6

FUNCIONALIDAD

La [Figura 111](#) muestra la topología en la cual se estarán aplicando diferentes tipos de protocolos de enrutamiento dinámico fusionándolo con BGP

- Se deberá establecer un protocolo de enrutamiento dinámico para cada Sistema Autónomo como lo muestra la topología implementando el protocolo IPV6.

- En los Routers de Borde de cada Sistema Autónomo se configurará BGP para la comunicación entre diferentes Sistemas Autónomos implementando el protocolo IPV6.
- En los Routers de Borde de cada Sistema Autónomo se redistribuirá rutas entre BGP y el protocolo de enrutamiento que se estableció en ese Sistema Autónomo.

COMANDOS DE AYUDA

Comando	Descripción
Comandos de RIP	
ipv6 rip -proceso-id enable	Habilita enrutamiento RIP con IPV6 en una interfaz
ipv6 router rip - proceso-id	Habilita enrutamiento RIP con IPV6 en el Router
Comandos de OSPF	
ipv6 ospf 30 area 1	Habilita enrutamiento OSPF con IPV6 en una interfaz
ipv6 router ospf 30	Habilita enrutamiento OSPF con IPV6 en el Router
log-adjacency-changes	Traza los cambios en estado de adyacencia
router-id	Asigna un id al Router OSPF
Comandos de IS-IS	
router isis	Entra en el modo de configuración del Router para IS-IS
address-family ipv6	Permite configurar ISIS con las familias de direcciones IPV6
ipv6 router isis	Permite IS-IS con IPV6 en una interfaz
net direccion-net	Configura Título Entidad de Red (NET) para IS-IS
Comandos de BGP	
router bgp número-sistema-autónomo	Entra en el modo de configuración del Router para BGP
address-family ipv6	Permite configurar BGP con las familias de direcciones IPV6
bgp log-neighbor-changes	Habilita el registro de los cambios de adyacencia de vecinos para monitorear la estabilidad del sistema de enrutamiento BGP y para ayudar a detectar problemas.
bgp router-id 2.2.2.2	Asigna un id al Router bgp
neighbor direccion-ip next-hop-self	Define al vecino como siguiente salto del vecino
neighbor direccion-ip remote-as número-sistema-autónomo	Establece una relación de vecino BGP

no synchronization	Deshabilita la publicación exclusivamente de redes con las cuales está en condiciones de comunicarse utilizando una ruta aprendida por un protocolo de enrutamiento interior
Comandos Generales	
clock rate <i>clock-rate</i>	Establece la velocidad de reloj para un equipo de comunicaciones de datos (DCE)
ip address <i>dirección-ip máscara de subred</i>	Asigna una dirección IP a una interfaz
interface <i>tipo número</i>	Cambios en el modo de configuración global al modo de configuración de interfaz
ipv6 unicast-routing	Permite enrutamiento IPV6
no auto-summary	No restaura la conducta por default de sumarización automática de rutas de subredes en rutas a nivel de red.
no shutdown	Habilita una interfaz
redistribute <i>protocolo [proceso-id] [metric {valor-metrica transparent}]</i>	Configura la redistribución en el protocolo especificado
redistribute static	Permite la publicación de una ruta estática cuando una interfaz no está definida
redistribute connected	Permite anunciar por el tipo de protocolo las interfaces directamente conectados que forman parte de la red del tipo de protocolo
show ip route	Muestra la tabla de enrutamiento IP
show running-config	Muestra el archivo de configuración activo

DATOS DE LOS DISPOSITIVOS

Routers					
Nombre	Direcciones IP				
	Se 0/0	Se0/1	Se0/2	Se0/3	Protocolo y S.A
R1	2001:db8:0:1::1/64	2001:db8:0:2::1/64	----	----	RIP (S.A 10)
R2	2001:db8:0:1::2/64	2001:db8:0:3::2/64	----	----	RIP y BGP (S.A 10)
R3	2001:db8:0:2::2/64	2001:db8:0:4::2/64	----	----	RIP y BGP (S.A 10)

R4	2001:db8:0:3::1/ 64	2001:db8:0:4::1/ 64	2001:db8:0:5::1/ 64	2001:db8:0:6::1/ 64	OSPF y BGP (S.A 30)
R5	2001:db8:0:5::1/ 64	2001:db8:0:7::1/ 64	2001:db8:0:9::1/ 64	-----	OSPF y BGP (S.A 30)
R6	2001:db8:0:6::2/ 64	2001:db8:0:7::2/ 64	2001:db8:0:8::1/ 64	-----	RIP y BGP (S.A 20)
R7	2001:db8:0:8::2/ 64	-----	-----	-----	RIP (S.A 20)
R8	2001:db8:0:9::2/ 64	2001:db8:0:a::1/ 64	2001:db8:0:b::1/ 64	-----	ISIS Y BGP (S.A 40)
R9	2001:db8:0:a::2/ 64	2001:db8:0:c::1/ 64	-----	-----	ISIS (S.A 40)
R10	2001:db8:0:b::2/ 64	2001:db8:0:c::2/ 64	-----	-----	ISIS (S.A 40)

Interfaz FastEthernet				
Nombre	Direcciones IP	Puerto local	Puerto remoto	Dirección de Host remoto
R1	2001:DB8:0:10::1/64			
R4	2001:DB8:0:20::1/64			
R7	2001:DB8:0:30::1/64			
R10	2001:DB8:0:40::1/64			
PC1	2001:DB8:0:10::2/64	30004	20004	127.0.0.1
PC2	2001:DB8:0:20::2/64	30005	20005	127.0.0.1
PC3	2001:DB8:0:30::2/64	30006	20006	127.0.0.1
PC4	2001:DB8:0:40::2/64	30007	20007	127.0.0.1

ENUNCIADO**Asignación de direcciones IP**

Se deberá asignar las direcciones de las PCs y a las interfaces FastEthernet de los Routers tal y como aparecen en el cuadro llamado Interfaz FastEthernet.

Nota: Se recomienda ejecutar primero El Software Virtual PC y luego GNS3, 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.

Ejemplo:

Si la maquina PC1 en Virtual PC tiene el puerto local 20001 y puerto remoto 30001 nosotros configuraremos una de las PC en GNS3 con puerto local 30001 y puerto remoto 20001 y solo esa máquina podrá usar ese puerto local y remoto y así de las misma manera haremos para las demás PCs en GNS3 fijándonos que puerto local y remoto tienen las PCs en Virtual PC. de esta manera conectamos Virtual PC con GNS3

Cabe mencionar que las direcciones de las PC se configuran en la consola de Virtual PC no en GNS3 a diferencia de los Routers.

- R1.
- R4.
- R7.
- R10.
- VPCS 1.
- VPCS 2.
- VPCS 3.
- VPCS 4.

De igual manera a los Routers se le estarán asignando las direcciones IP a las interfaces Seriales a como se muestra en el cuadro llamado Routers.

- R1.
- R2.
- R3.
- R4.
- R5.
- R6.
- R7.
- R8.
- R9.

- R10.

Configuración de protocolo de enrutamiento (RIP)

Para encaminar el tráfico entre cada una de las redes existente del Sistema Autónomo 10, se estará usando RIP como protocolo de enrutamiento, en cada uno de los Routers Pertenecientes a este Sistema Autónomo ver cuadro llamado Routers para ver cuáles son los Routers pertenecientes a este Sistema Autónomo (recuerde que se está utilizando el protocolo IPv6 ver cuadro de ayuda de comandos para una correcta configuración).

- R1.
- R2.
- R3.

Configuración de protocolo de enrutamiento (OSPF)

El protocolo público conocido como "Primero la ruta más corta" (OSPF) es un protocolo de enrutamiento de estado del enlace no patentado.

Se deberá de configurar este protocolo para encaminar el tráfico entre cada una de las redes existente del Sistema Autónomo 30 ver cuadro de Routers para ver cuáles son los Routers pertenecientes a este Sistema Autónomo y configurar el Protocolo antes mencionado (de nuevo recuerde que se está utilizando el protocolo IPv6 ver cuadro de ayuda de comandos para una correcta configuración).

- R4.
- R5.

Configuración de protocolo de enrutamiento (RIP)

Para encaminar el tráfico entre cada una de las redes existente del Sistema Autónomo 20, se estará usando RIP como protocolo de enrutamiento, en cada uno de los Routers Pertenecientes a este Sistema Autónomo ver cuadro de Routers para ver cuáles son los Routers pertenecientes a este Sistema Autónomo (recuerde que se está utilizando el protocolo IPv6 ver cuadro de ayuda de comandos para una correcta configuración).

- R6.
- R7.

Configuración de protocolo de enrutamiento (IS-IS)

IS-IS es un protocolo de IGP desarrollado por DEC, Y suscrito por la ISO en los años 80 como protocolo de Routing para OSI, IS-IS Y OSPF son similares Ambos son protocolos de Routing de estado de enlace. Ambos están basados en el algoritmo SPF basado a su vez en el algoritmo de Dijkstra.

