

UNIVERSIDAD MIGUEL HERNÁNDEZ DE ELCHE

ESCUELA POLITÉCNICA SUPERIOR DE ELCHE

GRADO EN INGENIERÍA DE TECNOLOGÍAS DE
TELECOMUNICACIÓN



"ESTUDIO COMPARATIVO DE
MECANISMOS DE INTERCONEXIÓN
ENTRE SISTEMAS AUTÓNOMOS PARA
L3VPN SOBRE MPLS"

TRABAJO FIN DE GRADO

Junio -2026

AUTOR: Adrián Padilla Rosales

DIRECTOR/ES: Pedro Juan Roig Roig
Cristina Bernad Cantó

ESTUDIO COMPARATIVO DE MECANISMOS DE INTERCONEXIÓN ENTRE SISTEMAS AUTÓNOMOS PARA L3VPN SOBRE MPLS

Resumen:

El presente Trabajo Fin de Grado estudia y valida distintos mecanismos de interconexión entre sistemas autónomos para el transporte de servicios de red privada virtual de capa 3 en entornos MPLS. Para ello, se han implementado en GNS3 tres escenarios estándar de interconexión, de acuerdo con la RFC 4364, denominados escenario A, escenario B y escenario C, con el objetivo de analizar sus diferencias de funcionamiento, su complejidad de configuración y su comportamiento desde el punto de vista del encaminamiento y de la conectividad extremo a extremo. El trabajo parte de una base teórica centrada en tecnologías como OSPF, MPLS, BGP y las VRF, incorporando además el estudio de distintas familias de direcciones de BGP, como IPv4 unicast, VPNv4 y IPv4 labeled-unicast, necesarias para comprender la lógica de cada escenario. A nivel práctico, cada escenario se ha desplegado sobre topologías con distinto número de routers de core, con el fin de observar tanto el caso más simple como una versión más representativa y extrapolable a redes con un número arbitrario de nodos intermedios. La validación se ha realizado mediante comandos de verificación, pruebas de conectividad y análisis del camino seguido por el tráfico entre extremos. Los resultados obtenidos permiten comparar las ventajas y limitaciones de cada alternativa y muestran cómo varía la forma de transportar la información de encaminamiento entre sistemas autónomos en función de la solución adoptada. En conjunto, el trabajo ofrece una visión práctica y razonada de tres modelos clásicos de interconexión inter-AS, prestando especial atención a su aplicación en laboratorios de emulación y a su valor formativo dentro del ámbito de las redes de telecomunicación.

Palabras clave:

MPLS, BGP, MP-BGP, VRF, OSPF, interconexión, L3VPN, GNS3, VPNv4, IPv4 labeled-unicast

COMPARATIVE STUDY OF INTERCONNECTION MECHANISMS BETWEEN AUTONOMOUS SYSTEMS FOR L3VPN OVER MPLS

Abstract:

This Final Degree Project studies and validates different interconnection mechanisms between autonomous systems for the transport of Layer 3 virtual private network services in MPLS environments. To achieve this, three standard interconnection scenarios have been implemented in GNS3, in accordance with RFC 4364, referred to as scenario A, scenario B and scenario C, with the aim of analysing their operational differences, configuration complexity and behaviour from the perspective of routing and end-to-end connectivity. The work starts from a theoretical basis focused on technologies such as OSPF, MPLS, BGP and VRF, also including the study of different BGP address families, such as IPv4 unicast, VPNv4 and IPv4 labeled-unicast, which are necessary to understand the logic of each scenario. At a practical level, each scenario has been deployed on topologies with a different number of core routers, in order to observe both the simplest case and a more representative version that can be extrapolated to networks with an arbitrary number of intermediate nodes. The validation has been carried out through verification commands, connectivity tests and analysis of the path followed by the traffic between endpoints. The results obtained make it possible to compare the advantages and limitations of each alternative and show how the way routing information is transported between autonomous systems varies depending on the solution adopted. Overall, the project provides a practical and reasoned view of three classic inter-AS interconnection models, paying special attention to their application in emulation laboratories and to their educational value within the field of telecommunication networks.

Keywords:

MPLS, BGP, MP-BGP, VRF, OSPF, interconnection, L3VPN, GNS3, VPNv4, IPv4 labeled-unicast.

AGRADECIMIENTOS

En primer lugar, quiero agradecer a mis tutores, Pedro y Cristina, por su orientación, paciencia y ayuda durante el desarrollo de este Trabajo Fin de Grado, ya que sus consejos han sido fundamentales para poder avanzar y dar forma a este proyecto. También quiero hacer una mención especial a todos los profesores de la carrera, porque de una forma u otra todos han aportado gran importancia en mi formación y me han ayudado a llegar hasta este momento.

También quiero agradecer profundamente a mis padres, porque sin ellos este camino habría sido mucho más difícil. Gracias por haberme apoyado económicamente, pero sobre todo por haber estado ahí emocionalmente, por confiar en mí incluso en los momentos en los que yo mismo dudaba y por darme siempre la tranquilidad necesaria para poder seguir adelante. Este trabajo también es resultado de todos esos esfuerzos silenciosos que muchas veces no se ven, pero que han estado presentes durante toda la carrera. Por todo ello, les estaré siempre agradecido.

Quiero dedicar también unas palabras a mis compañeros de clase, porque compartir este camino con ellos ha hecho que la experiencia universitaria sea mucho más llevadera y especial. En especial, quiero mencionar a Sergio y Mireya, y también a todo mi grupo de compañeros y amigos de la universidad, ya que gracias a ellos los días difíciles han sido más fáciles y los buenos momentos han sido todavía mejores. Haber podido vivir esta etapa rodeado de personas con las que estudiar, aprender, reír y apoyarnos mutuamente ha sido una de las partes más bonitas de estos años.

También quiero agradecer a mis amigos, empezando por César, y después a Aitana, Alberto, Álvaro, Andrea, Dani, Eli, Janine, Kautar, Lucía, María, Mary, Mercedes, Miriam, Rocío y Sergi. Gracias por estar presentes, por escucharme, por animarme y por formar parte de mi vida durante esta etapa. Aunque no pueda nombrar a todos para no alargar demasiado la lista, quiero hacer extensivo este agradecimiento a todas las personas que han estado ahí de alguna manera, porque cada conversación, cada apoyo y cada momento compartido también han formado parte de este proceso.

Por último, quiero dedicar este Trabajo Fin de Grado a la memoria de mi abuelo José. Este trabajo también va por él, por todo lo que significó y sigue significando para mí. El tiempo que pasé con él ha sido de los mejores momentos que he podido vivir, y su recuerdo siempre me acompañará. Esta dedicatoria es una forma de tenerlo presente en un momento tan importante de mi vida y de agradecerle, aunque ya no esté físicamente, todo el cariño y todos los recuerdos que me dejó.



ÍNDICE

1. INTRODUCCIÓN.....	14
1.1. Contexto del problema y motivación del trabajo.....	14
1.2. OSPF como protocolo interno de encaminamiento en el núcleo.....	15
1.3. MPLS y el transporte de servicios VPN de capa 3.....	15
1.4. BGP y extensiones multiprotocolo en redes L3VPN	17
1.5. VRF, RD y RT en redes L3VPN	18
1.6. Alternativas de interconexión entre sistemas autónomos.....	19
1.6.1. Opción A	20
1.6.2. Opción B.....	21
1.6.3. Opción C.....	22
1.7. Entorno de emulación y objetivos del trabajo	23
2. MATERIAL Y MÉTODOS	25
2.1. Router c3600	25
2.2. PCs.....	27
2.2.1. VPCS	27
2.2.2. Linux Core.....	28
2.3. Ethernet Switch	30
2.4. Escenario A	31
2.4.1. Escenario A con dos routers	33
2.4.2. Escenario A con tres routers	34
2.5. Escenario B.....	36
2.5.1. Escenario B con dos routers	37
2.5.2. Escenario B con tres routers	38
2.6. Escenario C.....	39
2.6.1. Escenario C con dos routers	41
2.6.2. Escenario C con tres routers	42
3. RESULTADOS Y DISCUSIÓN	44
3.1. Escenario A con dos routers	45
3.1.1. Pruebas de conectividad mediante ping	45
3.1.2. Trazado de ruta mediante traceroute	47
3.1.3. Verificación de rutas en la VRF	48
3.1.4. Verificación de BGP en la VRF	49

3.2 Escenario A con tres routers	51
3.2.1. Pruebas de conectividad mediante ping	51
3.2.2. Trazado de ruta mediante traceroute	52
3.2.3. Verificación de vecinos OSPF.....	53
3.2.4. Verificación de vecinos LDP.....	54
3.2.5. Verificación de la tabla de reenvío MPLS.....	55
3.2.6. Verificación de rutas en la VRF	55
3.2.7. Verificación de BGP en la VRF	56
3.3 Escenario B con dos routers	57
3.3.1. Pruebas de conectividad mediante ping	57
3.3.2. Trazado de ruta mediante traceroute	59
3.3.3. Verificación de la tabla de reenvío MPLS.....	59
3.3.4. Verificación de rutas en la VRF	60
3.3.5. Verificación de rutas VPNv4 en los ASBR.....	61
3.3.6. Verificación de rutas recibidas y anunciadas por BGP VPNv4	62
3.3.7. Comparación de tablas BGP VPNv4 en PE1 y ASBR1	63
3.3.8. Resumen de vecinos BGP VPNv4	64
3.4 Escenario B con tres routers	65
3.4.1. Pruebas de conectividad mediante ping	65
3.4.2. Trazado de ruta mediante traceroute	66
3.4.3. Verificación de vecinos OSPF.....	66
3.4.4. Verificación de vecinos LDP.....	67
3.4.5. Verificación de la tabla de reenvío MPLS.....	68
3.4.6. Verificación de rutas en la VRF	68
3.4.7. Verificación de rutas VPNv4 en los ASBR.....	69
3.5 Escenario C con dos routers	70
3.5.1. Pruebas de conectividad mediante ping	70
3.5.2. Trazado de ruta mediante traceroute	71
3.5.3. Verificación de rutas IPv4 etiquetadas en BGP.....	71
3.5.4. Verificación de la tabla de reenvío MPLS.....	72
3.5.5. Verificación de rutas en la VRF	73
3.5.6. Verificación de rutas VPNv4 en PE1	73
3.5.7. Resumen de vecinos BGP VPNv4	74

3.6 Escenario C con tres routers	75
3.6.1. Pruebas de conectividad mediante ping	75
3.6.2. Trazado de ruta mediante traceroute	76
3.6.3. Verificación de vecinos OSPF.....	76
3.6.4. Verificación de vecinos LDP.....	77
3.6.5. Verificación de la tabla de reenvío MPLS en P1.....	77
3.6.6. Verificación de rutas IPv4 etiquetadas en BGP.....	78
3.6.7. Verificación de la tabla de reenvío MPLS en ASBR1	78
3.6.8. Verificación de rutas en la VRF	79
3.6.9. Verificación de alcanzabilidad entre loopbacks de PE.....	80
3.6.10. Resumen de vecinos BGP VPNv4	81
3.7. Resumen comparativo de ventajas y desventajas de los escenarios.....	81
4. CONCLUSIONES.....	86
5. ANEXO	88
5.1. VPCS	88
5.1.1. PC1	88
5.1.2. PC2	89
5.2. LinuxCore.....	89
5.2.1. LinuxCore1	90
5.2.2. LinuxCore2.....	92
5.3. Router CE	93
5.3.1. Router CE1	93
5.3.2. Router CE2	95
5.4. Escenario A con dos routers	96
5.4.1. Router PE1.....	96
5.4.2. Router ASBR1	98
5.4.3. Router ASBR2.....	99
5.4.4. Router PE2.....	100
5.5. Escenario A con tres routers	101
5.5.1. Router PE1.....	102
5.5.2. Router P1	103
5.5.3. Router ASBR1	104
5.5.4. Router ASBR2.....	106
5.5.5. Router P2	107
5.5.6. Router PE2.....	108

5.6. Escenario B con dos routers	110
5.6.1. Router PE1.....	110
5.6.2. Router ASBR1	112
5.6.3. Router ASBR2.....	113
5.6.4. Router PE2.....	115
5.7. Escenario B con tres routers	116
5.7.1. Router PE1.....	117
5.7.2. Router P1	118
5.7.3. Router ASBR1	119
5.7.4. Router ASBR2.....	121
5.7.5. Router P2	122
5.7.6. Router PE2.....	123
5.8. Escenario C con dos routers	124
5.8.1. Router PE1.....	125
5.8.2. Router ASBR1	127
5.8.3. Router ASBR2.....	129
5.8.4. Router PE2.....	130
5.9. Escenario C con tres routers	132
5.9.1. Router PE1.....	132
5.9.2. Router P1	134
5.9.3. Router ASBR1	135
5.9.4. Router ASBR2.....	137
5.9.5. Router P2	138
5.9.6. Router PE2.....	139
6. ACRÓNIMOS	141
7. BIBLIOGRAFIA	142

ÍNDICE DE FIGURAS

Figura 1: Arquitectura básica de una L3VPN sobre MPLS con routers CE, PE y P (Cisco, 2018a).....	16
Figura 2: Topología física de la opción A de VPN de capa 3 entre proveedores (Juniper Networks, 2025).	20
Figura 3: Topología física de la opción B de VPN de capa 3 entre proveedores (Juniper Networks, 2025).	21
Figura 4: Topología física de la opción C de VPN de capa 3 entre proveedores (Juniper Networks, 2025).	23
Figura 5: Configuración de plantilla del router c3600 en GNS3.....	26
Figura 6: Configuración general de la plantilla Linux Core en GNS3.....	28
Figura 7: Opciones avanzadas de la plantilla Linux Core en GNS3.	29
Figura 8: Uso de switches Ethernet en los extremos del escenario de laboratorio.....	30
Figura 9: Topología del escenario A con dos routers de proveedor por lado.....	33
Figura 10: Topología del escenario A con tres routers de proveedor por lado.	35
Figura 11: Topología del escenario B con dos routers de proveedor por lado.....	37
Figura 12: Topología del escenario B con tres routers de proveedor por lado.....	38
Figura 13: Topología del escenario C con dos routers de proveedor por lado.....	41
Figura 14: Topología del escenario C con tres routers de proveedor por lado.....	43
Figura 15: Prueba de conectividad desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario A con dos routers.....	45
Figura 16: Pruebas de conectividad desde LinuxCore1 hacia LinuxCore2, PC2 y PC1 en el escenario A con dos routers.....	45
Figura 17: Pruebas de conectividad desde PC1, PC2 y LinuxCore2 hacia equipos locales y remotos en el escenario A con dos routers.	46
Figura 18: Trazado de ruta desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario A con dos routers.....	47
Figura 19: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario A con dos routers.	48
Figura 20: Tabla de rutas de la VRF EMPRESA en ASBR1 en el escenario A con dos routers.	49
Figura 21: Tabla BGP VPNv4 de la VRF EMPRESA en PE1 en el escenario A con dos routers.	49

Figura 22: Tabla BGP VPNv4 de la VRF EMPRESA en ASBR1 en el escenario A con dos routers.	50
Figura 23: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario A con tres routers.	51
Figura 24: Trazado de ruta desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario A con tres routers.	52
Figura 25: Vecinos OSPF en PE1, P1 y ASBR1 en el escenario A con tres routers.	53
Figura 26: Vecinos LDP en PE1, P1 y ASBR1 en el escenario A con tres routers.	54
Figura 27: Tabla de reenvío MPLS en P1 en el escenario A con tres routers.	55
Figura 28: Tablas de rutas de la VRF EMPRESA en PE1 y ASBR1 en el escenario A con tres routers.	55
Figura 29: Tabla BGP VPNv4 de la VRF EMPRESA en PE1 y ASBR1 en el escenario A con tres routers.	56
Figura 30: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario B con dos routers.	58
Figura 31: Trazado de ruta desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario B con dos routers.	59
Figura 32: Tabla de reenvío MPLS en ASBR1 en el escenario B con dos routers.	59
Figura 33: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario B con dos routers.	60
Figura 34: Tabla BGP VPNv4 en ASBR1 y ASBR2 en el escenario B con dos routers.	61
Figura 35: Rutas VPNv4 recibidas por ASBR1 desde PE1 en el escenario B con dos routers.	62
Figura 36: Rutas VPNv4 anunciadas por ASBR1 hacia ASBR2 en el escenario B con dos routers.	62
Figura 37: Tabla BGP VPNv4 en PE1 y ASBR1 en el escenario B con dos routers.	63
Figura 38: Resumen de vecinos BGP VPNv4 en PE1 y ASBR1 en el escenario B con dos routers.	64
Figura 39: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario B con tres routers.	66
Figura 40: Trazado de ruta desde CE1 hacia la loopback remota de CE2 en el escenario B con tres routers.	66