Se configurara el protocolo ISIS en el Sistema Autónomo 40 para encaminar el tráfico perteneciente a ese Sistema Autónomo ver cuadro de Routers para ver cuáles son los Routers pertenecientes a este Sistema Autónomo y configurar este protocolo.

- R8.
- R9.
- R10.

Configuración de Redistribución de rutas en los sistemas autónomos

La redistribución de rutas para que dos dispositivos (Routers o Switches de capa 3) intercambien información de enrutamiento es preciso, en principio, que ambos dispositivos 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.

Sin embargo, cuando un dispositivo aprende información de enrutamiento a partir de diferentes fuentes (por ejemplo. rutas estáticas o a través de diferentes protocolos) Cisco IOS permite que la información aprendida por una fuente sea publicada hacia otros dispositivos utilizando un protocolo diferente.

Por ejemplo, que una ruta aprendida a través de RIP sea publicada hacia otros dispositivos utilizando OSPF. 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 (otro protocolo, rutas estáticas o directamente conectadas).

5. Redistribución de rutas en el sistema autónomo 10

En el sistema autónomo 10 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.

Nota: recuerde que los Routers de bordes de cada sistema autónomo son los únicos que se establecerán dos tipos de protocolo de enrutamiento IP (BGP para comunicación con otros sistemas autónomos y el protocolo de enrutamiento interno del sistema autónomo).

Ver cuadro llamado Routers y la topología para observar cuales son los Routers de bordes pertenecientes a este sistema autónomo (los Routers que tienen dos protocolos de enrutamiento) y ver cuál es nuestro Router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, luego configurar la redistribución de rutas de BGP en el protocolo RIP protocolo interno del sistema autónomo y de RIP en BGP (ver cuadro de ayuda de comandos para la correcta configuración).

- R2.
- R3.

6. Redistribución de rutas en el sistema autónomo 30

En el sistema autónomo 30 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.

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 BGP y se deberá configurar como tal. Tendrá igual que ver cuáles son nuestros Routers vecinos conectados directamente a nuestros dos Routers de los otros sistemas autónomos conectado directamente al sistema autónomo 30, luego configurar la redistribución de rutas de BGP en el protocolo OSPF protocolo interno del sistema autónomo y de OSPF en BGP(ver cuadro ayuda de comandos para la correcta configuración).

- R4.
- R5.

7. Redistribución de rutas en el sistema autónomo 20

En el sistema autónomo 20 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

Ver cuadro de Routers y la topología para observar cuales son los Routers de bordes pertenecientes a este sistema autónomo (los Routers que tienen dos protocolos de enrutamiento) y ver cuál es nuestro Router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, luego configurar la redistribución de rutas de BGP en el protocolo EIGRP protocolo interno del sistema autónomo y de EIGRP en BGP

- R6.

8. Redistribución de rutas en el sistema autónomo 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

Ver cuadro de Routers y topología para observar cuales son los Routers de bordes pertenecientes a este sistema autónomo (los Routers que tienen dos protocolos de enrutamiento) y ver cuál es nuestro Router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, luego configurar la redistribución de rutas de BGP en el protocolo IS-IS protocolo interno del sistema autónomo y de IS-IS en BGP.

- R8.

Tiempo estimado de solución

- 6 horas

Preguntas de análisis

La siguiente línea de código está presente en la tabla de enrutamiento:

```
O 10.16.1.0/27 [110/129] via 192.168.1.5, 00:00:05, Serial0/0/1
```

1. ¿Qué indica el número 129 en este resultado?
 - a) El costo de este enlace tiene un valor de 129.
 - b) La frecuencia de reloj en esta interfaz serial está establecida en 129,000.
 - c) El Router de siguiente salto está a 129 saltos de distancia de este Router.
 - d) Esta ruta ha sido actualizada 129 veces en esta tabla de enrutamiento.

2. ¿Cuáles son las dos afirmaciones que describen EIGRP? (Elija dos opciones).
 - a) EIGRP se puede utilizar con Routers Cisco y con Routers que no son Cisco.
 - b) EIGRP envía updates disparados cada vez que hay un cambio en la topología que influye en la información de enrutamiento.
 - c) EIGRP tiene una métrica infinita de 16.
 - d) EIGRP envía una actualización parcial de la tabla de enrutamiento, que incluye solamente rutas que se han cambiado.
 - e) EIGRP transmite sus actualizaciones a todos los Routers en la red.

3. ¿Cuáles son las afirmaciones verdaderas con respecto a las métricas? (Elija dos opciones).
 - a) RIP utiliza el ancho de banda como métrica.
 - b) OSPF utiliza el retardo como métrica.
 - c) EIGRP utiliza el ancho de banda como métrica.
 - d) OSPF utiliza el costo basado en el ancho de banda como métrica.
 - e) RIP utiliza el retardo como métrica.
 - f) EIGRP utiliza el conteo de saltos solamente como métrica.

4. ¿Cuál de las siguientes opciones describe mejor la operación de los protocolos de enrutamiento por vector de distancia?
 - a) La única métrica que utilizan es el conteo de saltos.
 - b) Sólo envían actualizaciones cuando se agrega una nueva red.
 - c) Envían sus tablas de enrutamiento a los vecinos directamente conectados.
 - d) Inundan toda la red con actualizaciones de enrutamiento.

5. ¿Cuál es el propósito de un protocolo de enrutamiento?
 - a) Se utiliza para desarrollar y mantener tablas ARP

- b) Proporciona un método para segmentar y re ensamblar los paquetes de datos.
 - c) Permite que un administrador cree un esquema de direccionamiento para la red.
 - d) Permite que un Router comparta información acerca de redes conocidas con otros Routers.
 - e) Ofrece un procedimiento para codificar y decodificar datos en bits para el reenvío de paquetes.
6. Un administrador de red está evaluando RIP versus EIGRP para una red nueva. La red será sensible a la congestión y debe responder rápidamente a cambios de topología. ¿Cuáles son dos buenas razones para elegir EIGRP en lugar de RIP en este caso? (Elija dos opciones).
- a) EIGRP utiliza actualizaciones periódicas.
 - b) EIGRP sólo actualiza vecinos afectados.
 - c) EIGRP utiliza actualizaciones de broadcast.
 - d) Las actualizaciones de EIGRP son parciales.
 - e) EIGRP utiliza el eficiente algoritmo Bellman-Ford

PRÁCTICA N° 13: **DISEÑO Y ANÁLISIS DE REDES** **WLAN**

OBJETIVOS

General

- Integrar y configurar una solución inalámbrica en el marco de una organización empresarial.

Específicos

- Diseñar una topología de red con los componentes necesarios que participan de una red inalámbrica.
- Configurar una red inalámbrica por medio de sus puntos de acceso, estableciendo los SSID (Service Set Identifier) y el rango de direcciones dinámicas.
- Implementar VLAN en los grupos de trabajo, conformados por medios físicos e inalámbricos, de acuerdo con un servicio DHCP (Dynamic Host Configuration Protocol).

INTRODUCCIÓN

La implementación y realización de esta práctica, está orientado al desarrollo de una plataforma de una red empresarial, con base en el diseño e implementación de un sistema de comunicaciones basadas en Inter-VLANs, y algunos componentes necesarios que participan de una red inalámbrica.

Se pretende que los estudiantes sean capaces de dar solución a problemas de conectividad en el marco de una red inalámbrica y física empresarial, con acceso protegidos a algunos servicios implementando grupos de trabajos conformados por medios físicos e inalámbricos.

REQUERIMIENTOS

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.

➤ Software

- Simulador Cisco Packet Tracer 6.0 con los siguientes elementos:
 - 4 Servidores
 - 2 Routers serie 1841
 - 2 Routers Wireless
 - 1 Switches 2950
 - 1 Hub
 - 4 PCs
 - 2 Laptops
 - Cable directo
 - Cable cruzado

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica es necesario que el estudiante tenga conocimientos básicos de cómo manejar conectividad de manera inalámbrica y por medios físicos en el marco de una red empresarial; además conocimientos básicos de los siguientes temas:

- Creación de VLANs.
- Asignación dinámica de direcciones IP mediante DHCP.
- Hacer uso de Sub-Interfaces.
- Agregación de rutas estáticas.