Figura 41: Vecinos OSPF en P1 en el escenario B con tres routers.....	67
Figura 42: Vecinos LDP en P1 en el escenario B con tres routers.....	67
Figura 43: Tabla de reenvío MPLS en P1 en el escenario B con tres routers.	68
Figura 44: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario B con tres routers.	68
Figura 45: Tabla BGP VPNv4 en ASBR1 y ASBR2 en el escenario B con tres routers.	69
Figura 46: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario C con dos routers.	70
Figura 47: Trazado de ruta desde CE1 hacia la loopback remota de CE2 en el escenario C con dos routers.	71
Figura 48: Rutas IPv4 etiquetadas en ASBR1 mediante BGP en el escenario C con dos routers.	71
Figura 49: Tabla de reenvío MPLS en ASBR1 en el escenario C con dos routers.	72
Figura 50: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario C con dos routers.	73
Figura 51: Tabla BGP VPNv4 en PE1 en el escenario C con dos routers.	73
Figura 52: Resumen de vecinos BGP VPNv4 en PE1 en el escenario C con dos routers.	74
Figura 53: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario C con tres routers.	75
Figura 54: Trazado de ruta desde CE1 hacia la loopback remota de CE2 en el escenario C con tres routers.	76
Figura 55: Vecinos OSPF en P1 en el escenario C con tres routers.....	76
Figura 56: Vecinos LDP en P1 en el escenario C con tres routers.....	77
Figura 57: Tabla de reenvío MPLS en P1 en el escenario C con tres routers.	77
Figura 58: Rutas IPv4 etiquetadas en ASBR1 mediante BGP en el escenario C con tres routers.	78
Figura 59: Tabla de reenvío MPLS en ASBR1 en el escenario C con tres routers.	78
Figura 60: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario C con tres routers.	79
Figura 61: Rutas globales hacia las loopbacks remotas de los PE en el escenario C con tres routers.	80

Figura 62: Resumen de vecinos BGP VPNv4 en PE1 en el escenario C con tres routers.	81
Figura 63: Error de consola Telnet QEMU durante la apertura de un equipo en GNS3.83	
Figura 64: Pérdida inicial de paquetes y posterior conectividad correcta en una prueba de ping.	84
Figura 65: Prueba de ping desde CE1 hacia la loopback remota de CE2 sin utilizar origen específico.....	85
Figura 66: Arranque e inicio de sesión en Linux Core.....	89



1. INTRODUCCIÓN

1.1. Contexto del problema y motivación del trabajo

El problema de fondo que aborda este trabajo es muy habitual en redes de operador. Una empresa puede tener varias sedes repartidas geográficamente y necesitar que todas intercambien rutas como si formasen parte de una única red privada, aunque en realidad estén conectadas a través de distintos proveedores o de distintos sistemas autónomos. La dificultad no está solo en dar conectividad, sino en hacerlo sobre una infraestructura compartida, manteniendo el aislamiento lógico entre clientes y evitando conflictos incluso cuando varios de ellos utilizan direccionamiento privado solapado. Precisamente esa es la base de las BGP/MPLS IP VPN definidas en RFC 4364, donde el proveedor adopta un modelo peer en el que los routers CE anuncian sus rutas a los PE y el backbone del operador se encarga de intercambiarlas sin que el cliente vea una malla overlay entre sedes (Rosen & Rekhter, 2006). Cisco describe el mismo modelo diciendo que una L3VPN sobre MPLS está formada por un conjunto de sedes que se interconectan a través de un núcleo MPLS del proveedor, con routers CE en el borde del cliente y routers PE en el borde del operador (Cisco, 2020a).

Dicho de una forma más práctica, el interés de estas arquitecturas está en que permiten transportar tráfico privado de muchos clientes sobre una red común sin que cada cliente tenga que desplegar su propia infraestructura WAN. MPLS añade un plano de transporte eficiente y BGP aporta el mecanismo de distribución de rutas con políticas. Cuando ambos elementos se combinan con VRF, Route Distinguisher y Route Target, el operador puede mantener separadas las tablas de encaminamiento de cada cliente y, al mismo tiempo, ofrecer conectividad extremo a extremo entre sedes remotas (Cisco, 2020a). Desde el punto de vista de este TFG, esa idea general se vuelve todavía más interesante cuando la VPN debe extenderse entre varios sistemas autónomos, porque ahí entran en juego distintas opciones de interconexión con compromisos diferentes entre simplicidad, escalabilidad y complejidad operativa. RFC 4364 dedica precisamente su sección 10 a los multi-AS backbones y presenta varias soluciones ordenadas de menor a mayor escalabilidad (Rosen & Rekhter, 2006).

1.2. OSPF como protocolo interno de encaminamiento en el núcleo

Aunque buena parte del interés del trabajo gira alrededor de BGP y de las L3VPN, el papel de OSPF dentro del núcleo del proveedor sigue siendo esencial. RFC 2328 define OSPFv2 como un protocolo de estado de enlace diseñado para funcionar dentro de un único sistema autónomo. Cada router mantiene una base de datos idéntica de la topología del AS y, a partir de ella, calcula un árbol de caminos mínimos. El propio estándar insiste en que OSPF fue concebido para recalcular rutas con rapidez ante cambios topológicos y para hacerlo con un uso contenido del tráfico de control (Moy, 1998). Cisco resume esa misma idea destacando su rápida convergencia, su buena escalabilidad y su empleo extendido como IGP en redes IP (Cisco, 2022a).

En el contexto de una red MPLS de operador, OSPF no se utiliza para transportar directamente las rutas de los clientes, sino para dar coherencia al underlay del proveedor. Es decir, se encarga de que los routers del core conozcan las loopbacks y los prefijos internos necesarios para que luego MPLS y BGP puedan construir servicios sobre una base estable (Cisco, 2022b). Esta separación entre encaminamiento interno del proveedor y distribución de rutas VPN es una idea importante, porque ayuda a entender por qué en topologías mínimas la necesidad de un IGP puede parecer menos evidente, mientras que en núcleos más realistas, con varios routers intermedios, un protocolo interno como OSPF pasa a ser prácticamente imprescindible. Por eso tiene sentido que en este trabajo la comparación entre topologías simples y topologías con más routers de core no sea solo una cuestión estética, sino una manera de ver cuándo el papel del IGP se vuelve realmente necesario.

1.3. MPLS y el transporte de servicios VPN de capa 3

La segunda pieza fundamental del trabajo es MPLS. RFC 3031 define la arquitectura de Multiprotocol Label Switching y establece el marco general de un sistema de reenvío basado en etiquetas (Rosen et al., 2001). En la documentación de Cisco aparece descrito de una manera muy similar, como una tecnología que permite a empresas y proveedores ofrecer múltiples servicios avanzados sobre una única infraestructura. Lo importante aquí no es quedarse solo con la idea clásica de “conmutación rápida”, sino entender que MPLS terminó convirtiéndose en la base de servicios como las VPN de capa

3, las VPN de capa 2 o la ingeniería de tráfico (Cisco, s. f.). En otras palabras, MPLS dejó de ser solamente una mejora de reenvío y pasó a ser una plataforma de transporte para servicios de operador.

Cuando esa arquitectura se aplica a una L3VPN, la lógica de funcionamiento se aclara bastante. Cisco explica que el router PE recibe rutas del CE, las transforma en rutas VPNv4, les añade etiquetas adecuadas y las intercambia con otros PE mediante MP-BGP. A partir de ahí, el core transporta el tráfico usando MPLS, de modo que las sedes del cliente se comunican de forma privada sobre una red compartida. La misma guía de Cisco insiste en que, para que esto funcione, todos los dispositivos del core deben soportar encaminamiento MPLS y que, dentro de la red del proveedor, la distribución de etiquetas suele apoyarse en LDP (Cisco, 2022b). Esto encaja bien con la idea que vertebra el trabajo, porque permite separar tres planos que conviene no mezclar: el plano interno del proveedor, resuelto con IGP y MPLS, el plano de servicio del cliente, modelado con VRF, y el plano de control de la VPN, resuelto con BGP multiprotocolo.

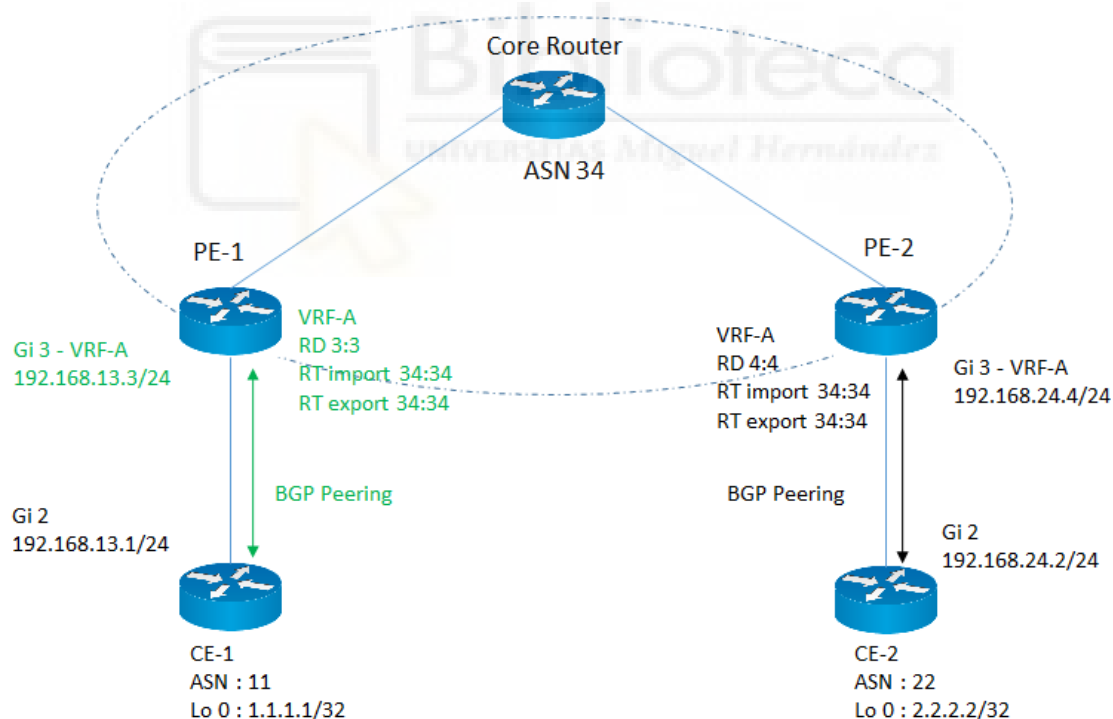


Figura 1: Arquitectura básica de una L3VPN sobre MPLS con routers CE, PE y P (Cisco, 2018a).

En la Figura 1 se observa una arquitectura general de una L3VPN sobre MPLS en la que aparecen dos routers CE situados en los extremos, dos routers PE en el borde del proveedor y un router de core en el interior del backbone. Esta representación es útil

porque permite identificar de forma muy visual cómo se reparte la función de cada equipo dentro del servicio. Los routers CE representan las redes del cliente, los routers PE actúan como punto de entrada del proveedor al encaminamiento de esas sedes y el router del núcleo se encarga del transporte dentro de la red MPLS. Además, la figura introduce de forma clara la existencia de una VRF en los PE y muestra también los parámetros RD y RT, lo que ayuda a relacionar desde el principio la arquitectura física con la separación lógica de rutas propia de una L3VPN.

La Figura 1 también permite entender que las sedes del cliente no se conectan directamente entre sí, sino que lo hacen a través del backbone del operador. Desde el punto de vista del servicio, esto resulta importante porque muestra que el proveedor ofrece conectividad privada sobre una infraestructura compartida, manteniendo separadas las tablas de encaminamiento de cada cliente mediante VRF (Cisco, 2018a). Por eso esta figura encaja bien como punto de partida antes de entrar en más detalle en conceptos como VPNv4, MP-BGP o las alternativas de interconexión entre sistemas autónomos, ya que fija primero la estructura básica sobre la que se apoya todo lo demás.

1.4. BGP y extensiones multiprotocolo en redes L3VPN

BGP es el protocolo que se utiliza para intercambiar rutas entre sistemas autónomos. A diferencia de OSPF, que trabaja dentro de un mismo sistema autónomo, BGP está pensado para establecer relaciones de encaminamiento entre dominios administrativos distintos (Rekhter et al., 2006). Por eso, dentro del contexto de este trabajo, BGP resulta fundamental para entender cómo se intercambia la información entre redes diferentes y cómo se construyen servicios VPN que deben atravesar más de un sistema autónomo.

Cuando se trabaja con VPN de capa 3, BGP no se limita únicamente al intercambio clásico de rutas IPv4 unicast. Para poder transportar otro tipo de información, como rutas VPNv4, es necesario utilizar sus extensiones multiprotocolo. Gracias a ellas, BGP puede anunciar distintas familias de direcciones y no solo prefijos IPv4 tradicionales (Bates et al., 2007). Este detalle es especialmente importante en este TFG, porque varias de las diferencias entre unas soluciones inter-AS y otras no dependen tanto del uso de BGP en sí, sino del tipo de información que se intercambia en cada caso.

En este sentido, interesan especialmente tres familias de direcciones. La primera es IPv4 unicast, que representa el uso más básico y habitual de BGP. La segunda es VPNv4, utilizada en entornos MPLS L3VPN para transportar prefijos de clientes de forma diferenciada mediante el uso de RD y etiquetas MPLS (Rosen & Rekhter, 2006). La tercera es IPv4 labeled-unicast, que permite anunciar prefijos IPv4 junto con su etiqueta correspondiente y resulta especialmente útil en las soluciones inter-AS más escalables. (Rekhter & Rosen, 2001)

Visto de forma general, la lógica es la siguiente. Un router PE aprende una ruta desde un CE en formato IPv4 normal. Si esa ruta pertenece a una VPN, el router la adapta al contexto del servicio y la distribuye mediante las extensiones adecuadas de BGP. A partir de ahí, la forma en la que esa información se propaga entre sistemas autónomos depende de la alternativa de interconexión utilizada. Por eso, comprender el papel de BGP y de sus distintas familias de direcciones resulta imprescindible para interpretar correctamente los escenarios que se analizan más adelante.

1.5. VRF, RD y RT en redes L3VPN

Si MPLS aporta el transporte y BGP se encarga de distribuir la información de encaminamiento, la pieza que realmente permite separar unos clientes de otros dentro de la red del proveedor es la VRF. En una L3VPN, cada cliente, o cada conjunto de sedes que deben compartir rutas, dispone de una tabla propia en el router PE. Esto permite que el proveedor transporte tráfico de varios clientes sobre una misma infraestructura sin mezclar sus rutas, incluso cuando utilizan direcciones privadas iguales o solapadas. Cisco describe esta idea indicando que cada sitio del cliente se asocia a una VRF (Cisco, 2022b), mientras que RFC 4364 basa toda la arquitectura de las BGP/MPLS IP VPN en tablas separadas de encaminamiento y reenvío dentro de los routers PE (Rosen & Rekhter, 2006).

A partir de ahí aparecen dos conceptos que conviene diferenciar bien. El primero es el Route Distinguisher o RD. Su función no es decidir qué sitios pertenecen a una VPN, sino convertir un prefijo IPv4 normal en una ruta única dentro del espacio VPNv4. Dicho de otra manera, el RD sirve para distinguir rutas que podrían coincidir en direccionamiento, algo muy habitual cuando distintos clientes usan redes privadas

parecidas (Rosen & Rekhter, 2006). Cisco lo resume de forma muy clara cuando explica que el RD añade un valor de 8 bytes al prefijo IPv4 para crear un prefijo VPN IPv4 único (Cisco, 2020a), y RFC 4364 desarrolla esa misma lógica al definir la estructura de las rutas VPNv4.

El segundo concepto es el Route Target o RT. Aquí la función es distinta, porque mientras que el RD hace única la ruta, el RT decide en qué VRF puede importarse o exportarse esa información. En la práctica, esto significa que el RT es el mecanismo que permite construir la pertenencia lógica a una VPN. RFC 4364 explica esta idea mediante el uso de comunidades extendidas (Rosen & Rekhter, 2006), y Cisco lo presenta indicando que cada VRF dispone de políticas de importación y exportación asociadas a los Route Target (Cisco, 2020a). Una forma sencilla de recordarlo sería decir que el RD diferencia rutas y el RT controla quién puede recibirlas dentro del servicio.

Visto de forma conjunta, puede decirse que VRF, RD y RT forman el núcleo lógico de una MPLS L3VPN. La VRF separa tablas de encaminamiento, el RD evita colisiones entre prefijos y el RT controla cómo se comparten las rutas dentro del servicio. Cuando estas tres piezas se combinan con MPLS y MP-BGP, el proveedor puede ofrecer conectividad privada entre sedes remotas sobre una red común, manteniendo al mismo tiempo el aislamiento entre clientes. Esta idea resulta especialmente importante en este trabajo, porque después ayuda a entender por qué unas soluciones inter-AS son más simples, otras más escalables y otras reparten la complejidad de forma diferente entre los routers frontera y los routers PE.

1.6. Alternativas de interconexión entre sistemas autónomos

Conviene precisar que, cuando en este apartado se habla de Opción A, Opción B y Opción C, se hace referencia a las alternativas estándar de interconexión inter-AS descritas en la literatura técnica y en la documentación de fabricantes. No se trata todavía de los escenarios concretos implementados en este TFG, aunque lógicamente existe relación entre ambas partes. La adaptación práctica al entorno de laboratorio se desarrollará más adelante, en el apartado de material y métodos.

1.6.1. Opción A

El capítulo 10 de la RFC 4364 presenta las soluciones multi-AS en orden de escalabilidad creciente, y ese detalle no es menor, porque encaja exactamente con la lógica comparativa del trabajo. La primera alternativa consiste en enlazar VRF con VRF en los routers frontera de AS (Rosen & Rekhter, 2006). Es la solución que Juniper y Cisco denominan Opción A, Juniper la define como una conexión interproveedor VRF-to-VRF en los ASBR y la describe como la solución más simple y a la vez la menos escalable (Juniper Networks, 2025). Cisco coincide con esa idea y añade un detalle muy útil para este TFG, que entre ASBR el tráfico circula como paquetes IP sin etiqueta sobre una interfaz asociada a VRF. Esta opción tiene una gran ventaja didáctica, porque deja muy visible la lógica de la VPN, pero resulta costosa cuando el número de VPNs crece, ya que obliga a mantener una VRF por VPN en los routers frontera y, normalmente, una sesión de routing por cada una de ellas (Cisco, 2024a).

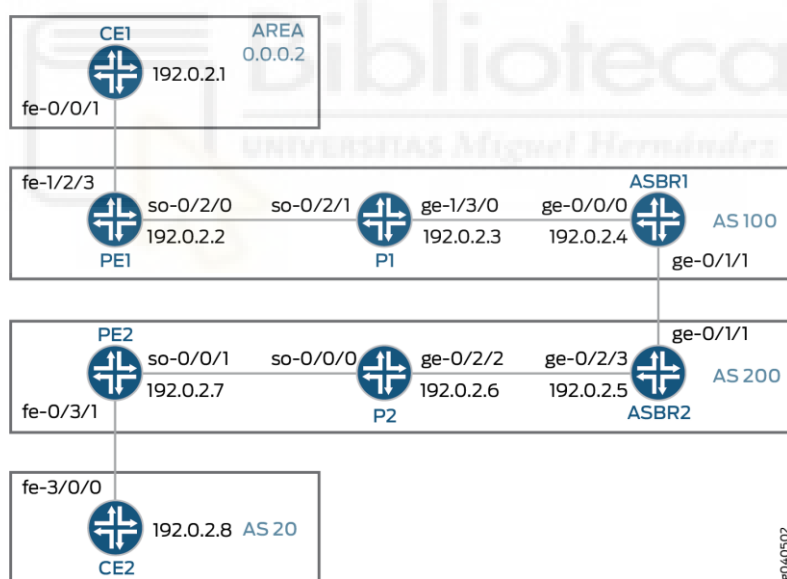


Figura 2: Topología física de la opción A de VPN de capa 3 entre proveedores (Juniper Networks, 2025).

En la Figura 2 se representa la topología física asociada a la Opción A. En ella se observa que cada proveedor mantiene su propia red y que la interconexión entre sistemas autónomos se realiza a través de los routers frontera ASBR. Esta imagen resulta útil porque permite visualizar la idea básica de esta alternativa, que consiste en prolongar la lógica del servicio VPN hasta el borde entre proveedores mediante una relación VRF a

VRF. Aunque la topología física no parece especialmente compleja, lo importante aquí no es tanto el número de equipos como el hecho de que cada VPN debe mantenerse de forma explícita en la frontera, lo que convierte a esta opción en una solución sencilla desde el punto de vista conceptual, pero limitada cuando el número de servicios crece. Vista así, la figura ayuda a entender por qué la Opción A suele considerarse la alternativa más fácil de interpretar y, al mismo tiempo, la menos escalable.

1.6.2. Opción B

La segunda alternativa es la Opción B. En este caso, los ASBR ya no intercambian rutas IPv4 dentro de una VRF enlazada punto a punto, sino que redistribuyen rutas VPNv4 etiquetadas entre sistemas autónomos (Rosen & Rekhter, 2006). Cisco lo resume diciendo que Opción B proporciona distribución de rutas VPNv4 entre ASBR y que estos routers reciben tráfico MPLS en la interfaz inter-AS (Cisco, 2016). Juniper añade que los ASBR mantienen las rutas VPN-IPv4 en la base de información de encaminamiento y las etiquetas asociadas en la base de reenvío, y la califica como una solución algo escalable, más razonable que la Opción A, aunque todavía por debajo de la siguiente (Juniper Networks, 2025). Desde el punto de vista conceptual, Opción B representa un paso claro en madurez técnica, porque evita replicar VRF por VPN en la frontera, pero a cambio obliga a que los ASBR entiendan y procesen información VPNv4.

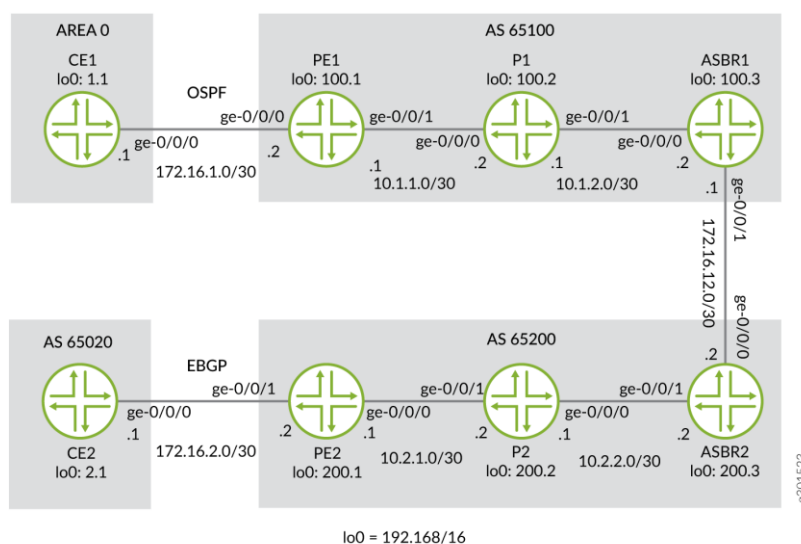


Figura 3: Topología física de la opción B de VPN de capa 3 entre proveedores (Juniper Networks, 2025).

En la Figura 3 se muestra la topología correspondiente a la Opción B. A diferencia del caso anterior, aquí la atención ya no se centra en una unión VRF a VRF entre ASBR, sino en el hecho de que los routers frontera pasan a intercambiar rutas VPNv4 etiquetadas entre sistemas autónomos. La imagen sigue mostrando una separación clara entre los distintos AS, con routers CE en los extremos, PE en el borde del proveedor y ASBR en la frontera, pero el cambio importante no está tanto en la disposición física como en la información que intercambian los equipos de borde entre proveedores. Por eso esta figura debe entenderse como una evolución de la opción anterior. Se reduce la necesidad de replicar la lógica de cada VPN directamente en la frontera, pero a cambio los ASBR deben ser capaces de mantener y procesar información VPNv4. De este modo, la imagen sirve para reforzar la idea de que la Opción B representa un punto intermedio entre simplicidad y escalabilidad.

1.6.3. Opción C

La tercera alternativa, Opción C, es la más escalable de las tres en la documentación consultada. RFC 4364 la describe como una redistribución multihop EBGp de rutas VPN-IPv4 etiquetadas entre AS de origen y destino, combinada con redistribución EBGp de rutas IPv4 etiquetadas entre AS vecinos (Rosen & Rekhter, 2006). El punto clave es que, en esta solución, las rutas VPNv4 no se mantienen ni se distribuyen en los ASBR, sino que estos se limitan a anunciar rutas internas etiquetadas, típicamente /32 de loopback de PE, de manera que se construye un camino MPLS de extremo a extremo entre routers PE situados en AS distintos (Cisco, 2018b). Juniper la presenta como la alternativa más escalable (Juniper Networks, 2025) y Cisco la explica en su documentación técnica apoyándose precisamente en ese uso de sesiones multihop y en la necesidad de verificar después el intercambio de etiquetas y la conectividad con ping y traceroute (Cisco, 2018b). En otras palabras, se trata de una solución más exigente desde el punto de vista del plano de control, pero que reduce la carga funcional de los ASBR y se adapta mejor a entornos interproveedor de mayor tamaño.

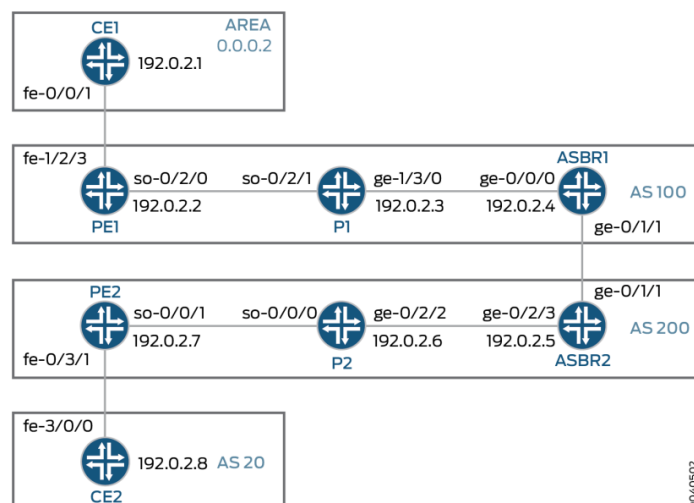


Figura 4: Topología física de la opción C de VPN de capa 3 entre proveedores (Juniper Networks, 2025).

En la Figura 4 se observa la topología física asociada a la Opción C. A primera vista puede parecer muy similar a la de la Opción A, y precisamente por eso conviene explicarla bien, ya que en esta alternativa la diferencia principal no está en el dibujo físico, sino en la lógica de control que se construye sobre él. En este caso, los ASBR dejan de encargarse de mantener las rutas VPNv4 y se limitan a intercambiar rutas internas etiquetadas que permiten construir el transporte MPLS entre sistemas autónomos, mientras que la información VPN queda fuera de la frontera directa. Esto hace que la red sea más exigente desde el punto de vista del plano de control, pero también más adecuada para entornos interproveedor de mayor tamaño. Por tanto, la Figura 4 no solo ilustra una topología, sino que ayuda a ver que la Opción C conserva una estructura física parecida a otras alternativas, aunque redistribuye la complejidad de una forma más eficiente y escalable.

1.7. Entorno de emulación y objetivos del trabajo

Aunque el núcleo teórico del trabajo se apoya en RFC y documentación de fabricante, la validación práctica se realiza en GNS3, y por eso merece la pena justificar aquí su elección. La documentación oficial de GNS3 lo describe como una herramienta de uso extendido para emular, configurar, probar y depurar redes virtuales y reales, capaz de ejecutar desde topologías pequeñas en un portátil hasta laboratorios de mayor tamaño distribuidos en varios servidores. También destaca que es software libre y que

actualmente soporta dispositivos de distintos fabricantes y distintos tipos de appliances. Este carácter abierto y multiproveedor hace que resulte especialmente apropiado para un TFG orientado a redes de operador, donde interesa combinar routers, nodos Linux y equipos de borde en un mismo laboratorio (GNS3, s. f.-a).

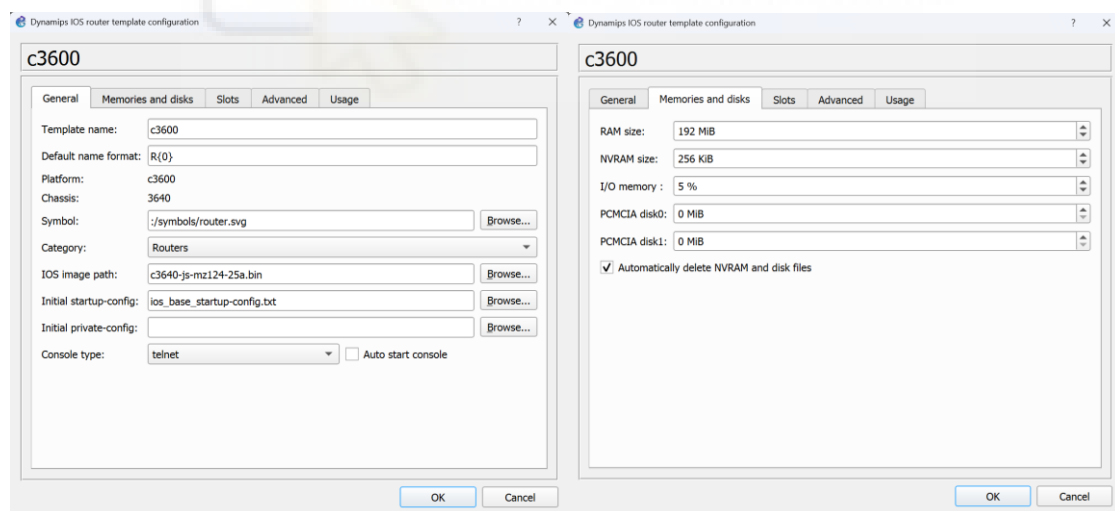
Ese valor práctico no aparece solo en la documentación oficial del proyecto, sino que también se observa en trabajos universitarios previos. Un TFG de la Universitat Politècnica de València sobre un entorno MPLS basado en GNS3 lo emplea precisamente para diseñar y configurar redes MPLS, estudiar BGP y construir VPN de nivel 3 (Rodríguez, 2019), mientras que otro trabajo de la Universidad de Valladolid concluye que GNS3 permite reproducir prácticas de diseño y configuración de redes y lo considera especialmente útil por la variedad de escenarios que puede simular. No son fuentes normativas ni sustituyen a los RFC, pero sí sirven para reforzar una idea razonable dentro del contexto académico: GNS3 no es una elección improvisada, sino una herramienta que ya ha sido utilizada en entornos docentes e incluso en trabajos fin de grado para validar tecnologías de encaminamiento y servicios VPN (González Ramírez, 2022), además de en determinadas asignaturas a lo largo de la carrera.