TOPOLOGÍA

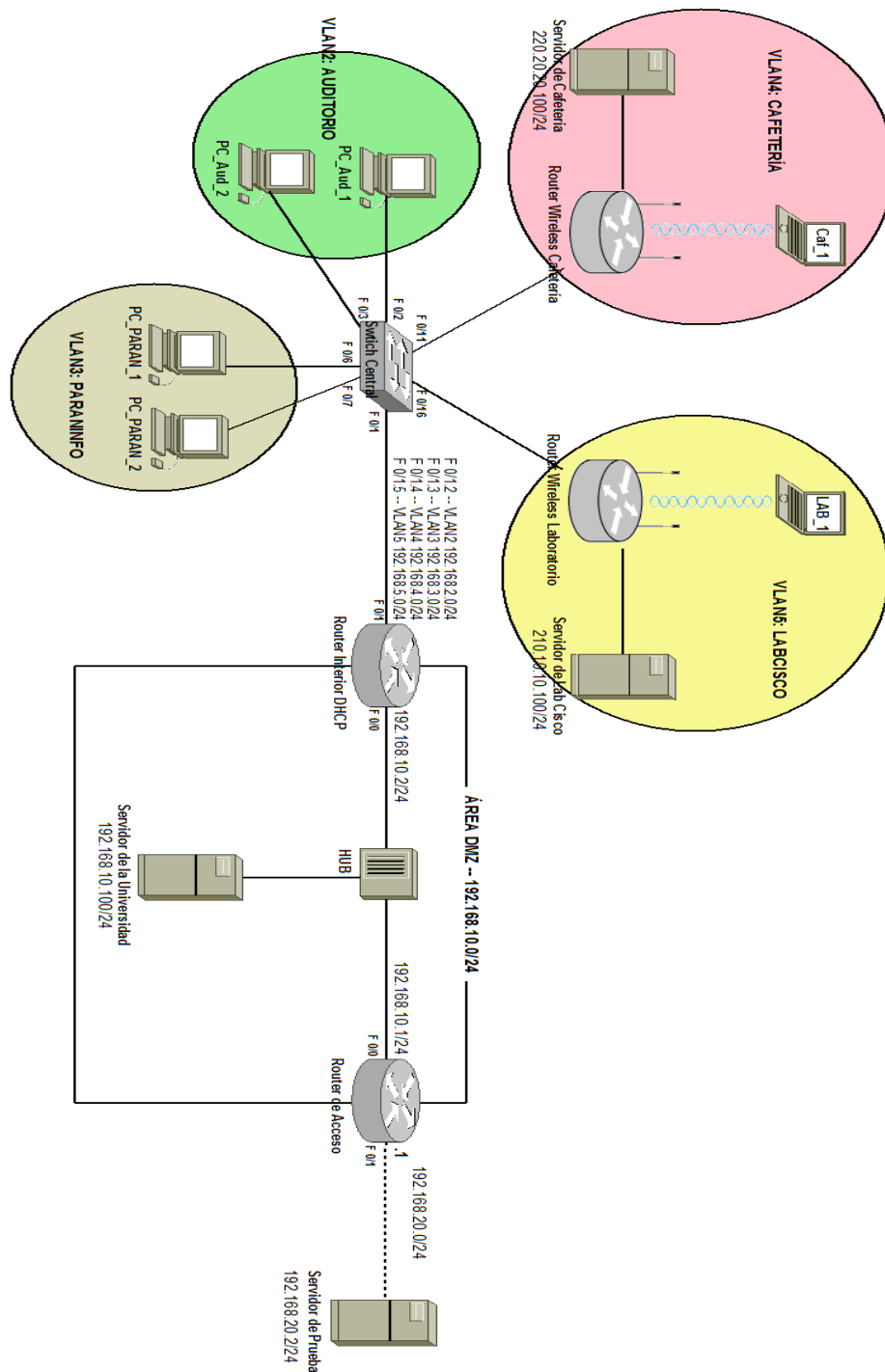


Figura 112. Topología de práctica de Diseño y análisis de una red WLAN

FUNCIONAMIENTO

La Figura 112 muestra la topología de una red empresarial, donde lo que se pretende es configurar e integrar una solución inalámbrica demostrando la escalabilidad que puede existir a la hora de interconectar equipos de manera inalámbrica así como también por medios físicos (Cables):

- Se deberán crear cuatro VLANs, asignándoles direcciones IP dinámicamente a los equipos finales perteneciente a cada una de ellas mediante el protocolo DHCP.
- Los equipos de las cuatro VLANs creadas podrán tener acceso al Servidor de la Universidad como al Servidor de Prueba.
- Las PC que están conectadas de manera inalámbricas necesitaran alguna manera de autenticación para poderse conectar.

DATOS DE LOS DISPOSITIVOS

VLANs		
Número	Nombre	Dirección IP
2	AUDITORIO	192.168.2.1/24
3	PARANINFO	192.168.3.1/24
4	CAFETERÍA	192.168.4.1/24
5	LAB_CISCO	192.168.5.1/24

Routers		
Nombre	Dirección IP	
	FastEthernet 0/0	FastEthernet 0/1
Router Interior DHCP	192.168.10.2/24	-----
Router de Acceso	192.168.10.1/24	192.168.20.1/24

Router Interior DHCP	
Sub-Interfaz	Dirección IP
FastEthernet 0/1.2	192.168.2.1/24
FastEthernet 0/1.3	192.168.3.1/24
FastEthernet 0/1.4	192.168.4.1/24
FastEthernet 0/1.5	192.168.5.1/24

Switch Central		
Interfaces	Modo	VLAN
FastEthernet 0/1	Troncal	--
FastEthernet 0/2-5	Acceso	2
FastEthernet 0/6-10	Acceso	3
FastEthernet 0/11-15	Acceso	4
FastEthernet 0/16-20	Acceso	5

Servidores		
Nombre	Dirección IP	Gateway
Servidor de Cafetería	220.20.20.100/24	220.20.20.1
Servidor Lab_Cisco	210.10.10.100/24	210.10.10.1
Servidor de la Universidad	192.168.10.100/24	192.168.10.1
Servidor de Prueba	192.168.20.2/24	192.168.20.1

ENUNCIADO

Crear la topología en el simulador Cisco PacketTracer 6.0 mostrada en la figura que representa la topología de esta, recordando usar los equipos y la serie de cada uno ellos mencionado en los requerimientos Hardware, así como también el tipo de cable que se usaran para la conexión entre estos.

Asignación de dirección IP

Asignar dirección IP a las interfaces de los siguientes equipos:

- Router Interior DHCP.
- Router de Acceso.
- Asignar direcciones IP estáticamente a los servidores a como se muestra en el cuadro llamado Servidores

Creación de VLANs

En el Switch Central crear las VLANs a como se detallan en cuadro nombrado VLANs.

En el Switch Central configurar las interfaces pertenecientes a cada una de las VLANs tal y como aparece en el cuadro nombrado Switch Central.

Asignar direcciones IP a cada una de las sub-interfaces en el Router Interior DHCP para que cada una de las VLANs pueda tomar sus respectivas direcciones IP.

Servicio DHCP

Indicar a las PC que deberán tomar sus respectivas direcciones IP mediante DHCP.

Agregación de rutas estáticas o rutas por defectos

Para que en la topología de red los equipos que se encuentran dentro de la red interna (LAN), puedan comunicarse con la parte externa (WAN) es decir con el Servidor de Prueba; será necesario anunciar las rutas a cada uno de los enrutadores, por lo que se tendrá que agregar rutas estáticas o rutas por defectos.

Agregar rutas estáticas en el Router Interior DHCP y el Router de Acceso para poder establecer comunicación de la LAN con la WAN.

- Router Interior DHCP.
- Router de Acceso.

Configuración de la Red Inalámbrica

Verificar que los Router Linksys hayan recibido correctamente su dirección IP mediante DHCP de acuerdo a la VLAN perteneciente a cada uno.

Habilitar el servicio DHCP en los Router Linksys de la siguiente manera:

- a) La dirección IP inicial que estará asignando el Router Wireless Academia Cisco a los clientes que se conecten a él será a partir de 210.10.10.2 con máscara 255.255.255.0 y tendrá un máximo de 50 usuarios.
- b) La dirección IP inicial que estará asignando el Router Wireless Cafetería a los clientes que se conecten a él será a partir de 220.20.20.2 con máscara 255.255.255.0 y tendrá un máximo de 50 usuarios.

Configuración de los SSID

- a) En el Router Wireless Cafetería configurar el SSID como CAFETERIA.
- b) En el Router Wireless Academia Cisco configurar el SSID como LABCISCO.

Configuración Segura

- a) Para los usuarios que quieran acceder a la red LABCISCO (DHCP: 192.168.5.0/24), se requiere la contraseña EST_TELEMATICA, de acuerdo con WPA (Wi-Fi Protected Access) personal o de clave pre-compartida, configurar el Router Wireless Academia Cisco para que sea posible, así como también la PC cliente que tendrá que acceder.
 - En el Router Wireless Academia Cisco.
 - En la PC cliente.

Configuración Remota

- a) Para el Router Wireless Cafetería habilitar configuración remota del router, mediante HTTP a la dirección IP 192.168.4.2
 - Use name: admin
 - Password: telematica_123
- b) En el Router Wireless Cafetería.
- c) En la PC Cliente.

Tiempo estimado de solución

- 2 horas

Preguntas de Análisis

- 1. ¿Qué ventajas posee el administrado de una red al poder realizar configuraciones remotas de manera inalámbrica?
- 2. ¿Cuál es la importancia de configurar un SSID?
- 3. Según la figura que muestra la topología a la que se le dio solución ¿Qué configuraciones tendría que hacer el administrador de la red para tener lo más posible?

PRÁCTICA N° 14: **POLITICAS DE SEGURIDAD EN REDES DE COMUNICACIONES**

OBJETIVOS

General

- Integrar una solución de seguridad sobre la base de una intranet, habitual en muchas organizaciones empresariales.

Específicos

- Comprender las necesidades de seguridad que demandan las comunicaciones actuales, en el contexto del sector empresarial.
- Analizar y configurar las políticas de acceso a los recursos de una red de datos, desde la perspectiva de las ACL (Access Control List), tanto estándar como extendidas.
- Analizar y configurar mecanismos de acceso remoto seguro, como SSH (Secure SHell), respecto al método de acceso tradicional, Telnet.
- Analizar e implementar la seguridad AAA (Authentication, Authorization and Accounting), de acuerdo con la configuración de servidores RADIUS
- Analizar y configurar una VPN (Virtual Private Network), para acceder a una red privada desde una infraestructura pública.

INTRODUCCIÓN

La implementación y realización de esta práctica, está orientado al desarrollo de una plataforma de una red empresarial, con base en el diseño e implementación de un sistema de comunicaciones basadas en Inter-VLANs, y algunos protocolos de seguridad.

Se pretende que los estudiantes sean capaces de dar solución a problemas de seguridad de una red LAN ante red WAN y puedan realizar el ejercicio detallado a continuación.

REQUERIMIENTOS

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.

➤ Software

- Simulado Cisco Packet Tracer 6.0 con los siguientes elementos:
 - 3 Servidores.
 - 6 Routers serie 1841.
 - 1 Routers Wireless.
 - 3 Switches 2950.
 - 1 Hub.
 - 4 PCs.
 - 3 Laptops.

- Cable directo
- Cable serial DTE.

CONOCIMIENTOS PREVIOS

Para la correcta realización de esta práctica es necesario que el estudiante tenga conocimientos básicos de cómo manejar la seguridad interna de la red de una empresa, de protocolos y servicios que puedan usar, además conocimientos básicos de los siguientes temas:

- Creación de VLANs.
- Asignación dinámica de direcciones IP mediante DHCP.
- Implementación de protocolos de enrutamiento (RIP).
- Hacer uso de Sub-Interfaces.

TOPOLOGÍA

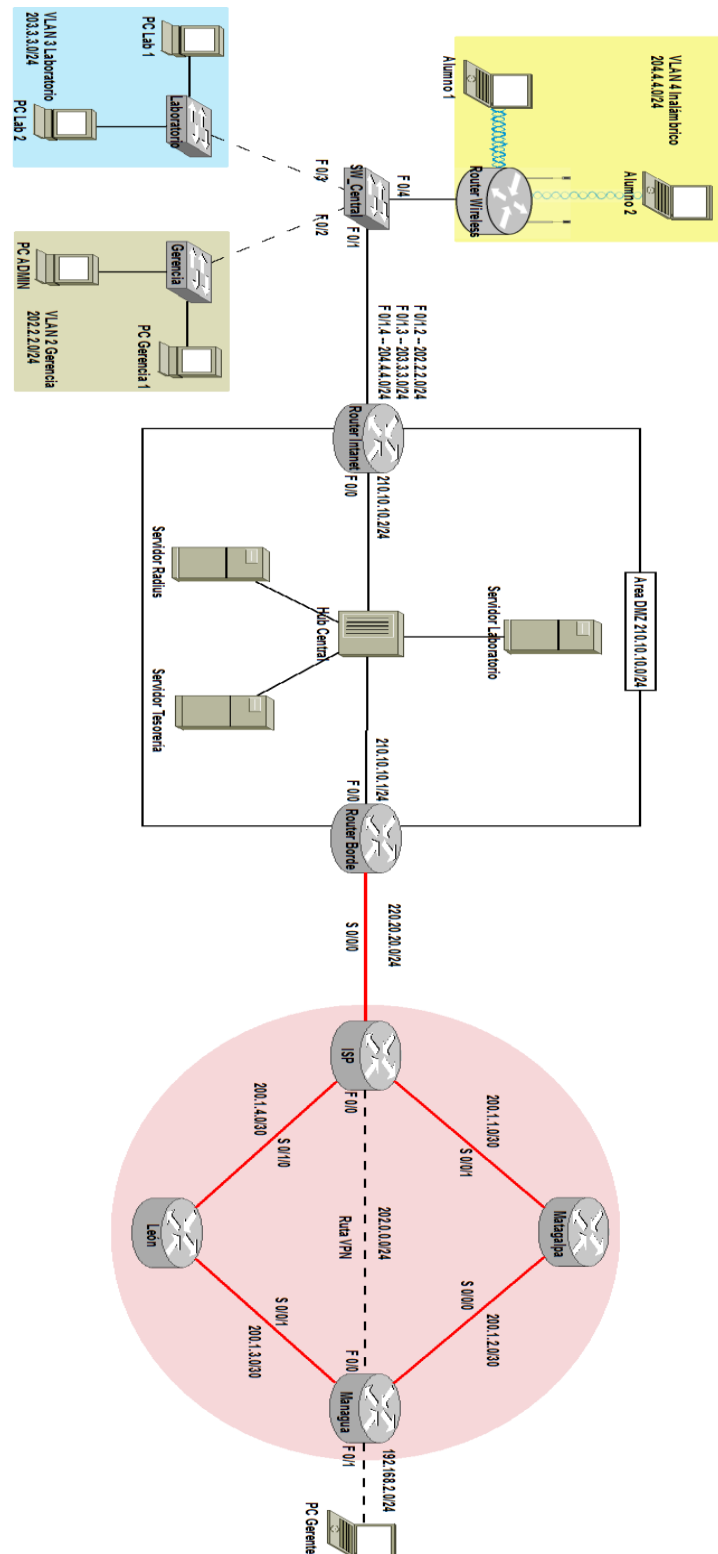


Figura 113. Topología de práctica de Políticas de seguridad en redes de comunicaciones

FUNCIONALIDAD

La Figura 113 muestra la topología de la plataforma de una red empresarial, en la que se pretende que por medio de algunos protocolos y servicios (SSH, AAA ACL, VPN) se logre tener una red con un uso seguro.

Se deberá configurarse protocolo de enrutamiento dinámico (RIP), para establecer comunicación entre los router exteriores a la red, además que las conexiones inalámbricas se realicen de manera segura mediante un servidor Radius.

Se deberá configurar un router como DHCP para asignarles direcciones IP de manera dinámica a cada una de las PC pertenecientes a la LAN, no obstante los servidores deberán recibir su dirección IP de manera estática.

COMANDOS DE AYUDA

Comando	Descripción
authentication pre-share	Autenticación basada en clave compartida
crypto ipsec security-association lifetime seconds <i>86400</i>	Tiempo en segundo de validez de la clave compartida
crypto ipsec transform-set <i>Conjunto_dominio</i> esp-aes esp-sha-hmac	Seleccionar para el conjunto dominio el protocolo ESP
crypto isakmp enable	Habilitar ISAKMP o Internet Security
crypto isakmp key <i>Clave</i> address <i>Dirección_IP</i>	Comparte la clave con el par especificado
crypto isakmp policy 1	Habilitar política de seguridad
crypto key generate rsa	Generar clave de cifrado
enable secret <i>Clave</i>	Habilitar clave modo privilegiado
encryption aes	Habilitar cifrado simétrico AES
hash sha	Habilita el algoritmo hash SHA
ip domain-name <i>Nombre_de_dominio</i>	Agregar nombre de dominio
ip ssh time-out <i>Tiempo_en_segundos</i>	tiempo de permanencia en la conexión segura
set peer <i>Dirección_IP</i>	Definir el otro extremo IP del mapa cifrado
ssh -l <i>Nombre_Usuario</i> <i>Dirección_IP</i>	Acceso a equipo remoto
transport input/output ssh	Tipo de transporte SSH
username <i>Nombre_usuario</i> password <i>Contraseña</i>	Nombre de usuario SSH y contraseña

DATOS DE LOS DISPOSITIVOS

Routers					
Nombre	Interfaces FastEthernet (/24)		Interfaces Serial (/30)		
	F 0/0	F 0/1	S 0/0/0	S 0/0/1	S 0/1/0
Router-Intranet	210.10.10.2	Sub Interfaces	-----	-----	-----
Router-Borde	210.10.10.1	-----	220.20.20.2	-----	-----
ISP	-----	-----	220.20.20.1	200.1.1.1	200.1.4.1
Matagalpa	-----	-----	200.1.2.1	200.1.1.2	-----
Managua	-----	-----	200.1.2.2	200.1.3.1	-----
León	-----	-----	200.1.4.2	200.1.3.1	-----

Router-Intranet		
Sub Interfaz	Dirección IP	VLAN
Fa0/1.1	201.1.1.1/24	1
Fa0/1.2	202.2.2.1/24	2
Fa0/1.3	203.3.3.1/24	3

Servidores	
Nombre	Dirección IP
Servidor Web Laboratorio	210.10.10.200/24
Servidor Radius	210.10.10.100/24
Servidor Tesorería	210.10.10.150/24

VLANs		
Número	Nombre	Dirección IP
1	Default	201.1.1.1/24
2	Laboratorio	202.2.2.1/24
3	Tesorería	203.3.3.1/24

PC	
Nombre	Dirección IP
Laptop del Gerente	192.168.100.2/24

ENUNCIADO

Crear la topología en el simulador Cisco PacketTracer 6.0 mostrada en la Figura 113 recordando usar los equipos y la serie de cada uno ellos mencionado en los requerimientos Hardware, así como también el tipo de cable que se usaran para la conexión entre estos.