A partir de todo lo anterior, los objetivos de esta investigación quedan bastante bien definidos. En primer lugar, se pretende estudiar con una base técnica sólida las tecnologías que hacen posible una L3VPN entre sedes remotas, prestando especial atención a OSPF, MPLS, BGP, MP-BGP, VRF y a las familias de direcciones IPv4 unicast, VPNv4 e IPv4 labeled-unicast. En segundo lugar, se busca comparar las principales alternativas de interconexión entre sistemas autónomos, analizando qué cambia en cada una desde el punto de vista del intercambio de rutas, del transporte del tráfico y de la escalabilidad. En tercer lugar, el trabajo quiere validar esos conceptos en laboratorio, primero en topologías simples y después en topologías con un núcleo más representativo, para observar cómo evoluciona el papel de OSPF y MPLS cuando el número de routers del core deja de ser trivial. Finalmente, se persigue que la comparación entre los escenarios no se limite a “funciona o no funciona”, sino que permita argumentar por qué una solución es más sencilla, otra más flexible y otra más escalable. Ese enfoque coincide con la estructura de trabajo acordada en tutoría, donde la introducción se reserva para el estado de la técnica y la parte experimental se traslada a Material y métodos y a Resultados y discusión.

2. MATERIAL Y MÉTODOS

2.1. Router c3600

Antes de describir los escenarios A, B y C, considero necesario presentar el router base utilizado en el laboratorio, ya que toda la implementación práctica parte de una misma plantilla de emulación dentro de GNS3. En este trabajo se empleó la serie Cisco 3600, concretamente una plantilla c3600 con chasis 3640, cargada mediante Dynamips y asociada a una imagen IOS 12.4. La elección de esta plataforma respondió sobre todo a criterios de disponibilidad, estabilidad y coherencia con un entorno académico de prácticas, donde resulta más importante comprender el comportamiento de las tecnologías que trabajar con appliances modernos de mayor consumo. La documentación oficial de GNS3 indica precisamente que las imágenes c3640 y c3660 son de las más recomendables dentro de Dynamips, siempre que se utilicen con la memoria adecuada y con un valor Idle-PC correcto, lo que refuerza su idoneidad para laboratorios de este tipo (GNS3, s. f.-b). Además, cabe añadir que se trata de un router que ya había utilizado previamente durante determinadas asignaturas de la carrera.



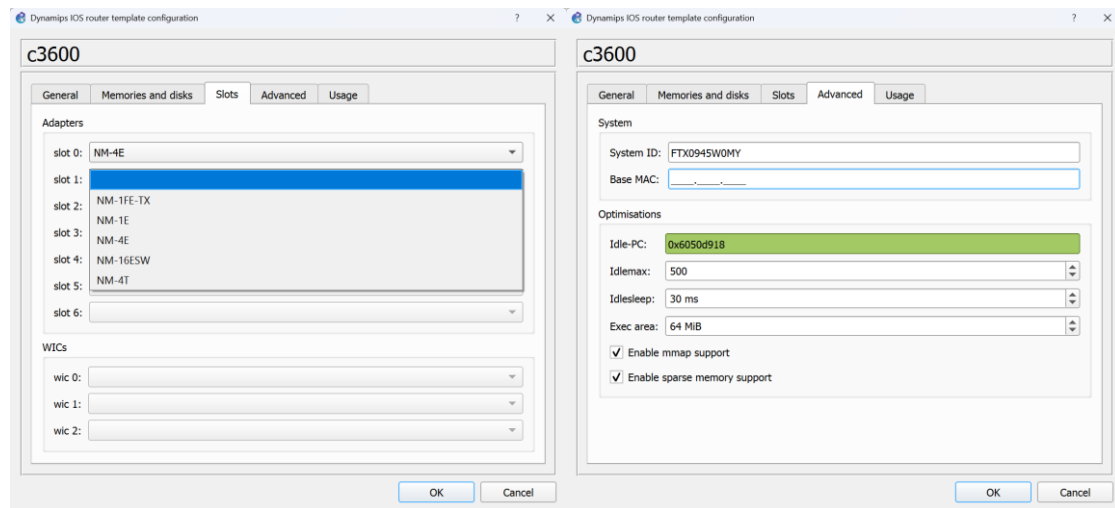


Figura 5: Configuración de plantilla del router c3600 en GNS3.

En la Figura 5 observamos la configuración utilizada para la plantilla c3600 dentro del entorno GNS3. En ella se aprecian la pestaña general, la configuración de memoria, la selección de módulos de interfaz y el apartado avanzado del emulador. A partir de esta plantilla se definió un entorno homogéneo para todos los routers del laboratorio, usando consola telnet, una memoria RAM de 192 MiB, una NVRAM de 256 KiB, el módulo NM-4E en el slot 0 y un valor Idle-PC fijado manualmente para reducir la carga innecesaria del procesador durante la emulación. En este caso, se optó por un módulo Ethernet para mantener una configuración sencilla y homogénea en todos los escenarios, aunque la propia plantilla permitía seleccionar otras alternativas de interfaz, como módulos Ethernet, FastEthernet o GigaEthernet, en función de las necesidades del montaje. Esta imagen es especialmente útil porque documenta de forma visual que el router no se añadió como un dispositivo genérico, sino como una plantilla previamente ajustada para que todos los escenarios partieran de una base común, estable y repetible.

Metodológicamente, utilizar siempre la misma familia de routers permitió mantener unas condiciones de partida homogéneas y centrar la comparación en el diseño de los escenarios, y no en cambios de plataforma. Dicho de otra forma, al mantener la misma plantilla c3600 para los nodos CE, PE, P y ASBR, la comparación entre opciones se centró en la lógica de diseño y no en cambios de plataforma. Además, la serie Cisco 3600 fue concebida como una familia modular, con capacidad para incorporar distintos tipos de interfaces y adaptarse a diferentes funciones dentro de la red, algo que encaja bien con el planteamiento del laboratorio. En la documentación histórica de Cisco se

indica que el modelo 3640 dispone de cuatro ranuras para módulos de red y que la serie 3600 fue diseñada como una solución modular para servicios de routing y conectividad WAN y LAN (Cisco, 2001), por lo que su uso en este trabajo no solo fue práctico, sino también coherente con la variedad de papeles que cada router debía asumir en la topología.

2.2. PCs

2.2.1. VPCS

Dentro del laboratorio, los nodos VPCS se utilizaron como estaciones finales ligeras para representar equipos de usuario conectados a los routers CE. Su función principal no fue añadir complejidad al escenario, sino ofrecer un extremo de red sencillo desde el que poder definir direccionamiento IP, puerta de enlace y conectividad básica dentro de cada dominio. Esta elección encaja bien con el planteamiento del trabajo, ya que en el apartado de Material y métodos interesa describir cómo se construyó el entorno de pruebas y no añadir sistemas finales más pesados de lo necesario. En ese sentido, VPCS resultó especialmente útil porque permitió disponer de hosts simples, rápidos de desplegar y suficientemente funcionales para los objetivos del laboratorio.

Desde el punto de vista técnico, la propia documentación oficial de GNS3 define VPCS como un simulador de PC ligero que soporta funciones básicas como DHCP y ping, y destaca además que consume únicamente 2 MB de RAM por instancia y que no requiere cargar una imagen adicional. Esto lo convierte en una solución muy adecuada para topologías docentes o de validación, donde interesa reservar la mayor parte de los recursos del equipo anfitrión para los routers y no para los extremos de usuario. En la práctica, su integración dentro de GNS3 también resulta cómoda porque aparece por defecto en la categoría de dispositivos finales, de manera que puede añadirse directamente al proyecto sin configuración previa compleja (GNS3, s. f.-c).

La utilización de VPCS fue útil porque permitió representar de forma sencilla los extremos de usuario sin incorporar herramientas o servicios adicionales que no aportaban valor al objetivo del laboratorio. Al ser un nodo muy ligero, su papel quedó reducido a

actuar como host final dentro de la subred correspondiente. Gracias a ello, la construcción del escenario se mantuvo clara y el interés del trabajo siguió recayendo sobre la red y no sobre el sistema operativo del equipo final.

2.2.2. Linux Core

Frente a los nodos VPCS, en este trabajo también se utilizaron instancias Linux Core como equipos finales más completos dentro del laboratorio. Su incorporación permitió representar un host realista sin abandonar el entorno de emulación de GNS3, lo que resultó útil para construir los extremos de red de una forma más cercana a un equipo de propósito general. En la práctica, estos nodos Linux se conectaron a los switches Ethernet no configurables situados en ambos extremos de la topología, compartiendo red local con los VPCS y actuando como uno de los equipos finales asociados a cada CE. De este modo, el laboratorio combinó un terminal muy ligero, como VPCS, con un host más flexible basado en Linux, lo que enriqueció la construcción del escenario sin desviar la atención del comportamiento de la red. Además, en la documentación interna de los escenarios ya se recoge expresamente que uno de los PCs conectados a cada switch extremo sería un Linux Core, con direccionamiento dentro de la red 172.17.1.0/24 en un lado y 172.17.2.0/24 en el otro, lo que confirma que su uso formó parte del diseño previsto desde el principio.

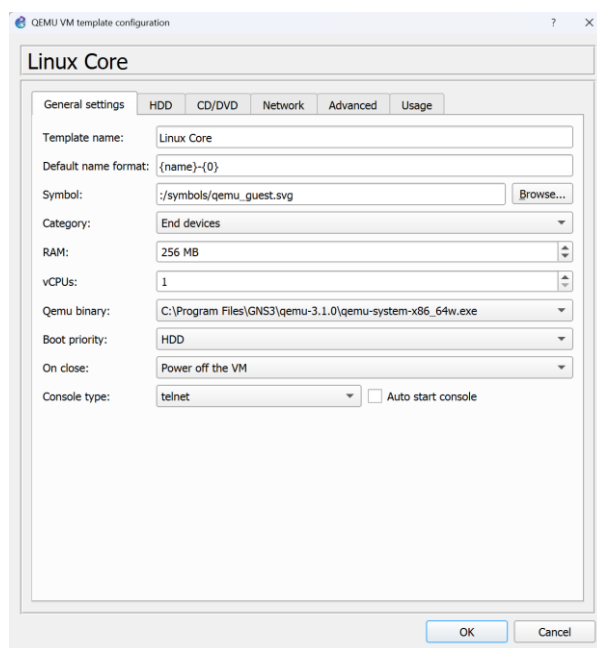


Figura 6: Configuración general de la plantilla Linux Core en GNS3.

Observamos en la Figura 6 la configuración general de la plantilla Linux Core dentro de GNS3. En ella se aprecia que el nodo se definió como una máquina QEMU, con 256 MB de RAM, una vCPU, arranque desde disco duro, consola telnet y apagado directo de la máquina virtual al cerrar el nodo. Esta configuración se eligió con el objetivo de mantener un host sencillo, suficientemente ligero para el trabajo y fácil de integrar en los distintos escenarios. La captura también muestra que el binario de QEMU se asoció directamente desde la instalación local de GNS3, lo que refuerza la idea de que no se trabajó con un contenedor improvisado, sino con una plantilla configurada de forma estable para poder reutilizarla varias veces dentro del proyecto.

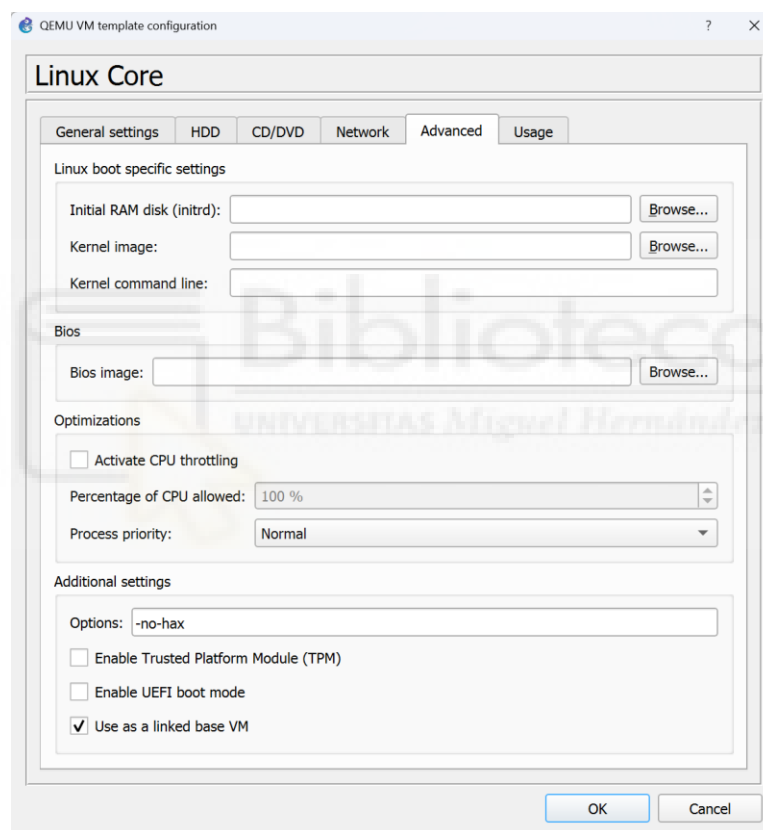


Figura 7: Opciones avanzadas de la plantilla Linux Core en GNS3.

En la Figura 7 podemos visualizar uno de los puntos que más costó durante la instalación de Linux Core, ya que fue necesario ajustar manualmente varias opciones avanzadas hasta conseguir que la plantilla funcionara de forma correcta. En concreto, se activó la opción de usar la máquina como linked base VM y se añadió el parámetro no-hax dentro del apartado de opciones adicionales. Este ajuste no fue un detalle menor, porque precisamente una parte de la dificultad de la instalación estuvo en conseguir que

el nodo Linux arrancara de forma estable dentro de QEMU sin generar problemas adicionales en el entorno de emulación. Desde un punto de vista metodológico, esta experiencia también tiene valor dentro del apartado de Material y métodos, porque muestra que la construcción del trabajo no consistió solo en arrastrar dispositivos al escenario, sino en ajustar previamente cada plantilla para asegurar que el entorno final fuese reutilizable, homogéneo y operativo en todos los montajes posteriores.

2.3. Ethernet Switch

En el laboratorio también se emplearon switches Ethernet no programables integrados en GNS3 para interconectar los equipos finales con los routers CE en cada extremo de la topología. Su función dentro del montaje fue sencilla pero necesaria, ya que permitieron agrupar en una misma red local a los nodos VPCS, a los Linux Core y al router de cliente correspondiente, reproduciendo así un pequeño segmento LAN de acceso sin añadir complejidad de configuración adicional. Esta elección resultó adecuada para el trabajo porque el objetivo no era estudiar el funcionamiento interno de un switch gestionable, sino disponer de un elemento de capa 2 que facilitara la conexión física de los hosts finales dentro de cada escenario. La propia documentación de GNS3 explica que su switch Ethernet integrado es un dispositivo simulado y no un equipo que ejecute un sistema operativo real, lo que encaja perfectamente con el uso que se le dio en este TFG, donde interesaba mantener el foco en los routers y en la interconexión entre sistemas autónomos (GNS3, s. f.-a).

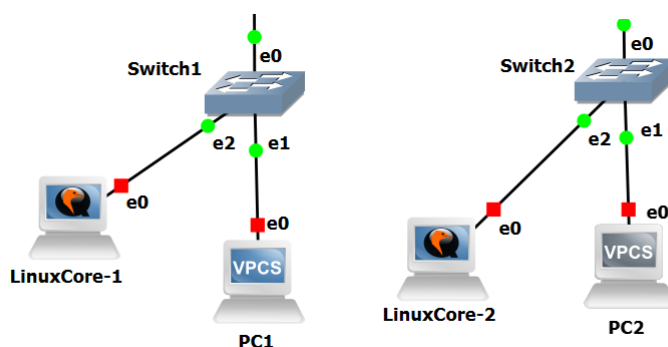


Figura 8: Uso de switches Ethernet en los extremos del escenario de laboratorio.

Podemos observar en la Figura 8 que en cada uno de los extremos del escenario aparece un switch conectado directamente al router CE y, a su vez, a dos equipos finales, uno de tipo Linux Core y otro de tipo VPCS. Esta disposición se repitió de forma equivalente en los distintos montajes porque permitía mantener una estructura homogénea en ambos lados de la red y facilitaba la construcción del laboratorio sin necesidad de introducir configuraciones específicas de conmutación. Desde un punto de vista metodológico, el uso de estos switches ayudó a separar claramente el acceso local del resto de la arquitectura, dejando a los routers la responsabilidad del encaminamiento y a los dispositivos de capa 2 la mera función de interconexión dentro de la LAN. De esta forma, el laboratorio se mantuvo ordenado, repetible y coherente con el propósito del trabajo, que no era analizar switching avanzado, sino validar la construcción de los escenarios inter-AS.