Asignación de Dirección IP

Asignar dirección IP (ver tabla correspondiente de cada uno) a las interfaces de los siguientes equipos:

- a) Servidor Web Laboratorio.
- b) Servidor Radius.
- c) Servidor de Tesorería.
- d) Router-Intranet.
- e) Router de Acceso.
- f) ISP.
- g) Managua.
- h) Matagalpa.
- i) León.
- j) Laptop del Gerente.

Creación de VLANs

En el Switch-Central crear las VLANs a como se detallan en la tabla nombrada VLANs.

Asignar direcciones IP a cada una de las sub-interfaces en el router C para que cada una de las VLANs pueda tomar sus respectivas direcciones IP de manera dinámica.

Configurar las interfaces pertenecientes que van desde el Switch-Central hacia los Switches de Laboratorio y Tesorería para cada uno de estos pertenezca a su VLAN perteneciente.

Asignación de direcciones IP mediante DHCP a las PC y al Router Acceso Inalámbrico

Configurar un pool en el Router-Intranet con los nombres pertenecientes a cada una de las VLANs, indicando el rango de direcciones IP la cual se estarán asignando de manera dinámica a cada uno de los equipos finales (PC y Router Wireless o de Acceso Inalámbrico) que realice una petición de DHCP.

Indicar a las PCs que deberán realizar una petición de DHCP, para que obtengan su dirección IP de manera dinámica.

Configuración de protocolo de enrutamiento (RIP)

Para encaminar el tráfico entre cada una de las redes que existen entre cada uno de los routers externos, se estará usando RIP como protocolo de enrutamiento, y debe ser configurado en los siguientes equipos:

- a) Router ISP.

- b) Router Matagalpa.
- c) Router Managua.
- d) Router León.

Una vez establecido RIP como protocolo de enrutamiento en los Routers, debe existir comunicación entre los equipos de las diferentes redes, verificar si dicha comunicación existe.

Configuración del Servidor Radius

Crear un cliente que será nuestro Router de Acceso Inalámbrico que utilizara el servidor Radius para la autenticación y los usuarios serán los hosts finales que necesitarán de un usuario y contraseña (Laptop Alumno 1 y Laptop Alumno 2).

Nota:

- a) *Para el cliente a crear debe agregar la dirección IP del Router de Acceso Inalámbrico, establecer el número de puerto Radius y especificar que el tipo de servicio que será Radius.*
- b) *Para los usuarios deberá crear uno para cada equipo que se deberá conectar.*

Configurar al cliente que usara el servicio del Servidor Radius (Router de Acceso Inalámbrico)

Establecer el método de cifrado que utilizará y luego indicarle la dirección IP el servidor Radius que estará usando, usando como SSID WIFIZONE.

Configuración de los host finales

Configurar en cada host el nombre de usuario y la contraseña.

Editar el perfil del host. Seleccionamos el perfil y le damos clic en Edit.

Se tendrá que seleccionar el SSID del Router de Acceso Inalámbrico que es al que nos queremos conectar y seleccionamos Advanced Setup.

Verificar que en efecto sea el SSID del Router de Acceso Inalámbrico que es al que nos queremos conectar (Lo podemos editar, si ese fuese el caso) y luego damos clic en Next.

Seleccionamos el tipo de encriptación que usa el Router de Acceso Inalámbrico y damos clic en Next.

Ingresar el usuario y contraseña que se añadió anteriormente a la base de datos del Servidor-Radius (El número de usuarios en Radius dependerá del número de clientes que se conecten a Router de Acceso Inalámbrico).

Creación y configuración de Listas de Control de Acceso (ACL)

Por medio de listas de acceso extendidas bloquear la comunicación entre las InterVLANs de la siguiente manera:

- a) VLAN 1 y VLAN 2 pueden comunicarse entre sí, pero no con la VLAN 3 y ambas tendrán acceso a la red externa.
- b) VLAN 3 no podrá comunicarse con VLAN 1 VLAN 2, pero si tendrá acceso a la red externa.

Configuración de SSH

Únicamente el PC del gerente tendrá acceso SSH al Router-Intranet, por lo que tendrá que configurarle el servicio y crear un usuario llamado Gerente_1 y la contraseña será Gerente123@.

Configuración de la VPN

Habilitar un túnel VPN para que el gerente pueda acceder desde el exterior o una infraestructura pública (PC Gerente) a la red del área DMZ. Con las siguientes características:

- El tipo de cifrado a usar es AES.
- El algoritmo es hash SHA.
- Para habilitar la fortaleza del intercambio de claves usar Deffie-Hellman.
- La clave a compartir es: TELEMÁTICA.
- Se creara un conjunto denominado ESTUDIANTES, aplicando el protocolo ESP.
- Aplicar 86400 segundos de validez.
- Crear una lista de acceso para identificar el tráfico que hará uso del túnel VPN.
- El mapa de cifrado se llamara TEL con un número de secuencia de 1000.

Nota: La configuración del túnel VPN se hará en los router más cercanos a los clientes VPN, en este escenario de prueba estos son (Router_Borde y Managua).

Tiempo estimado de solución

- 6 horas

Preguntas de análisis

1. ¿Cuál es la importancia del uso de SSH?
2. ¿Cuáles equipos fue necesario configurar para crear el túnel VPN?
3. ¿A la hora de configurar el túnel VPN, es necesario que las claves compartidas sean las mismas o perfectamente se podría hacer el intercambio de datos mediante claves diferentes?
4. ¿Cuál es el mecanismo o que comando usa (en este simulador) para verifícas si las claves se están compartiendo?
5. ¿Cuál es la diferencia de usar SSH y un túnel VPN?
6. ¿Qué otras medidas de seguridad usarías para tener una red lo más segura posible?

PRÁCTICA N° 15: **GESTIÓN DE SERVICIOS DE APLICACIÓN**

OBJETIVOS

General

- Configurar servicios de red por medio de una granja de servidores sobre la base de una intranet, habituales en muchas organizaciones empresariales.

Específicos

- Configurar el sistema de nombres de dominio (DNS), para habilitar la petición de servicios web, por medio de HTTP.
- Configurar el servicio de correo electrónico vía SMTP, a fin de implementar el envío de mensajes en el contexto de un dominio de correo electrónico.
- Implantar un servicio de transferencia de archivos vía FTP, con objeto de compartir los datos disponibles en una organización.
- Configurar un mecanismo de sincronización de red mediante el protocolo NTP, de gestión de red vía SNMP, así como por medio de la activación de Syslog, un sistema de registro de las actividades más notables detectadas en la red administrada.

INTRODUCCIÓN

Hoy en día en las redes de comunicaciones existen múltiples protocolos en el nivel de aplicación que ofrecen servicios necesarios a los clientes o usuarios finales.

En esta práctica se pretende configurar algunos protocolos y tecnologías de máxima importancia, los cuales operan en las redes de la actualidad en lo que corresponde al nivel de aplicación, tales protocolos son: HTTP, FTP, DNS, SMTP, SNMP, NTP y Syslog. Esto con el fin de que el estudiante adquiera los conocimientos básicos en el área de administración de servicios de redes.

REQUERIMIENTOS

➤ Hardware

- Computadora con los siguientes requisitos:
 - Procesador mínimo de velocidad de 2.1 GHz
 - Memoria RAM de 1 GB.

➤ Software

- Simulador Cisco Packet Tracer 6.0 con los siguientes elementos:
 - 6 Routers 1841.
 - 4 Switches 2960.
 - 7 Servidores.
 - 2 Laptops.
 - 4 PCs.

CONOCIMIENTOS PREVIOS

Para la realización de ésta práctica es necesario que el alumno tenga conocimientos acerca del funcionamiento básico de los protocolos bases de la capa de aplicación, y de cómo operan de manera conjunta en una intranet de una empresa que ofrezca ya sea servicio de Correo, FTP, Web entre otros.

TOPOLOGÍA

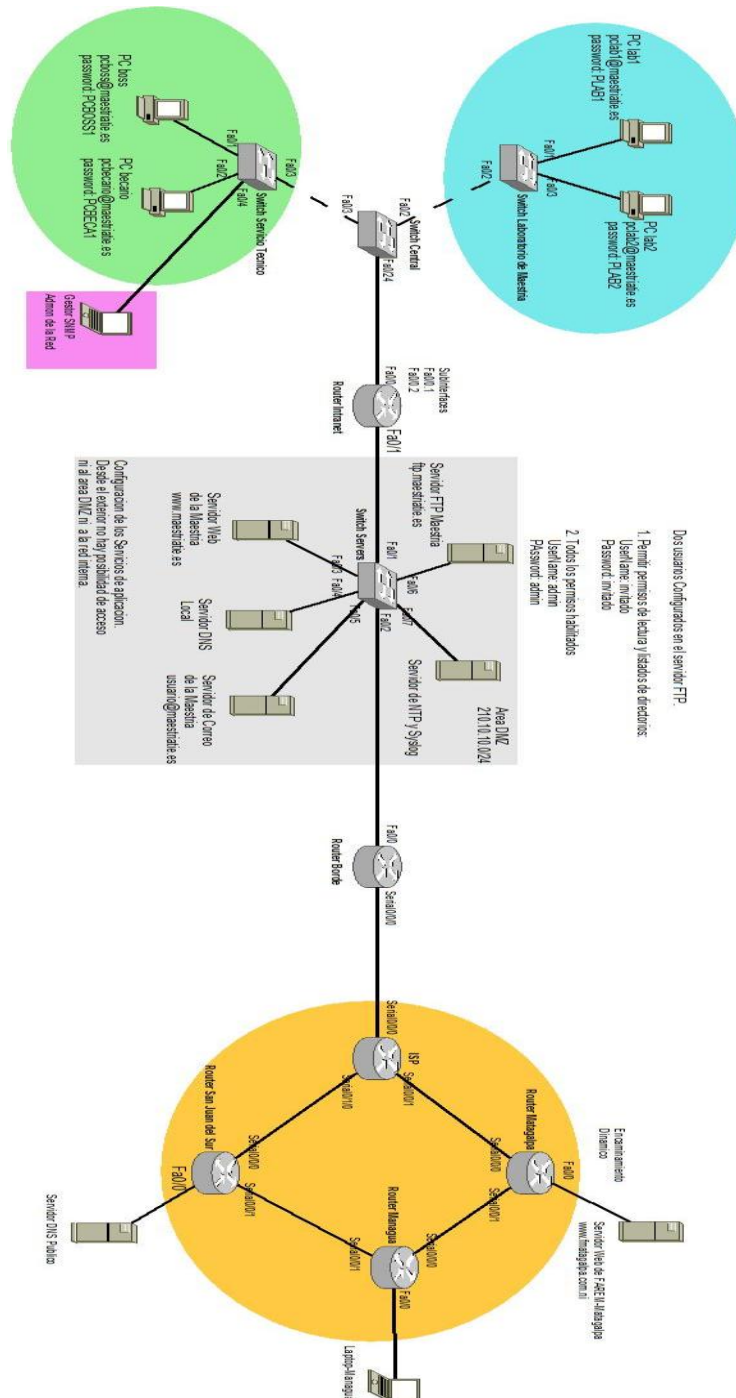


Figura 114. Topología de práctica de Gestión de servicios de aplicación

FUNCIONALIDAD

La Figura 114 muestra la topología que en esta práctica se estará configurando, dicha topología simula una red intranet de una empresa privada, la cual ofrece los siguientes servicios:

- Servicio Web para el alojamiento del sitio Web de la empresa.
- Servicio de FTP para el alojamiento de archivos del personal de la empresa.
- Servicio de Correo electrónico.
- Servicio de DNS para la traducción de dominios.
- Servicio de gestión de red, para monitorear los dispositivos que conforman la red.

Además de lo antes mencionado, también los miembros de la empresa tendrán la facultad de poder navegar a internet y visitar sitios web alojados en otros servidores. Se crearán diferentes VLANs, y habrá un administrador que se encargará de administrar y llevar el control del funcionamiento de los dispositivos de la red.

COMANDOS DE AYUDA

Comando	Descripción
ip dhcp pool {nombre de vlan}	Crea un pool para el servicio dhcp
network {dirección de red}	Dirección de red de un pool.
ip dhcp excluded-address [dirección IP]	Excluye una dirección IP cuando el pool comienza a asignar a los clientes.
default-router {dirección IP}	Indica un default gateway a los clientes del pool
dns-server {dirección IP DNS}	Indica un servidor DNS a los clientes del pool
ip route {origen IP – mascara} {siguiente salto IP}	Establece una ruta estática.
access-list {número de lista} {deny permit} {dirección de red} {mascara wildcard}	Establece una lista de acceso.
ip nat {inside outside}	Establece una interfaz operando NAT.
ip nat inside source list {número de lista de acceso} interface {interfaz de salida} overload	Establece la operación de traducción de las redes a la interfaz de salida.
ntp authentication-key [key] md5 [password]	Habilita el protocolo NTP indicándole la llave y password del servidor de NTP.
ntp server 210.10.10.120 key 777	Indica la dirección del servidor NTP y su llave
snmp-server community {nombre de la comunidad} {RO RW}	Crea una comunidad SNMP para el monitoreo.

DATOS DE LOS DISPOSITIVOS

Switches capa 2		
Equipo	ID de VLAN	Puertos access
Switch Laboratorio de Maestría	2	Fa0/1, Fa0/3
Switch Servicio Técnico	3	Fa0/1, Fa0/2, Fa0/4
Switch Central	-----	-----
Switch Server	-----	-----

Routers			
Equipo	Interfaz	Dirección IP	Dirección de subred
Router intranet	Fa0/1	210.10.10.2/24	210.10.10.0/24
Router Borde	Fa0/0	210.10.10.1/24	210.10.10.0/24
	Serial0/0/0	215.15.15.1/24	215.15.15.0/24
ISP	Serial0/0/0	215.15.15.2/24	215.15.15.0/24
	Serial0/0/1	216.16.16.1/24	216.16.16.0/24
	Serial0/1/0	217.17.17.1/24	217.17.17.0/24
Router Matagalpa	Serial0/0/0	216.16.16.2/24	216.16.16.0/24
	Serial0/0/1	218.18.18.1/24	218.18.18.0/24
	Fa0/0	193.200.20.1/24	193.200.20.0/24
Router Managua	Serial0/0/0	218.18.18.2/24	218.18.18.0/24
	Serial0/0/1	219.19.19.2/24	219.19.19.0/24
	Fa0/0	193.146.57.1/24	193.146.57.0/24
Router San Juan del Sur	Serial0/0/0	217.17.17.2/24	217.17.17.0/24
	Serial0/0/1	219.19.19.1/24	219.19.19.0/24
	Fa0/0	192.128.50.1/24	192.128.50.0/24

Subinterfaces				
Equipo	Subinterfaz	Dirección IP	Dirección de subred	ID de VLAN
Router Intranet	Fa0/0.1	202.2.2.1/24	202.2.2.0/24	2
	Fa0/0.2	203.3.3.1/24	203.3.3.0/24	3

Servidores			
Equipo	Dirección IP	Gateway	Dirección de Subred

WEB Maestría	210.10.10.200/24	210.10.10.2/24	210.10.10.0/24
FTP Maestría	210.10.10.150/24	210.10.10.2/24	210.10.10.0/24
DNS Local	210.10.10.100/24	210.10.10.2/24	210.10.10.0/24
Correo Maestría	210.10.10.250/24	210.10.10.2/24	210.10.10.0/24
NTP y Syslog	210.10.10.120/24	210.10.10.2/24	210.10.10.0/24
Web de FAREM - Matagalpa	193.200.20.23/24	193.200.20.1/24	193.200.20.0/24
DNS público	192.128.50.144/24	192.128.50.1/24	192.128.50.0/24

Usuarios de Correo	
Nombre de Usuario	Contraseña
pclab1	PLAB1
pclab2	PLAB2
pcboss	PCBOSS1
pcbecario	PCBECA1

Dominios	
Nombre de dominio	Dirección IP
www.maestriatie.es	210.10.10.200
www.fmatagalpa.com.ni	193.200.20.23
maestriatie.es	210.10.10.250
ftp.maestriatie.es	210.10.10.150

PC			
Equipo	Dirección IP	Gateway	Dirección de Subred
Laptop Managua	193.146.57.32/24	193.146.57.1/24	193.146.57.0/24

ENUNCIADO

Asignación de direcciones IP a los dispositivos

Para dar inicio a esta práctica, se deberá establecer de manera correcta las direcciones IP correspondiente a cada una de las interfaces o subinterfaces de los dispositivos (Routers, Laptop-Managua, Servidores) internos y externos que la ameriten. Dicha información se provee en los cuadros correspondientes a estos equipos.

Creación de VLANs

En esta fase se implementarán o se crearán las VLANs que en este caso pertenecerán a los departamentos de laboratorio de maestría y de servicio técnico.

Para cumplir con lo antes mencionado, en los dispositivos llamados Switch laboratorio de maestría y Switch servicio técnico configurar los puertos configurar modo access con su respectivo ID de VLANs a como se indica en el cuadro Switches de capa 2.

Configuración de Subinterfaces

En el router intranet se deberá de crear las subinterfaces pertenecientes a cada VLANs, de manera que el tráfico generado por los usuarios y que este lleve etiqueta de cualquier VLANs pueda ser encaminado mediante una misma interfaz física con la implementación de subinterfaces. Para cumplir con la tarea antes mencionada ver cuadro subinterfaces que muestra las subinterfaces en éste dispositivo con sus respectivas direcciones IP y el ID de VLAN.