2.4. Escenario A

Una vez descritos los equipos empleados en el trabajo, este apartado se centra en la construcción del escenario A, cuya base teórica ya fue explicada en el apartado 1.6. En términos generales, esta alternativa responde al planteamiento clásico de interconexión VRF a VRF entre routers frontera de distintos sistemas autónomos, de manera que la VPN se prolonga hasta el borde entre proveedores sin intercambiar VPNv4 en la frontera inter-AS. Tanto la documentación de Cisco como la de Juniper coinciden en presentar esta opción como la más sencilla desde el punto de vista conceptual, aunque también como la menos escalable cuando el número de VPN crece, precisamente porque obliga a mantener información específica del servicio en los equipos de borde (Cisco, 2022c), (Juniper Networks, 2025).

En este trabajo, el escenario A se desarrolló en dos versiones distintas. La primera corresponde a una topología más reducida, con dos routers de proveedor por lado, pensada para representar la lógica mínima del mecanismo. La segunda amplía esa misma idea añadiendo un router interno en cada sistema autónomo de proveedor, lo que permite construir una red más cercana a un entorno de operador sin alterar el principio básico del escenario A. En ambas variantes se mantuvo la misma distribución general de extremos, con un router CE en cada cliente, un switch Ethernet no configurable en cada LAN y dos

hosts finales por lado, uno basado en Linux Core y otro en VPCS, de forma que la comparación entre versiones se apoyara sobre una estructura homogénea.

Esta doble aproximación no se eligió al azar. La versión reducida permite aislar con claridad la lógica de la interconexión VRF a VRF entre ASBR, mientras que la versión ampliada introduce ya un núcleo interno de proveedor en el que aparecen mecanismos de alcance interno y transporte que no son necesarios en una topología demasiado simple. En consecuencia, se optó por no seguir incrementando el número de routers en este escenario, puesto que con la aparición de un router de tránsito ya queda representada una topología interna más completa que la versión inicial. En otras palabras, esta ampliación ya permite mostrar el cambio principal del escenario sin necesidad de repetir el mismo esquema con más nodos.

El escenario A se ha desarrollado utilizando una única VRF con el objetivo de simplificar la representación del mecanismo y centrar la atención en su funcionamiento básico. No obstante, esta misma lógica sería extensible a un número mayor de VRF, siempre que se replique la configuración correspondiente en los routers PE y en los ASBR. Para ello, sería necesario definir cada nueva VRF en los equipos de borde del proveedor y asociarla después a las interfaces que conectan los PE con los CE, donde habitualmente se emplea una interfaz física por cada VRF. Del mismo modo, también habría que aplicar esas VRF en la interconexión entre ASBR, donde lo más habitual es utilizar una única interfaz física y crear sobre ella una subinterfaz lógica para cada servicio. De esta forma, aunque en este escenario se haya trabajado con una sola VRF por claridad expositiva, el planteamiento utilizado mantiene validez conceptual para entornos con varios clientes o varios servicios diferenciados. Cabe recalcar que el resto de los escenarios también se definirá una única VRF en los PE para simplificar las configuraciones.

Conviene señalar que en este apartado no se pretende entrar en el detalle exhaustivo de cada línea de configuración, ya que eso haría la lectura demasiado pesada y desviaría la atención del verdadero objetivo de esta sección, que es explicar cómo se ha construido cada escenario y qué lógica sigue su diseño. Por este motivo, la configuración completa de los routers, junto con los fragmentos más específicos del montaje, se incluirá en el anexo correspondiente, donde podrá consultarse de forma ordenada y sin interrumpir

el hilo principal de la memoria. Del mismo modo, las comprobaciones de funcionamiento y la validación de que cada escenario opera correctamente no se desarrollan aquí, sino en el apartado de Resultados y discusión, que es donde corresponde analizar las pruebas realizadas, interpretar su significado y comparar el comportamiento de las distintas alternativas implementadas.

2.4.1. Escenario A con dos routers

La versión más simple del escenario A se construyó con una topología formada, en cada lado del proveedor, por un router PE y un router ASBR, además de los routers CE situados en los extremos del cliente. En esta implementación, la interconexión entre sistemas autónomos se resolvió sin MPLS entre AS y sin intercambio de rutas VPNv4 en la frontera, apoyándose en sesiones IPv4 dentro de la VRF del servicio. Dicho de otra manera, el escenario se planteó como una cadena de encaminamiento por VRF entre CE1, PE1, ASBR1, ASBR2, PE2 y CE2, manteniendo la VPN ligada al contexto de la tabla VRF y no a un intercambio multiprotocolo entre fronteras. Esta forma de trabajo encaja con la descripción clásica del escenario A como una solución back-to-back VRF o VRF-to-VRF, y también coincide con la documentación de configuración utilizada para este montaje concreto.

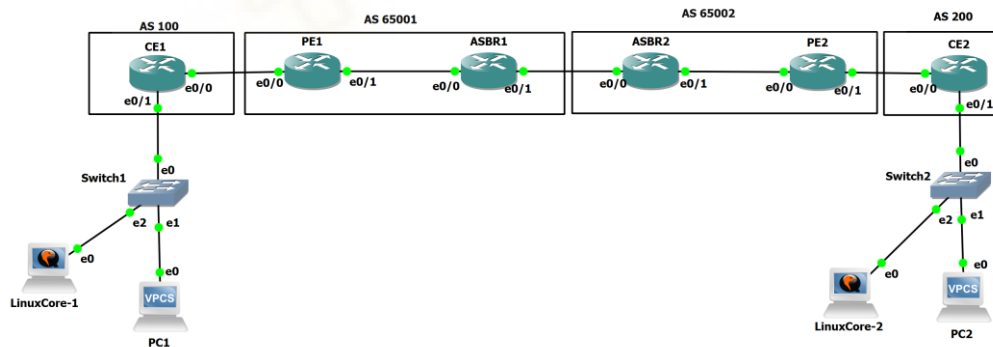


Figura 9: Topología del escenario A con dos routers de proveedor por lado.

En la Figura 9 puede verse la versión reducida del escenario A. La imagen muestra cuatro dominios claramente diferenciados, con AS 100 y AS 200 en los extremos de cliente y AS 65001 y AS 65002 como redes de proveedor. Dentro de esa estructura, CE1 y CE2 simbolizan las sedes del cliente, PE1 y PE2 realizan la entrada del servicio en la red del operador y ASBR1 y ASBR2 resuelven la interconexión entre ambos proveedores. También se aprecia que cada router CE conecta, mediante una interfaz Ethernet vertical,

con un switch Ethernet no configurable al que se asocian un nodo Linux Core y un VPCS. En el plano visual, la topología es lineal y limpia, algo que ayuda a entender que en esta primera versión se buscó una representación mínima del mecanismo, evitando añadir routers internos que todavía no eran necesarios para explicar el funcionamiento básico del escenario A. Además, en la implementación mostrada en GNS3 se trabajó con interfaces Ethernet identificadas como e0 y e1, lo que encaja con la plantilla realmente utilizada en el trabajo.

La construcción de esta versión siguió una lógica deliberadamente contenida. Primero se definieron los extremos de cliente y sus redes locales, conectando los equipos finales a cada CE a través del switch Ethernet. Después se desplegaron los routers de proveedor y se asignó a los enlaces relevantes el contexto de la VRF del servicio, de forma que el tráfico del cliente avanzara por la red sin salir de ese ámbito lógico. En esta topología todavía no era necesario introducir un núcleo interno de proveedor con funciones de tránsito, por lo que tampoco hacía falta reproducir mecanismos adicionales de distribución de etiquetas o de encaminamiento interior entre routers intermedios. Eso convierte a esta variante en un punto de partida muy útil dentro del trabajo, porque permite mostrar con bastante claridad la estructura mínima del escenario A antes de pasar a una versión más completa.

2.4.2. Escenario A con tres routers

La segunda versión del escenario A mantiene la misma filosofía general de la anterior, pero añade un router interno de tránsito en cada sistema autónomo de proveedor. Así, la cadena deja de ser CE-PE-ASBR y pasa a estructurarse como CE-PE-P-ASBR en un lado y ASBR-P-PE-CE en el otro. Este cambio no modifica el principio inter-AS del escenario A, ya que entre ASBR1 y ASBR2 se sigue trabajando con una interconexión VRF a VRF y sin MPLS entre sistemas autónomos, pero sí transforma el interior de cada proveedor en una red más próxima a un núcleo real (Cisco, 2022c). En la documentación del propio escenario se describe precisamente esta separación entre lo que ocurre dentro de cada AS y lo que ocurre en la frontera inter-AS, dejando claro que la ampliación no altera la naturaleza del escenario A, sino que la desarrolla en un entorno más representativo.

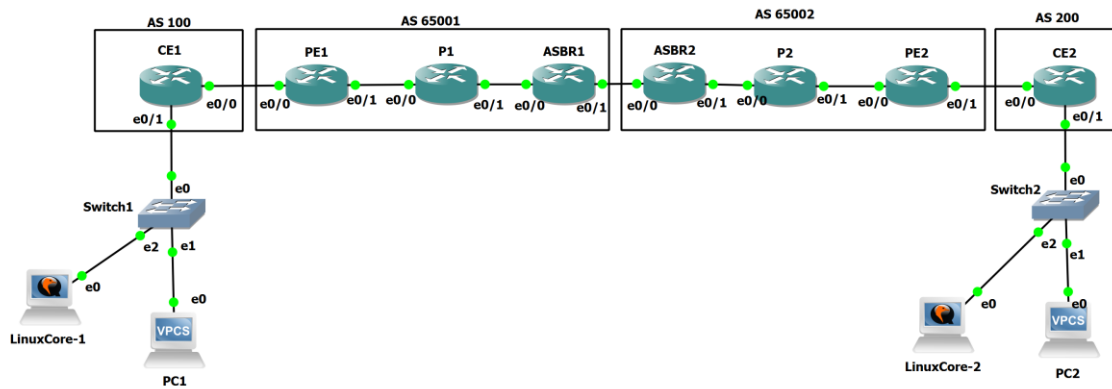


Figura 10: Topología del escenario A con tres routers de proveedor por lado.

La Figura 10 recoge la versión ampliada del escenario. A simple vista, la topología conserva la misma estructura general que la figura anterior, con los mismos extremos de cliente, los mismos switches Ethernet y la misma disposición simétrica entre ambos lados. Sin embargo, aquí ya aparecen P1 y P2 dentro de AS 65001 y AS 65002, respectivamente, actuando como routers internos de tránsito dentro del núcleo del proveedor y encargándose de transportar el tráfico entre los equipos de borde y los routers frontera, lo que convierte a cada proveedor en una pequeña red con tránsito interno. Esa diferencia es importante, aunque el escenario siga pareciendo similar, porque la presencia de un router intermedio obliga a separar con mayor claridad la parte de acceso, la parte de borde y la parte de núcleo dentro de cada AS. Por eso, aunque visualmente ambas variantes compartan casi la misma forma, internamente no responden al mismo grado de complejidad.

La razón de incorporar esta segunda versión fue precisamente esa. En cuanto aparece un router de tránsito entre el PE y el ASBR, el escenario deja de ser una simple concatenación de equipos adyacentes y pasa a requerir un mecanismo de alcance interno dentro del proveedor, además de un método de reenvío coherente a través del núcleo (Cisco, 2022b). En el trabajo se consideró suficiente detenerse en esta versión y no seguir aumentando el número de routers, porque una vez que el AS ya contiene tránsito interno, añadir más nodos repite la misma lógica básica de distribución de alcance y reenvío por saltos. Dicho de forma más simple, esta variante ya permite representar el momento en que la red del proveedor deja de ser trivial y empieza a comportarse como un pequeño core, que era exactamente lo que interesaba mostrar antes de pasar a los resultados.

2.5. Escenario B

Tras haber descrito el escenario A, el siguiente paso del trabajo consiste en construir el escenario B, cuya base teórica ya fue presentada en el apartado 1.6. A diferencia de la alternativa anterior, aquí la interconexión entre sistemas autónomos ya no se plantea como una prolongación VRF a VRF entre routers frontera, sino como un intercambio de rutas VPNv4 etiquetadas entre ASBR mediante eBGP. La documentación de Cisco describe esta opción indicando que los puertos de los ASBR deben estar preparados para recibir tráfico MPLS y que estos routers intercambian rutas VPNv4 y etiquetas VPN por eBGP (Cisco, 2020b), mientras que Juniper la presenta como una solución más escalable que el escenario A al evitar la necesidad de definir una VRF independiente por cada VPN común entre proveedores (Juniper Networks, 2025).

En este trabajo, el escenario B también se desarrolló en dos versiones. La primera mantiene una topología reducida, con un PE y un ASBR por lado dentro de cada proveedor, mientras que la segunda añade un router interno de tránsito en cada AS para obtener una red más próxima a un entorno de operador. Aunque visualmente ambas versiones resultan muy parecidas a las del escenario A, la diferencia importante no se encuentra en el dibujo físico, sino en la lógica interna del servicio. En escenario B, el enlace entre ASBR deja de funcionar como un simple traspaso IPv4 dentro de una VRF y pasa a convertirse en un punto de intercambio de rutas VPNv4 etiquetadas, lo que cambia por completo el papel de la frontera entre proveedores (Cisco, 2020b).

Precisamente por esa similitud visual con el escenario A, en este apartado interesa dejar claro que no basta con observar la topología para comprender el cambio entre escenarios. Los extremos de cliente, los switches Ethernet y la disposición general de los routers se mantienen prácticamente iguales, pero el comportamiento del plano de control ya no es el mismo. Por ello, esta sección se centra en explicar cómo se construyó cada variante del escenario B y qué sentido tiene su diseño dentro del trabajo, reservando el análisis detallado de las configuraciones y la validación de funcionamiento para los apartados que correspondan más adelante en la memoria.

2.5.1. Escenario B con dos routers

La primera versión del escenario B se planteó a partir de una topología reducida, compuesta en cada proveedor por un router PE y un router ASBR, además de los correspondientes routers CE en los extremos del cliente. La diferencia respecto al escenario A no reside en el número de equipos, sino en la forma en la que se establece la continuidad del servicio entre sistemas autónomos. En esta variante, los ASBR se conectan mediante una interfaz global preparada para tráfico MPLS y establecen entre sí una sesión eBGP a través de la cual intercambian rutas VPNv4 etiquetadas. Cisco indica expresamente que en el escenario B el ASBR funciona también como un PE hacia su propio AS, que no mantiene VRFs por cada servicio común en la frontera y que conserva todas o parte de las rutas VPNv4 que deben pasar al otro AS (Cisco, 2024b).

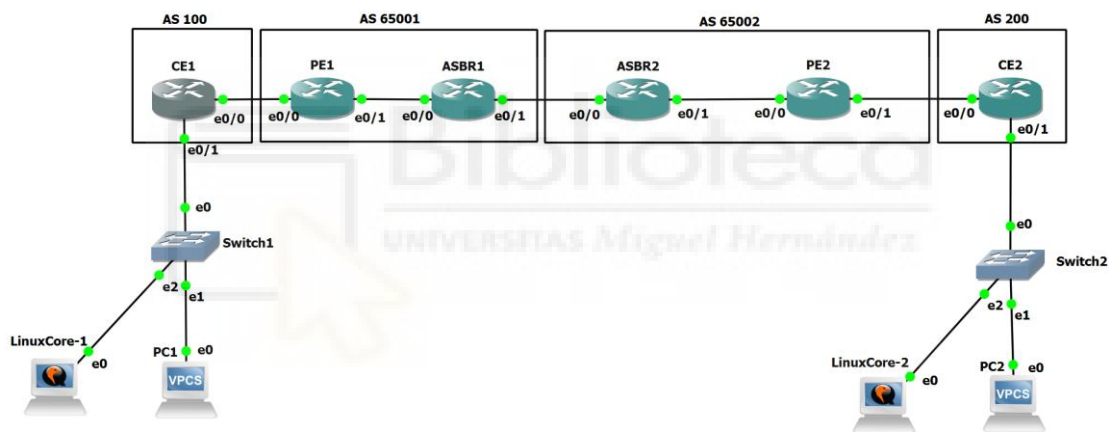


Figura 11: Topología del escenario B con dos routers de proveedor por lado.

La Figura 11 recoge esta versión reducida del escenario B. A simple vista, la imagen conserva la misma disposición general que la figura equivalente del escenario A, con los clientes en los extremos, los dos sistemas autónomos de proveedor en la zona central y un switch Ethernet no configurable en cada LAN de acceso. También siguen apareciendo los nodos Linux Core y VPCS conectados a cada CE, de modo que el trabajo mantiene una estructura homogénea en todos los escenarios. Sin embargo, aunque el dibujo físico apenas cambie, aquí la frontera inter-AS ya no representa una cesión VRF a VRF, sino un enlace global sobre el que se distribuyen rutas VPNv4 y etiquetas, que es precisamente lo que define la naturaleza del escenario B (Cisco, 2024b).

La construcción de esta versión siguió una lógica progresiva. Primero se mantuvieron los extremos de cliente tal y como ya habían sido definidos en el escenario A, con sus segmentos Ethernet y sus equipos finales. Después se desplegó la parte de proveedor, cuidando que el intercambio entre PE y ASBR dentro de cada AS mantuviera la coherencia del servicio y que la unión entre ASBR reflejara el mecanismo propio del escenario B. Juniper resume esta alternativa como una solución intermedia en términos de escalabilidad, más razonable que el escenario A cuando el número de VPN aumenta, precisamente porque evita trasladar la lógica completa de cada VRF a la frontera entre proveedores (Juniper Networks, 2025). Esa idea encaja bien con el sentido de esta versión reducida, que sirve para mostrar el mecanismo esencial del escenario B antes de pasar a una red interna más completa.

2.5.2. Escenario B con tres routers

La segunda versión del escenario B mantiene la misma filosofía general de la anterior, pero añade un router de tránsito dentro de cada sistema autónomo de proveedor. Así, la estructura deja de quedar reducida a CE-PE-ASBR y pasa a organizarse como CE-PE-P-ASBR en un lado y ASBR-P-PE-CE en el otro. Esta ampliación no altera el principio inter-AS del escenario B, ya que entre ASBR1 y ASBR2 se sigue trabajando con intercambio de rutas VPNv4 etiquetadas sobre un enlace global (Cisco, 2024b), pero sí transforma el interior de cada proveedor en una pequeña red con tránsito interno. Con ello, el escenario deja de representar únicamente una unión entre bordes y empieza a reflejar una organización más cercana a la de un proveedor real.

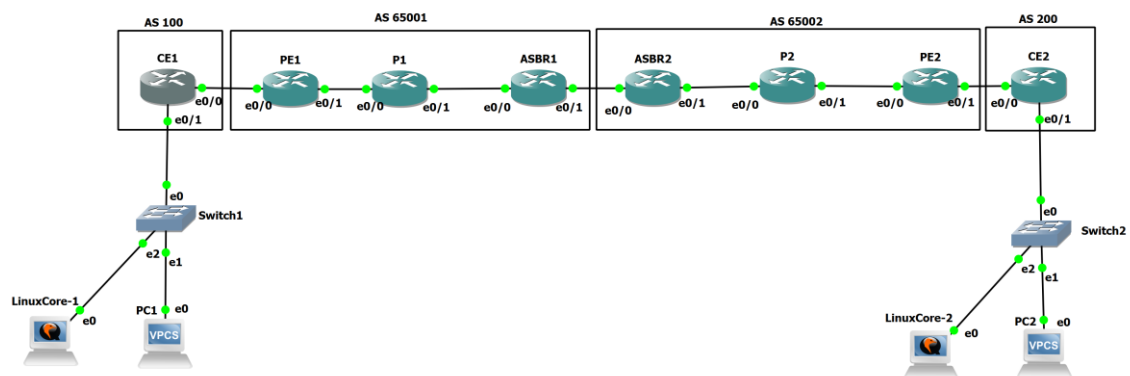


Figura 12: Topología del escenario B con tres routers de proveedor por lado.

La Figura 12 muestra esta versión ampliada. Igual que ocurría en el paso del escenario A simple a la ampliada, la topología conserva casi la misma forma general, con los mismos extremos de cliente, los mismos switches y una disposición simétrica entre ambos lados. La diferencia principal vuelve a estar en la presencia de los routers de tránsito P1 y P2 dentro de cada AS, que introducen una separación más clara entre acceso, borde y núcleo. Por eso, aunque visualmente esta imagen pueda recordar mucho a la del escenario A con tres routers, internamente responde a otra lógica, ya que aquí el intercambio entre proveedores se apoya en VPNv4 etiquetado y no en una relación VRF a VRF entre ASBR (Cisco, 2024b).

La razón de incorporar esta segunda versión fue la misma que en el escenario A: pasar de una topología muy simple a otra en la que ya exista tránsito interno dentro del proveedor. Una vez aparece ese nodo intermedio, el escenario deja de ser una conexión directa entre extremos de borde y pasa a representar una pequeña red de operador con núcleo interno. Cisco sitúa esta opción como una solución donde los ASBR intercambian VPNv4 entre sí (Cisco, 2024b), y Juniper la presenta como una alternativa más escalable que la A, pero todavía por debajo de la C (Juniper Networks, 2025). En ese sentido, esta versión ampliada resulta suficiente para reflejar el salto relevante dentro del trabajo, sin necesidad de seguir aumentando el número de routers y repetir una lógica que ya queda representada en esta topología.

2.6. Escenario C

Tras haber desarrollado los escenarios A y B, el siguiente paso del trabajo consistió en construir el escenario C, que es la alternativa más exigente de las tres desde el punto de vista del plano de control. Su idea básica no es prolongar la VPN hasta la frontera entre proveedores ni intercambiar directamente las rutas VPNv4 entre ASBR, sino separar con mayor claridad el transporte y el servicio. En esta opción, la información VPNv4 se mantiene en los extremos adecuados del servicio, mientras que entre sistemas autónomos se construye un camino etiquetado que permita sostener ese transporte de extremo a extremo. RFC 4364 presenta precisamente esta solución como la más escalable de las alternativas multi-AS (Rosen & Rekhter, 2006), y Juniper la describe también como la opción de mayor escalabilidad dentro de este grupo (Juniper Networks, 2025).

En este trabajo, el escenario C también se desarrolló en dos versiones. La primera parte de una topología reducida, con un PE y un ASBR por lado dentro de cada proveedor, mientras que la segunda añade un router de tránsito en cada AS para obtener una red más próxima a un entorno de operador. Igual que ocurría en el escenario B, el dibujo general puede parecerse bastante al de las opciones anteriores, ya que los clientes permanecen en los extremos, los proveedores en la zona central y la estructura física sigue siendo bastante simétrica. Sin embargo, aquí la diferencia importante tampoco está en la forma del esquema, sino en cómo se organiza el intercambio de rutas y etiquetas para construir el servicio. En otras palabras, el escenario puede parecer visualmente similar, pero internamente responde a una lógica distinta.

Además, en esta implementación concreta fue necesario adaptar la forma de construir el transporte etiquetado al entorno realmente utilizado en el trabajo. En la documentación técnica aportada para este escenario se indica que, al trabajar con imágenes antiguas como c3600, no siempre aparece disponible la address-family ipv4 labeled-unicast, por lo que se recurrió a la lógica clásica de RFC 3107 mediante el uso de send-label sobre vecinos IPv4 unicast. Esa decisión no altera el principio general del escenario C, sino que lo adapta a la plataforma realmente empleada, manteniendo la distribución de etiquetas en BGP para las rutas de transporte y reservando el intercambio VPNv4 para la parte correspondiente del servicio (Rekhter & Rosen, 2001). Por tanto, este escenario no solo es más complejo por diseño, sino también porque obliga a cuidar mejor cómo se materializa esa teoría en la imagen IOS realmente utilizada.

Precisamente por esa complejidad añadida, en este apartado conviene explicar con algo más de detalle cómo se construyó cada variante y qué sentido tiene dentro del trabajo. La versión reducida permite aislar el mecanismo esencial del escenario C sin introducir todavía routers internos de tránsito, mientras que la versión ampliada añade ya un pequeño núcleo dentro de cada proveedor y permite representar una situación más cercana a una red real. De este modo, el trabajo no se limita a mostrar que el escenario C existe como posibilidad teórica, sino que intenta reflejar cómo cambia su implementación cuando se pasa de una topología simple a otra en la que el transporte interno del proveedor deja de ser trivial.

2.6.1. Escenario C con dos routers

La variante más simple del escenario C se implementó sobre una topología básica en la que cada proveedor dispone de un router PE y un router ASBR, mientras que en los extremos se sitúan los routers CE asociados al cliente. A primera vista, esta disposición no difiere demasiado de la utilizada en los escenarios anteriores, pero la lógica interna ya no es la misma. Aquí el servicio se apoya en una sesión VPNv4 de extremo a extremo entre los PE, mientras que la parte de transporte entre dominios se resuelve mediante el intercambio de rutas IPv4 con etiqueta entre los equipos que sostienen la conectividad inter-AS. RFC 4364 describe esta opción precisamente como una combinación entre redistribución multihop de rutas VPN-IPv4 etiquetadas entre AS de origen y destino e intercambio de rutas IPv4 etiquetadas entre AS vecinos (Rosen & Rekhter, 2006).

En la implementación concreta de este trabajo, esa idea se materializó con una adaptación importante. La documentación del escenario indica que la imagen IOS empleada en la serie c3600 no ofrecía de forma explícita la address-family ipv4 labeled-unicast, por lo que el transporte se resolvió siguiendo el estilo RFC 3107 mediante vecinos IPv4 unicast con send-label, manteniendo al mismo tiempo la sesión VPNv4 directa entre PE1 y PE2. Esto significa que el escenario no se apoyó en una sintaxis moderna de labeled-unicast, sino en una forma igualmente válida y coherente con la plataforma disponible. En términos prácticos, esa decisión permitió reproducir la lógica del escenario C sin apartarse de las limitaciones reales del entorno de emulación utilizado en el trabajo.

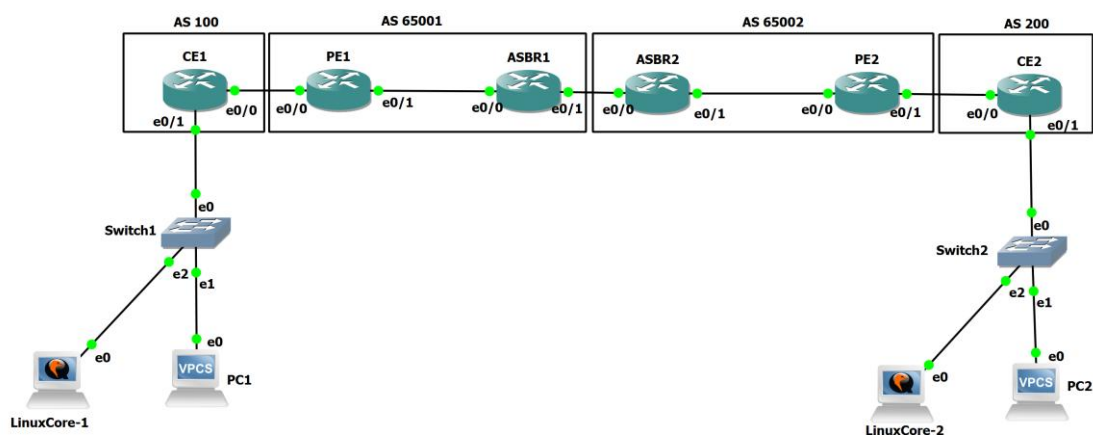


Figura 13: Topología del escenario C con dos routers de proveedor por lado.

Aunque la Figura 13 conserve una disposición visual muy similar a la del escenario B reducido, su interpretación interna es distinta. Los extremos de cliente siguen quedando representados por CE1 y CE2, los routers PE continúan ocupando el borde del proveedor y los ASBR mantienen su papel como frontera entre sistemas autónomos. Sin embargo, aquí la frontera inter-AS ya no debe entenderse como un punto donde se prolonga una VRF ni como un lugar donde simplemente se redistribuyen rutas VPNv4 entre ASBR, sino como una parte del transporte etiquetado que permite que la sesión VPNv4 entre los PE tenga soporte de extremo a extremo (Rosen & Rekhter, 2006). Esa diferencia es precisamente la que justifica tratar este escenario como una categoría propia dentro del trabajo.

La utilidad de esta versión reducida está en que permite mostrar el mecanismo esencial del escenario C sin introducir todavía un núcleo interno de proveedor con nodos de tránsito. De ese modo, el escenario conserva una topología relativamente limpia, pero ya incorpora el salto conceptual importante respecto a las opciones anteriores, que es la separación más clara entre el transporte inter-AS y el servicio VPN propiamente dicho. Por tanto, esta primera variante funciona como una base adecuada para entender el funcionamiento del escenario C antes de pasar a una topología más completa.

2.6.2. Escenario C con tres routers

La segunda versión del escenario C mantiene la misma filosofía general de la anterior, pero añade un router de tránsito dentro de cada sistema autónomo de proveedor. Así, la estructura deja de quedar reducida a CE-PE-ASBR y pasa a organizarse como CE-PE-P-ASBR en un lado y ASBR-P-PE-CE en el otro. Esta ampliación no altera el principio inter-AS del escenario C, ya que el servicio sigue apoyándose en una sesión VPNv4 de extremo a extremo entre los PE y el transporte continúa descansando sobre el intercambio de rutas IPv4 etiquetadas entre los equipos adecuados, pero sí transforma el interior de cada proveedor en una red más próxima a un núcleo real (Rosen & Rekhter, 2006). Con ello, el escenario deja de representar únicamente una unión entre bordes y empieza a reflejar una organización más cercana a la de un proveedor real.



Figura 14: Topología del escenario C con tres routers de proveedor por lado.

La Figura 14 mantiene una estructura similar a las topologías anteriores con tres routers, con los clientes en los extremos, los switches de acceso y una distribución simétrica entre ambos sistemas autónomos. La diferencia principal está en la presencia de P1 y P2 dentro de cada AS, que actúan como routers de tránsito interno y separan mejor las funciones de acceso, borde y núcleo. De esta forma, aunque visualmente el escenario se parezca a los anteriores, internamente representa una versión más completa del escenario C.

La incorporación de P1 y P2 permite observar mejor cómo se combinan el alcance interno del proveedor, el transporte etiquetado y la sesión VPNv4 extremo a extremo entre los PE. Por tanto, esta versión ampliada no se plantea solo para añadir más routers, sino para comprobar cómo se comporta el escenario C cuando el proveedor dispone de un núcleo interno más realista.

En ese sentido, la versión con tres routers por AS se consideró suficiente para reflejar el salto relevante dentro del trabajo. A partir de este punto, añadir más routers no introduciría un cambio conceptual importante, sino que ampliaría una lógica que ya queda representada en esta topología, como ocurre en los escenarios A y B. Esta variante, por tanto, permite recoger la estructura esencial del escenario C en una red más realista, manteniendo al mismo tiempo una complejidad asumible dentro de la memoria y del entorno de emulación utilizado.