Servicio de DHCP

En este apartado en el router intranet se deberá crear dos pools llamados vlan2 y vlan3, esto con el objetivo de que cada usuario perteneciente a cada una de las VLANs reciba las direcciones IP mediante DHCP. Recuerde establecer a cada pool una puerta de enlace que en este caso será la dirección de la subinterfaz, al igual que un dominio DNS para la resolución de nombres a direcciones.

Configuración de protocolo de enrutamiento (estático)

En el Router intranet se deberá establecer una ruta por defecto, la cual encaminará el tráfico generado por cualquiera de las VLANs de la intranet con destino al Router de Borde.

En el Router de Borde se deberá establecer rutas por defecto que permitan el paso del tráfico proveniente del exterior con destino solamente a los usuarios de las VLANs, de igual manera, es este dispositivo se creara una ruta por defecto que permita el paso del trafico proveniente de la intranet con destino al exterior (INTERNET).

Configuración de protocolo de enrutamiento (RIP)

En los enrutadores externos (Router ISP, Router San Juan del Sur, Router Matagalpa, y Router Managua) implementar RIP para el encaminamiento de tráfico entre las redes externas y la red LAN.

Configuración de NAT

En el router de borde se implementará la tecnología NAT sobrecargado para la traducción de direcciones, de manera que los usuarios de la intranet salgan al exterior mediante una única dirección pública.

Se crearán las listas de acceso indicando las direcciones locales internas que se deberán traducir.

Especificar la interfaz interna y marcarla como conectada al interior.

Especificar la interfaz externa y marcarla como conectada al exterior.

Asocie la lista de acceso, especificando la interfaz de salida para la traducción.

Configuración de Servicios de Aplicación

En este apartado de la práctica para cumplir con el objetivo más importante, se estarán configurando servicios de la capa de aplicación tales como web, correo, FTP, DNS, NTP y SNMP.

Servicio Web

En el servidor web de la red LAN llamado servidor web de la maestría y en el servidor web externo llamado Servidor web de FAREM-Matagalpa, asignar su respectiva dirección IP, default Gateway y la dirección del DNS a como se muestra en el cuadro llamado Servidores. Recuerde que para los usuarios de la intranet el DSN será el local y para los usuarios externos el DNS será el público.

En el Servidor Web de la Maestría editar el index.html y establecer un letrero de bienvenida que diga Bienvenidos al laboratorio de COMPUTACIÓN, de la misma manera se hará con el servidor llamado Servidor Web FAREM-Matagalpa, este último con un letrero que diga Bienvenidos a FAREM-Matagalpa.

Servicio DNS

En el servidor DNS de la red LAN llamado Servidor DNS local y en el servidor DNS externo llamado Servidor DNS público, asignar su respectiva dirección IP y default Gateway a como se muestra en el cuadro Servidores.

Agregar a este equipo el nombre de dominio con su respectiva dirección IP de los servicios a los que se les estará brindando servicios de resolución de nombres como: FTP, Correo, y Web. Recuerde agregar también el dominio y dirección IP del servicio Web externo.

En el servidor DNS externo llamado DNS publico agregar el dominio y dirección IP del servicio Web de FAREM-Matagalpa para que los usuarios externos puedan visitar este sitio Web mediante un nombre de dominio.

NOTA: Los nombres de dominios con sus respectivas direcciones IP se proporcionan en el cuadro llamado dominios.

Servicio FTP

En el servidor FTP de la red LAN llamado servidor FTP maestría, asignar su respectiva dirección IP, default Gateway y la dirección del DNS local a como se muestra en el cuadro llamado Servidores.

Agregar a este equipo dos usuarios. El primero es el usuario que se llamará invitado y su contraseña será invitado, a éste usuario se le otorgará permisos de lectura y listado de directorios.

Al segundo usuario se llamará admin y su contraseña será admin, a éste usuario se le otorgará todos los permisos.

Servicio de Correo

En el servidor de correo de la red LAN llamado servidor de Email de la maestría, asignar su respectiva dirección IP, default Gateway y la dirección del DNS local a como se muestra en el cuadro llamado Servidores.

En este equipo se deberán agregar los usuarios que tendrán acceso al servicio de correo de la intranet. La información de estos usuarios se muestra en el cuadro usuarios de correo.

Para hacer las pruebas de funcionamiento para el envío y recepción de mensajes mediante correo electrónico, se deberá configurar en cada equipo final (PC) la aplicación E-mail que trae por defecto cada usuario final (PC) que pertenezcan de la red LAN.

Servicio de gestión de red (NTP, SNMP, Syslog)

En el servidor llamado Servidor de NTP y Syslog, asignar su respectiva dirección IP y default Gateway a como se muestra en el cuadro llamado Servidores.

En este equipo se deberá configurar NTP indicando una llave y una clave, en este caso el valor de la llave es 777 y el de la clave será telemática. Recuerde ajustar la fecha de manera correcta, al igual que también habilitar el servicio de Syslog.

En los equipos Router intranet y Router salida al exterior habilitar el protocolo NTP, indicando la llave y clave (777, telemática) que se estableció en el servidor NTP anteriormente, esto para lograr una sincronización de los equipos autenticándose con el servidor NTP.

Indicar en estos dos equipos (Router intranet y Router salida al exterior) la dirección IP del servidor de Syslog y la llave de autenticación, esto para generar mensajes de logs al servidor de cualquier anomalía que pase en estos equipos.

En los equipos Router intranet y Router Salida exterior que serán nuestros agentes SNMP, establecer 2 comunidades las cuales permitirán tanto hacer cambios en cada equipo al igual que solo hacer consultas mediante SNMP. Las comunidades serán MONITORIZACION (solo lectura) y AUDITORIA (lectura y escritura).

Para hacer las pruebas de monitorización de estos 2 equipos, en el equipo (laptop) llamado Gestor SNMP, configurar la aplicación MIB Browser que nos permitirá monitorizar los equipos Router intranet y Router Salida exterior.

En configuración avanzada de esta aplicación se debe indicar la dirección IP del equipo a monitorizar, el puerto por el cual trabajará cada equipo (161), establecer las dos comunidades creadas anteriormente en los agentes SNMP (Router intranet y Router Salida Exterior) y la versión SNMP que se estará ejecutando.

Indicar el OID mediante la MIB para obtener o establecer un valor al equipo gestionado.

Tiempo estimado de solución

➤ 4 horas

Preguntas de Análisis

1. Cite cinco aplicaciones de Internet no propietarias y los protocolos de la capa de aplicación que utilizan.
2. Para una sesión de comunicación entre dos host, ¿Cuál es el host cliente y cuál es el host servidor?
3. ¿Qué información utiliza un proceso que se ejecuta en un host para identificar un proceso que se ejecuta en otro host?
4. ¿Por qué se ejecutan HTTP, FTP, SMTP, POP3 e IMAP sobre TCP en vez de sobre UDP?
5. ¿Cuál es la diferencia entre HTTP persistente con entubamiento y sin entubamiento? ¿Cuál es el utilizado por HTTP/1.1?
6. ¿Por qué se dice que FTP envía información de control fuera de banda?
7. Desde la perspectiva del usuario, ¿cuál es la diferencia entre el modo descargar-y-borrar y descargar-y-guardar en POP3?
8. Cada host en internet tiene al menos un servidor de nombres local y uno autorizado. ¿Qué papel tiene cada uno de estos servidores en el DNS?
9. ¿Es posible que un servidor web y un servidor de correo de una organización tengan exactamente el mismo alias de un nombre de host (por ejemplo, Telematica.com)? ¿Cuál sería el tipo de RR que contiene el nombre del host del servidor de correo?
10. ¿Por qué se dice que HTTP envía los comandos de control en banda?
11. ¿Por qué se beneficia un gestor de red de tener herramientas de gestión?
12. ¿Cuáles son las cinco áreas de la gestión de red que define ISO?
13. ¿Cuál es la diferencia de la gestión de red y la de servicios?
14. ¿Cuáles son los siete tipos de mensajes que se utilizan en SNMP?
15. ¿Cuáles son los tres componentes principales en una arquitectura de gestión de red?

CAPÍTULO N° 4: ASPECTOS FINALES

3. CONCLUSIONES

Con la finalización de este trabajo monográfico, consideramos que hemos logrado cumplir con los objetivos propuestos, llegando a las siguientes conclusiones:

1. La combinación de temas teóricos con ejercicios prácticos, le proporciona al estudiante una base sólida para dar solución a las prácticas.
2. El diseño de un adecuado formato de prácticas facilitará a los estudiantes una correcta comprensión de la práctica a realizar.
3. La tarea de realizar cada una de las prácticas propuestas, se hará más fácil a los estudiantes debido al orden de complejidad con la cual se han desarrollado y organizado.
4. Se han abordados temas importante, que en su mayoría son puesto en marcha en el ámbito laboral en el área de redes.