3. RESULTADOS Y DISCUSIÓN

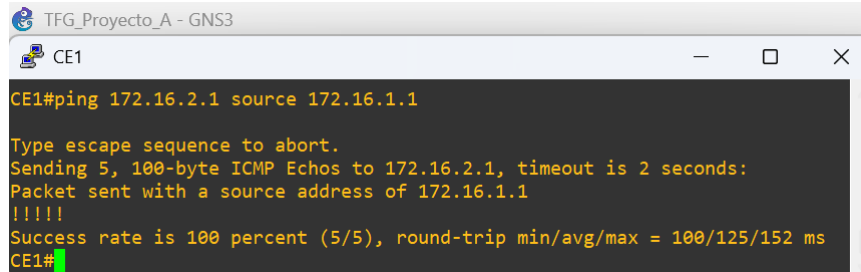
En este apartado se describirán y analizarán los resultados obtenidos siguiendo la metodología propuesta. Para ello, se presentarán las pruebas de conectividad y los comandos de verificación más relevantes de cada escenario, con el objetivo de comprobar que las configuraciones aplicadas permiten el intercambio de tráfico entre los extremos del cliente. En este apartado se explicarán las pruebas de conectividad y los comandos de verificación utilizados para comprobar el funcionamiento de cada escenario. En cada captura se indicará el tipo de prueba realizada, el equipo desde el que se ejecuta y la información que aporta al análisis. Además, en todas las capturas incluidas se muestra en la parte superior el nombre del proyecto abierto en GNS3, lo que permite asociar cada prueba con su escenario correspondiente y aporta mayor trazabilidad al proceso de validación realizado.

Igualmente, la estructura del apartado seguirá el mismo orden lógico que el desarrollo práctico del trabajo. Cada escenario tendrá su propio bloque de resultados, empezando por las pruebas de ping y continuando con los comandos show más importantes. Los ping servirán principalmente para demostrar la conectividad extremo a extremo, mientras que los comandos show permitirán justificar técnicamente por qué esa conectividad funciona, mostrando rutas en la VRF, vecinos BGP, intercambio VPNv4, etiquetas MPLS, sesiones LDP u OSPF según corresponda. De este modo, no solo se mostrará que el tráfico llega, sino también qué mecanismos internos permiten que ese resultado sea posible.

Al final del apartado se incluirá una sección adicional dedicada a comentarios y comprobaciones complementarias que han resultado importantes durante las pruebas. En ella se tratarán aspectos como la necesidad de utilizar ping con origen específico para representar correctamente el tráfico del cliente, la pérdida inicial de algunos paquetes mientras los protocolos terminan de converger y ciertos errores puntuales del entorno de emulación, como los relacionados con Telnet o QEMU. Además, se incorporará un resumen comparativo de los tres escenarios, indicando sus principales ventajas y desventajas a partir de los resultados obtenidos. Estos casos no forman parte del análisis individual de cada escenario, pero ayudan a interpretar mejor los resultados y a entender algunas situaciones que pueden aparecer durante la validación práctica en GNS3.

3.1. Escenario A con dos routers

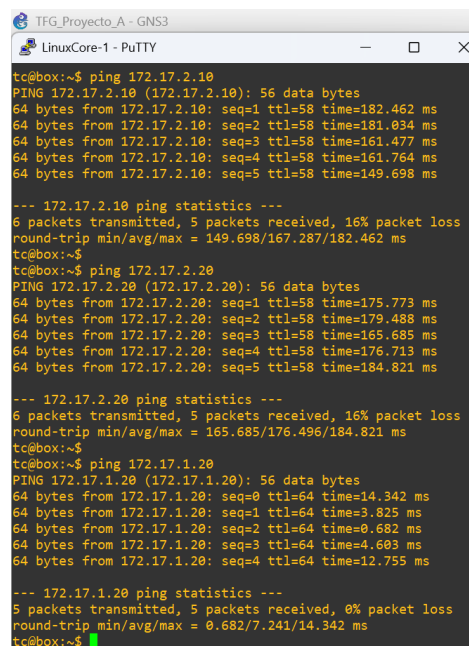
3.1.1. Pruebas de conectividad mediante ping



```
TFG_Proyecto_A - GNS3
CE1
CE1#ping 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
Packet sent with a source address of 172.16.1.1
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 100/125/152 ms
CE1#
```

Figura 15: Prueba de conectividad desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario A con dos routers.

En la Figura 15 se muestra una prueba realizada desde CE1 hacia la loopback remota de CE2, utilizando el comando `ping 172.16.2.1 source 172.16.1.1`. El uso del origen 172.16.1.1 permite comprobar la conectividad entre las redes lógicas de ambos extremos del cliente, y el resultado es satisfactorio, ya que se reciben los cinco paquetes enviados con una tasa de éxito del 100%. Más adelante, en el apartado 3.7, se comentará por qué este tipo de prueba puede comportarse de forma distinta cuando no se especifica un origen concreto.



```
TFG_Proyecto_A - GNS3
LinuxCore-1 - PuTTY
tc@box:~$ ping 172.17.2.10
PING 172.17.2.10 (172.17.2.10): 56 data bytes
64 bytes from 172.17.2.10: seq=1 ttl=58 time=182.462 ms
64 bytes from 172.17.2.10: seq=2 ttl=58 time=181.034 ms
64 bytes from 172.17.2.10: seq=3 ttl=58 time=161.477 ms
64 bytes from 172.17.2.10: seq=4 ttl=58 time=161.764 ms
64 bytes from 172.17.2.10: seq=5 ttl=58 time=149.698 ms

--- 172.17.2.10 ping statistics ---
6 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 149.698/167.287/182.462 ms
tc@box:~$
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=1 ttl=58 time=175.773 ms
64 bytes from 172.17.2.20: seq=2 ttl=58 time=179.488 ms
64 bytes from 172.17.2.20: seq=3 ttl=58 time=165.685 ms
64 bytes from 172.17.2.20: seq=4 ttl=58 time=176.713 ms
64 bytes from 172.17.2.20: seq=5 ttl=58 time=184.821 ms

--- 172.17.2.20 ping statistics ---
6 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 165.685/176.496/184.821 ms
tc@box:~$
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=0 ttl=64 time=14.342 ms
64 bytes from 172.17.1.20: seq=1 ttl=64 time=3.825 ms
64 bytes from 172.17.1.20: seq=2 ttl=64 time=0.682 ms
64 bytes from 172.17.1.20: seq=3 ttl=64 time=4.603 ms
64 bytes from 172.17.1.20: seq=4 ttl=64 time=12.755 ms

--- 172.17.1.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 0.682/7.241/14.342 ms
tc@box:~$
```

Figura 16: Pruebas de conectividad desde LinuxCore1 hacia LinuxCore2, PC2 y PC1 en el escenario A con dos routers.

En la Figura 16 se recogen varias pruebas desde LinuxCore1, cuya dirección es 172.17.1.10. En primer lugar, se comprueba la conectividad con LinuxCore2, situado en 172.17.2.10, y después con PC2, situado en 172.17.2.20. En ambos casos el resultado es correcto, ya que no se pierde ningún paquete. Sin embargo, los tiempos de respuesta son elevados en comparación con una comunicación local, porque el tráfico debe atravesar toda la topología del escenario A, pasando por CE1, PE1, ASBR1, ASBR2, PE2 y CE2 antes de llegar a la red remota.

En esa misma figura también aparece una prueba desde LinuxCore1 hacia PC1, cuya dirección es 172.17.1.20. En este caso los tiempos son mucho menores, ya que ambos equipos pertenecen a la misma LAN del extremo izquierdo y no necesitan recorrer la red de proveedor ni cruzar hacia el otro sistema autónomo. Esta diferencia entre los tiempos de respuesta confirma que el escenario distingue correctamente entre tráfico local dentro de una misma sede y tráfico remoto entre sedes distintas.

The image shows three terminal windows from a GNS3 simulation. The top-left window is 'PC1 - PuTTY' showing ping tests from 172.17.2.10 to 172.17.2.10, 172.17.2.20, and 172.17.1.10. The top-right window is 'PC2 - PuTTY' showing ping tests from 172.17.1.10 to 172.17.1.10, 172.17.1.20, and 172.17.2.10. The bottom window is 'LinuxCore2 - PuTTY' showing ping tests from 172.17.1.10 to 172.17.1.10, 172.17.1.20, and 172.17.2.20. Each test shows 5 packets with varying response times and TTL values.

```

PC1> ping 172.17.2.10
84 bytes from 172.17.2.10 icmp_seq=1 ttl=58 time=179.786 ms
84 bytes from 172.17.2.10 icmp_seq=2 ttl=58 time=180.388 ms
84 bytes from 172.17.2.10 icmp_seq=3 ttl=58 time=181.499 ms
84 bytes from 172.17.2.10 icmp_seq=4 ttl=58 time=181.892 ms
84 bytes from 172.17.2.10 icmp_seq=5 ttl=58 time=181.487 ms

PC1> ping 172.17.2.20
172.17.2.20 icmp_seq=1 timeout
172.17.2.20 icmp_seq=2 timeout
84 bytes from 172.17.2.20 icmp_seq=3 ttl=58 time=165.450 ms
84 bytes from 172.17.2.20 icmp_seq=4 ttl=58 time=180.453 ms
84 bytes from 172.17.2.20 icmp_seq=5 ttl=58 time=180.628 ms

PC1> ping 172.17.1.10
84 bytes from 172.17.1.10 icmp_seq=1 ttl=64 time=2.449 ms
84 bytes from 172.17.1.10 icmp_seq=2 ttl=64 time=1.023 ms
84 bytes from 172.17.1.10 icmp_seq=3 ttl=64 time=1.977 ms
84 bytes from 172.17.1.10 icmp_seq=4 ttl=64 time=1.013 ms
84 bytes from 172.17.1.10 icmp_seq=5 ttl=64 time=2.619 ms

PC2> ping 172.17.1.10
84 bytes from 172.17.1.10 icmp_seq=1 ttl=58 time=180.354 ms
84 bytes from 172.17.1.10 icmp_seq=2 ttl=58 time=180.650 ms
84 bytes from 172.17.1.10 icmp_seq=3 ttl=58 time=180.204 ms
84 bytes from 172.17.1.10 icmp_seq=4 ttl=58 time=179.712 ms
84 bytes from 172.17.1.10 icmp_seq=5 ttl=58 time=179.943 ms

PC2> ping 172.17.1.20
172.17.1.20 icmp_seq=1 timeout
172.17.1.20 icmp_seq=2 timeout
84 bytes from 172.17.1.20 icmp_seq=3 ttl=58 time=165.294 ms
84 bytes from 172.17.1.20 icmp_seq=4 ttl=58 time=180.702 ms
84 bytes from 172.17.1.20 icmp_seq=5 ttl=58 time=180.924 ms

PC2> ping 172.17.2.10
84 bytes from 172.17.2.10 icmp_seq=1 ttl=64 time=2.882 ms
84 bytes from 172.17.2.10 icmp_seq=2 ttl=64 time=2.430 ms
84 bytes from 172.17.2.10 icmp_seq=3 ttl=64 time=4.074 ms
84 bytes from 172.17.2.10 icmp_seq=4 ttl=64 time=2.193 ms
84 bytes from 172.17.2.10 icmp_seq=5 ttl=64 time=2.391 ms

tc@box:~$ ping 172.17.1.10
PING 172.17.1.10 (172.17.1.10): 56 data bytes
64 bytes from 172.17.1.10: seq=0 ttl=58 time=184.003 ms
64 bytes from 172.17.1.10: seq=1 ttl=58 time=180.372 ms
64 bytes from 172.17.1.10: seq=2 ttl=58 time=184.032 ms
64 bytes from 172.17.1.10: seq=3 ttl=58 time=160.807 ms
64 bytes from 172.17.1.10: seq=4 ttl=58 time=178.655 ms

--- 172.17.1.10 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 160.807/177.573/184.032 ms
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=1 ttl=58 time=150.998 ms
64 bytes from 172.17.1.20: seq=2 ttl=58 time=152.961 ms
64 bytes from 172.17.1.20: seq=3 ttl=58 time=179.676 ms
64 bytes from 172.17.1.20: seq=4 ttl=58 time=168.456 ms
64 bytes from 172.17.1.20: seq=5 ttl=58 time=158.477 ms

--- 172.17.1.20 ping statistics ---
5 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 150.998/162.113/179.676 ms
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=0 ttl=64 time=14.027 ms
64 bytes from 172.17.2.20: seq=1 ttl=64 time=17.725 ms
64 bytes from 172.17.2.20: seq=2 ttl=64 time=4.425 ms
64 bytes from 172.17.2.20: seq=3 ttl=64 time=6.717 ms
64 bytes from 172.17.2.20: seq=4 ttl=64 time=0.718 ms

--- 172.17.2.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 0.718/8.722/17.725 ms
tc@box:~$

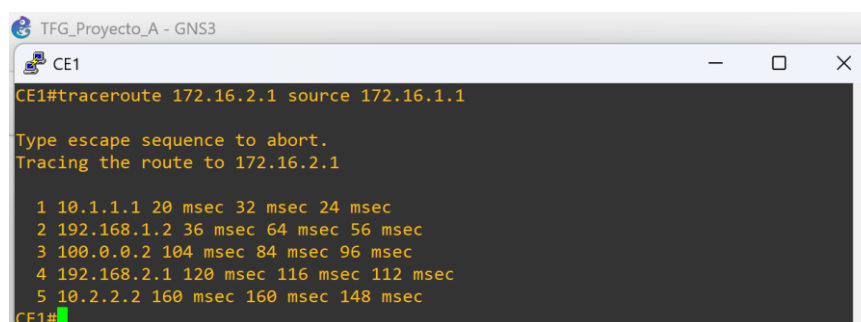
```

Figura 17: Pruebas de conectividad desde PC1, PC2 y LinuxCore2 hacia equipos locales y remotos en el escenario A con dos routers.

La Figura 17 reúne pruebas adicionales desde PC1, PC2 y LinuxCore2. En PC1 se observa que, al realizar ping hacia equipos remotos, puede aparecer un primer intento fallido antes de que las respuestas empiecen a recibirse correctamente. Este comportamiento es normal en entornos Ethernet, ya que antes de enviar tráfico IP hacia el siguiente salto es necesario resolver la dirección de enlace correspondiente mediante ARP. RFC 826 explica que, cuando no existe una entrada previa en la tabla de resolución, el equipo debe generar una solicitud ARP para obtener la dirección Ethernet asociada al destino o al siguiente salto, lo que puede provocar que el primer paquete no llegue a completarse a tiempo (Plummer, 1982).

En el resto de las capturas de la Figura 17 se aprecia que PC2 y LinuxCore2 también alcanzan correctamente los equipos del extremo contrario. Los tiempos hacia redes remotas vuelven a ser superiores a los tiempos locales, lo cual es coherente con el recorrido que debe seguir el tráfico a través de los routers del escenario. En el caso de LinuxCore2, algunos primeros tiempos pueden aparecer algo más altos que los siguientes. No lo interpretaría como un fallo de configuración, sino como un comportamiento razonable durante las primeras pruebas, provocado por la resolución inicial de direcciones, la actualización de cachés y la propia carga del entorno emulado. Lo importante para la validación es que, una vez establecida la comunicación, no se observa pérdida de paquetes y la conectividad entre sedes queda demostrada.

3.1.2. Trazado de ruta mediante traceroute



```
TFG_Proyecto_A - GNS3
CE1
CE1#traceroute 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Tracing the route to 172.16.2.1
 0 10.1.1.1 20 msec 32 msec 24 msec
 1 192.168.1.2 36 msec 64 msec 56 msec
 2 100.0.0.2 104 msec 84 msec 96 msec
 3 192.168.2.1 120 msec 116 msec 112 msec
 4 10.2.2.2 160 msec 160 msec 148 msec
CE1#
```

Figura 18: Trazado de ruta desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario A con dos routers.

En la Figura 18 podemos observar el resultado de traceroute 172.16.2.1 source 172.16.1.1, ejecutado desde CE1 hacia la loopback de CE2. El recorrido confirma que el

tráfico sigue el camino esperado dentro del escenario A con dos routers de proveedor por lado. El primer salto corresponde a PE1, identificado por la dirección 10.1.1.1. Después aparece ASBR1 mediante 192.168.1.2, seguido de ASBR2 mediante 100.0.0.2. A continuación, el tráfico llega a PE2 con 192.168.2.1 y finalmente alcanza CE2 en 10.2.2.2. Este trazado coincide con la estructura física y lógica configurada para el escenario A, donde la comunicación entre proveedores se mantiene dentro de la VRF del servicio.

La utilidad de esta prueba es que no solo confirma que el destino responde, sino que también permite comprobar por qué camino avanza el tráfico. Por tanto, el traceroute complementa a las pruebas de ping, ya que permite relacionar la conectividad final con los saltos intermedios que forman parte del escenario (Cisco, 2022b).