Con todos los aspectos antes mencionados se puede afirmar que se ha desarrollado un documento sencillo y practico que resuelve las problemáticas o necesidades planteadas en la Definición de Problemas del presente trabajo.

4. RECOMENDACIONES

Las recomendaciones que se describen a continuación son a base de una futura actualización del documento.

1. Los temas que este documento presenta son considerados de suma importancia para el desarrollo de aprendizaje de los alumnos de la carrera de Ingeniería en Telemática en lo que concierne al área de redes de computadores, sin embargo, deben actualizarse continuamente con nuevos temas de importancia.
2. Es necesario que el documento siempre lleve una secuencia lógica de las prácticas propuestas, esto con el fin de que el alumno lleve un control gradual de las tecnologías y protocolos empleados.
3. En razón a la importancia del surgimiento del protocolo IPv6, es necesario que para una futura actualización de este documento las prácticas a desarrollar implementen conjuntamente protocolo IPv4 como IPv6, con el objetivo de que el alumno tenga una visión de cómo operan dichos protocolos.
4. En lo que corresponde a la parte de desarrollo de las prácticas en este documento, es necesario siempre hacer un evaluado al final de cada una, de manera que el maestro y el alumno puedan medir el grado de aprendizaje adquirido en la teoría así como en la práctica.

5. BIBLIOGRAFÍA

➤ Libros físicos consultados

- Kurose, J.F. (SF). Redes de Computadores Un Enfoque Descendente Basado en Internet.
- Stalling, S, William. 2004. Comunicaciones y Redes de Computadores. 7 ed. Madrid, PEARSON EDUCACION, S.A.
- Sampieri, R,H; Collado, C, F; Baptista, L, P. 2010. Metodología de la investigación. 5 ed. McGrawHill

➤ Libros de la Web consultados

- Farrel, A. (SF). The Internet And ITS Protocols.). (Disponible en: <http://books.google.com.ni/books?hl=es&lr=&id=LtBegQowqFsC&oi=fnd&pg=PP2&dq=MKP++The+Internet+and+Its+Protocols,+A+Comparative+Approach&ots=QiiWvnEQMb&sig=dWIE2BNSIzRkYc3FZce3KN6UITE#v=onepage&q=MKP%20%20The%20Internet%20and%20Its%20Protocols%2C%20A%20Comparative%20Approach&f=false>)
- Hucaby, D. (2004). CCNP BCMSN Exam Certification Guide. Indianapolis USA.). (Disponible en: <http://docstore.mik.ua/cisco/pdf/routing/Cisco%20Press%20%20CCNP%20BCMSN%20Exam%20Certification%20Guide.pdf>)

➤ Tesis consultadas

- Larios Acuña, A.E; Mayorga Castellón, I.R; Moreira Cárcamo, B.M. (21 de mayo del 2008). Prácticas de laboratorios para la asignatura de redes de ordenadores II. León Nicaragua.

➤ Documentos de la Web consultados

- Allied Telesis. (SF). How to Configure VRRP. (Disponible en: <http://ptgmedia.pearsoncmg.com/images/0201715007/samplechapter/srikanthch02.pdf>)
- Cisco, Systems. Cisco Networking Academy Program, CCNA 3 Y 4, Módulo 1 CCNA3 (RIP). 3 ed. (Disponible en : <https://rapidshare.com/#!download|806p6|1370261463|CCNA-3-4.pdf|23000|0|0|1|referer-7B3ABFC2D0BCF13D1AD9FE3B94944F6B>)
- Cisco, Systems. Cisco Networking Academy Program, CCNA 3 Y 4. Módulo 2 CCNA3 (OSPF). 3 ed. (Disponible en : <https://rapidshare.com/#!download|806p6|1370261463|CCNA-3-4.pdf|23000|0|0|1|referer-7B3ABFC2D0BCF13D1AD9FE3B94944F6B>)

- Cisco, Systems. Cisco Networking Academy Program, CCNA 3 Y 4, Módulo 3 CCNA3 (EIGRP). 3 ed. (Disponible en : <https://rapidshare.com/#!download|806p6|1370261463|CCNA-3-4.pdf|23000|0|0|1|referer-7B3ABFC2D0BCF13D1AD9FE3B94944F6B>)
- Cisco, Systems. Cisco Networking Academy Program, CCNA 3 Y 4, Módulo 3 CCNA4 (PPP). 3 ed. (Disponible en : <https://rapidshare.com/#!download|806p6|1370261463|CCNA-3-4.pdf|23000|0|0|1|referer-7B3ABFC2D0BCF13D1AD9FE3B94944F6B>)
- Cisco, Systems. Cisco Networking Academy Program, CCNA 3 Y 4, Módulo 5 CCNA4 (Frame Relay). 3 ed. (Disponible en : <https://rapidshare.com/#!download|806p6|1370261463|CCNA-3-4.pdf|23000|0|0|1|referer-7B3ABFC2D0BCF13D1AD9FE3B94944F6B>)
- Cisco. Cisco PIX Firewall and VPN Configuration Guide, Version 6.3. San Jose, CA (Disponible en: <http://www.cisco.com/en/US/docs/security/pix/pix63/release/notes/pixrn634.pdf>)
- Cisco System. (Febrero 26 del 2010). USA.). (Disponible en: http://www.cisco.com/en/US/docs/ios/ios_xe/ipapp/configuration/guide/ipapp_vrrp_xe.pdf)
- Cisco System. (SF). High Availability Configuration Guide VRRP). (Disponible en: http://www.h3c.com/portal/Technical_Support___Documents/Technical_Documents/Switches/H3C_S12500_Series_Switches/Configuration/Operation_Manual/H3C_S12500_CG-Release1825P016W180/12/201302/774306_1285_0.htm#_Toc347309552)
- Eduardo, C. Cabeza. Fundamentos de Routing, Capítulo 6 (IS-IS). 3 ed. (Disponible en : http://eduangi.com/wp-content/uploads/2009/07/ebook_fundamentos_de_routing.pdf)
- Fleury Vicencio, D.C. (Junio 2007). (Disponible en: <http://cdigital.uv.mx/bitstream/123456789/31219/1/cesareofleurivicencio.pdf>)
- Masapanta, N, J. 2013. Implementación de una central PBX IP en la empresa ECUADORIANPIPE. LTDA basada en plataforma de Software Libre. Tesis Tec Elect. Y Tel. Escuela Politécnica Nacional, Quito-Ecuador.(Disponible en: <http://bibdigital.epn.edu.ec/bitstream/15000/6019/1/CD-4771.pdf>)
- Moreno, J.I.; Soto, I; Larrabeiti, D. Protocolos de Señalización para el transporte de Voz sobre redes IP. Universidad Carlos III de Madrid. (Disponible en: <http://www.it.uc3m.es/~jmoreno/articulos/protocolssenalizacion.pdf>)
- NETGEAR, Inc. 2005. Virtual Private Networking Basics. Santa Clara, CA. (<http://documentation.netgear.com/reference/nld/vpn/pdfs/FullManual.pdf>)

- Ochoa, A, J. s.f. El protocolo SSH. Tesis Doc. Ciencias de Computación - ESIDE (Disponible en: <http://www.naguissa.com/archivos.php?path=11&accion=descarga&IDarchivo=58>)
- Patrick J. Conlan. Cisco Network Professional's Advanced Internetworking Guide, Capítulo 10 (HSRP), 3 ed. (Disponible en : <http://books.google.es/books?id=x33i7TwBhHQC&pg=PA769&dq=cisco+networking+professional%C2%B4s+advanced+internetworking+guide&hl=es&sa=X&ei=WpIWUorwHeKf2QWzloCYCw&ved=0CD4Q6AEwAA#v=onepage&q=cisco%20networking%20professional%C2%B4s%20advanced%20internetworking%20guide&f=false> (Cisco Network Professional's Advanced Internetworking)
- Palet, J. 2007. Introducción a IPv6. Isla Margarita. (Disponible en: http://lacnic.net/documentos/lacnicx/introduccion_ipv6_v11.pdf)
- Ros, T; Hortal, S; Pérez, A. (2008). Tecnología EtherChannel). (Disponible en: <http://etherchanneleupmt.googlecode.com/files/TRABAJO%20ETHERCHANNEL.pdf>)
- Sánchez, S, Pedro. 2006. Estudio y Desarrollo de Soluciones VoIP para la Conexión Telefónica entre Sedes Internacionales de una misma Empresa. Tesis Ing. en Tel. UNIVERSIDAD POLITÉCNICA DE CARTAGENA. (Disponible en: <http://repositorio.bib.upct.es/dspace/bitstream/10317/232/1/pfc1836.pdf>)
- Unidad Docente de Redes. CONFIGURACIÓN DE ROUTERS: LISTAS DE CONTROL DE ACCESO (ACLs); Departamento de Informática, Universidad de Castilla-La Mancha. (Disponible en: http://virtualbook.weebly.com/uploads/2/9/6/2/2962741/ac__list.pdf).