3.1.3. Verificación de rutas en la VRF

```

TFG_Proyecto_A - GNS3
PE1
PE1#show ip route vrf EMPRESA
Routing Table: EMPRESA
Codes: C - connected, S - static, R - RIP, M - mobile, B - BGP
        D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
        N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
        E1 - OSPF external type 1, E2 - OSPF external type 2
        i - IS-IS, su - IS-IS summary, L1 - IS-IS level-1, L2 - IS-IS level-2
        ia - IS-IS inter area, * - candidate default, U - per-user static route
        o - ODR, P - periodic downloaded static route

Gateway of last resort is not set

172.17.0.0/24 is subnetted, 2 subnets
B    172.17.1.0 [20/0] via 10.1.1.2, 00:20:49
B    172.17.2.0 [200/0] via 192.168.1.2, 00:20:49
172.16.0.0/24 is subnetted, 2 subnets
B    172.16.1.0 [20/0] via 10.1.1.2, 00:20:49
B    172.16.2.0 [200/0] via 192.168.1.2, 00:20:49
10.0.0.0/30 is subnetted, 1 subnets
C    10.1.1.0 is directly connected, Ethernet0/0
192.168.1.0/30 is subnetted, 1 subnets
C    192.168.1.0 is directly connected, Ethernet0/1
PE1#

```

Figura 19: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario A con dos routers.

Podemos observar en la Figura 19 que se muestra el comando `show ip route vrf EMPRESA` ejecutado en PE1. Al principio aparece el listado de códigos de la tabla de rutas, donde se indica qué letra corresponde a cada tipo de ruta. En este caso, lo más relevante se encuentra en las entradas marcadas con B, que identifican rutas aprendidas por BGP, y en las entradas marcadas con C, que corresponden a redes directamente conectadas. También aparece el mensaje `Gateway of last resort is not set`, lo que indica que no hay una ruta por defecto configurada en esa VRF (Cisco, 2022b). Esto no supone un problema en este escenario, ya que la conectividad necesaria se basa en rutas específicas hacia las redes del cliente y no en una salida genérica por defecto.

La tabla de PE1 confirma que este router conoce tanto las redes locales como las remotas de la VPN. Las redes 172.17.1.0/24 y 172.16.1.0/24 aparecen aprendidas a través de 10.1.1.2, que corresponde a CE1. En cambio, las redes 172.17.2.0/24 y 172.16.2.0/24 aparecen a través de 192.168.1.2, que corresponde a ASBR1. Esta diferencia es importante, porque demuestra que PE1 distingue correctamente entre las redes del cliente conectadas a su propio extremo y las redes que debe alcanzar a través del proveedor remoto.

```

ASBR1#show ip route vrf EMPRESA

Routing Table: EMPRESA
Codes: C - connected, S - static, R - RIP, M - mobile, B - BGP
        D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
        N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
        E1 - OSPF external type 1, E2 - OSPF external type 2
        i - IS-IS, su - IS-IS summary, L1 - IS-IS level-1, L2 - IS-IS level-2
        ia - IS-IS inter area, * - candidate default, U - per-user static route
        o - ODR, P - periodic downloaded static route

Gateway of last resort is not set

100.0.0.0/30 is subnetted, 1 subnets
C    100.0.0.0 is directly connected, Ethernet0/1
B    172.17.0.0/24 is subnetted, 2 subnets
B    172.17.1.0 [200/0] via 192.168.1.1, 00:33:47
B    172.17.2.0 [20/0] via 100.0.0.2, 00:33:47
B    172.16.0.0/24 is subnetted, 2 subnets
B    172.16.1.0 [200/0] via 192.168.1.1, 00:33:47
B    172.16.2.0 [20/0] via 100.0.0.2, 00:33:47
C    192.168.1.0/30 is subnetted, 1 subnets
C    192.168.1.0 is directly connected, Ethernet0/0
ASBR1#

```

Figura 20: Tabla de rutas de la VRF EMPRESA en ASBR1 en el escenario A con dos routers.

En la Figura 20 se muestra el mismo tipo de verificación en ASBR1. La tabla de la VRF EMPRESA confirma que ASBR1 conoce el enlace inter-AS 100.0.0.0/30 como red conectada y que también mantiene rutas BGP hacia las redes de ambos extremos del cliente. Las rutas del lado izquierdo, como 172.17.1.0/24 y 172.16.1.0/24, aparecen a través de 192.168.1.1, que corresponde a PE1. Las rutas del lado derecho, como 172.17.2.0/24 y 172.16.2.0/24, aparecen a través de 100.0.0.2, que corresponde a ASBR2. Por tanto, ASBR1 actúa como punto de continuidad entre el proveedor local y el proveedor remoto, manteniendo el tráfico dentro de la VRF propia del escenario A.

3.1.4. Verificación de BGP en la VRF

```

PE1#show ip bgp vpnv4 vrf EMPRESA
BGP table version is 5, local router ID is 1.1.1.1
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop        Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*> 172.16.1.0/24   10.1.1.2         0       0 100 i
*>i172.16.2.0/24  192.168.1.2      0       100  0 65002 200 i
*> 172.17.1.0/24  10.1.1.2         0       0 100 i
*>i172.17.2.0/24  192.168.1.2      0       100  0 65002 200 i
PE1#

```

Figura 21: Tabla BGP VPNv4 de la VRF EMPRESA en PE1 en el escenario A con dos routers.

En la Figura 21 observamos el comando `show ip bgp vpnv4 vrf EMPRESA` ejecutado en PE1. En la salida aparece el Route Distinguisher 65001:1, asociado a la VRF EMPRESA. También se observan los prefijos del cliente, concretamente 172.16.1.0/24, 172.16.2.0/24, 172.17.1.0/24 y 172.17.2.0/24. Los prefijos del lado local se alcanzan mediante 10.1.1.2, mientras que los prefijos remotos se alcanzan mediante 192.168.1.2. Esto confirma que PE1 ha aprendido correctamente las redes de ambos extremos y que BGP está transportando la información necesaria dentro del contexto de la VRF (Rosen & Rekhter, 2006).

```

ASBR1#show ip bgp vpnv4 vrf EMPRESA
BGP table version is 5, local router ID is 2.2.2.2
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop           Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*>i172.16.1.0/24   192.168.1.1         0      100      0 100 i
*> 172.16.2.0/24   100.0.0.2           0      100      0 65002 200 i
*>i172.17.1.0/24   192.168.1.1         0      100      0 100 i
*> 172.17.2.0/24   100.0.0.2           0      100      0 65002 200 i
ASBR1#

```

Figura 22: Tabla BGP VPNv4 de la VRF EMPRESA en ASBR1 en el escenario A con dos routers.

En la Figura 22 se recoge el mismo comando en ASBR1. El resultado es muy parecido al de PE1, ya que aparecen los mismos prefijos de cliente dentro de la VRF EMPRESA. La diferencia principal está en los siguientes saltos. En ASBR1, las rutas del lado izquierdo se alcanzan mediante 192.168.1.1, es decir, PE1, mientras que las rutas del lado derecho se alcanzan mediante 100.0.0.2, es decir, ASBR2. Esta comparación entre PE1 y ASBR1 se incluye en este primer escenario para mostrar que ambos routers disponen de una visión coherente de las rutas del servicio. En apartados posteriores, cuando una salida sea equivalente a otra ya comentada, bastará con indicar que mantiene el mismo comportamiento y destacar únicamente las diferencias relevantes.

En conjunto, las figuras 19, 20, 21 y 22 permiten justificar técnicamente los resultados de conectividad mostrados en las pruebas de ping y traceroute. No solo se comprueba que el tráfico llega al destino, sino que también se verifica que las rutas del cliente están presentes dentro de la VRF y que BGP conoce los prefijos de ambos extremos. Por tanto, el escenario A con dos routers queda validado tanto desde el punto de vista funcional como desde el punto de vista del plano de control.

3.2 Escenario A con tres routers

3.2.1. Pruebas de conectividad mediante ping

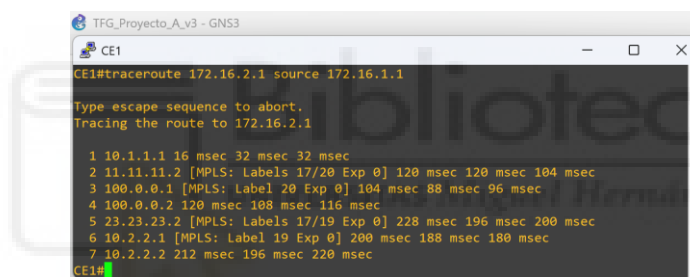
The screenshot displays five terminal windows from the GNS3 interface, showing the results of ping tests performed from different network devices in Scenario A. The top window, titled 'CE1', shows a successful ping to 172.16.2.1 with a 100% success rate. The middle-left window, 'LinuxCore-1 - PuTTY', shows ping tests to 172.17.2.10 and 172.17.2.20, both with 0% packet loss. The middle-right window, 'LinuxCore-2 - PuTTY', shows ping tests to 172.17.1.10 and 172.17.1.20, both with 16% packet loss. The bottom-left window, 'PC1 - PuTTY', shows ping tests to 172.17.2.10 and 172.17.1.10, with some timeouts and packet loss. The bottom-right window, 'PC2 - PuTTY', shows ping tests to 172.17.1.10 and 172.17.2.10, also with some timeouts and packet loss. The output includes details such as sequence numbers, TTL values, and round-trip times for each packet.

Figura 23: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario A con tres routers.

En la Figura 23 se agrupan las pruebas de conectividad principales del escenario A con tres routers por sistema autónomo. En la parte superior se observa el ping realizado desde CE1 hacia la loopback remota de CE2 mediante ping 172.16.2.1 source 172.16.1.1, obteniendo una tasa de éxito del 100%. En las capturas de LinuxCore1, LinuxCore2, PC1 y PC2 también se comprueba la conectividad entre los equipos finales de ambas sedes. Los tiempos hacia los equipos remotos son claramente superiores a los tiempos locales, ya que el tráfico debe atravesar el núcleo del proveedor, los routers P, los ASBR y el extremo remoto antes de alcanzar el destino.

En algunas pruebas aparecen uno o dos primeros paquetes sin respuesta, especialmente en los PC. Este comportamiento se puede relacionar con la resolución inicial de direcciones mediante ARP y con el establecimiento de las primeras entradas necesarias para reenviar el tráfico. En LinuxCore1 también se aprecia un caso con un 16% de pérdida hacia 172.17.2.20, aunque en la salida visible no aparece ningún timeout. Esto puede deberse a que, en LinuxCore, el ping continúa de forma indefinida hasta detenerlo manualmente con Ctrl + C, a diferencia de los PC, donde se envían cinco paquetes y el proceso termina solo. Para evitar esta situación, se podría limitar el número de paquetes directamente con un comando como `ping -c 5 172.17.2.20` (Linux man-pages project, 2026). Aun así, el resultado general es correcto, ya que las pruebas muestran conectividad entre los extremos locales y remotos del cliente.

3.2.2. Trazado de ruta mediante traceroute



```

TFG_Proyecto_A_v3 - GNS3
CE1
CE1#traceroute 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Tracing the route to 172.16.2.1
 0  10.1.1.1  16 msec 32 msec 32 msec
 1  11.11.11.2 [MPLS: Label 17/20 Exp 0] 120 msec 120 msec 104 msec
 2  100.0.0.1 [MPLS: Label 20 Exp 0] 104 msec 88 msec 96 msec
 3  100.0.0.2 120 msec 108 msec 116 msec
 4  23.23.23.2 [MPLS: Label 17/19 Exp 0] 228 msec 196 msec 200 msec
 5  10.2.2.1 [MPLS: Label 19 Exp 0] 200 msec 188 msec 180 msec
 6  10.2.2.2 212 msec 196 msec 220 msec
CE1#
  
```

Figura 24: Trazado de ruta desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario A con tres routers.

En la Figura 24 se recoge el resultado de `traceroute 172.16.2.1 source 172.16.1.1`, ejecutado desde CE1 hacia la loopback remota de CE2. El recorrido confirma que el tráfico atraviesa el escenario completo hasta llegar al extremo derecho. El primer salto corresponde a PE1 mediante 10.1.1.1, después aparece el tramo interno del proveedor con 11.11.11.2, y posteriormente se alcanza la zona de interconexión entre sistemas autónomos con 100.0.0.1 y 100.0.0.2. Finalmente, el tráfico continúa por el lado derecho hasta llegar a PE2 y CE2.

A diferencia del escenario A con dos routers, en este trazado aparecen referencias MPLS en varios saltos, como MPLS: Label 17/20, Label 20 o Label 17/19. Esto es coherente con la incorporación de routers P dentro de cada proveedor, ya que el tráfico pasa por un núcleo MPLS antes de llegar al siguiente extremo. Por tanto, el traceroute no

solo confirma la conectividad final, sino que también refleja que el camino seguido por los paquetes atraviesa los tramos internos etiquetados del escenario (Rosen et al., 2001).

3.2.3. Verificación de vecinos OSPF

The figure consists of three vertically stacked terminal windows from the GNS3 application, each showing the output of the 'show ip ospf neighbor' command on a different router.

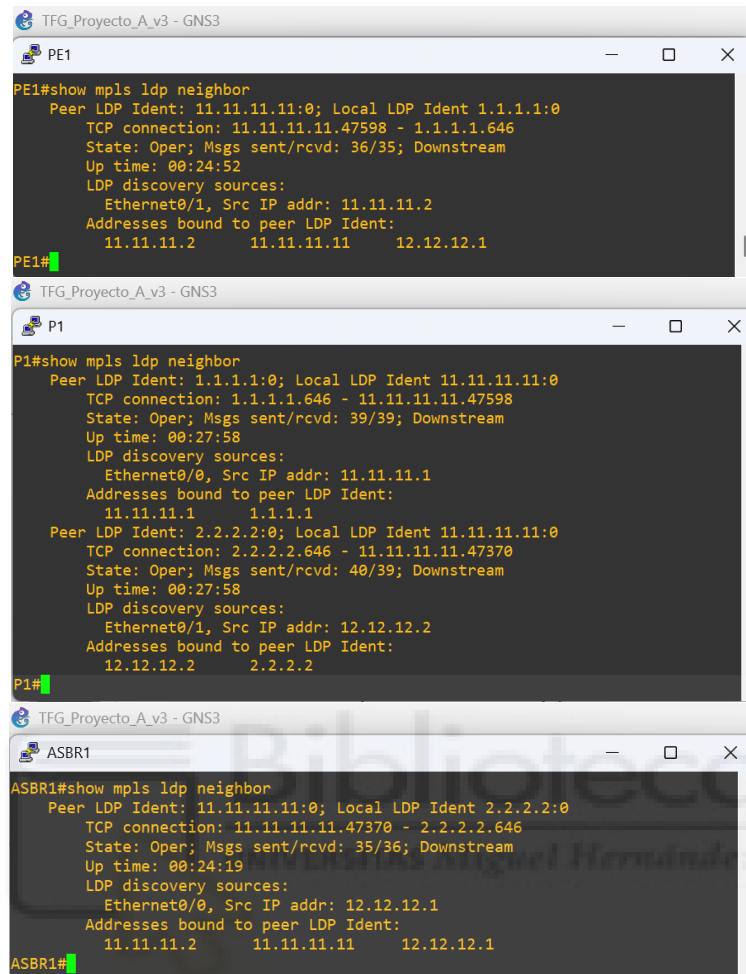
- Top window (PE1):** Shows one neighbor with ID 11.11.11.11, priority 1, state FULL/DR, dead time 00:00:33, address 11.11.11.2, and interface Ethernet0/1.
- Middle window (P1):** Shows two neighbors. The first has ID 2.2.2.2, priority 1, state FULL/BDR, dead time 00:00:35, address 12.12.12.2, and interface Ethernet0/1. The second has ID 1.1.1.1, priority 1, state FULL/BDR, dead time 00:00:34, address 11.11.11.1, and interface Ethernet0/0.
- Bottom window (ASBR1):** Shows one neighbor with ID 11.11.11.11, priority 1, state FULL/DR, dead time 00:00:38, address 12.12.12.1, and interface Ethernet0/0.

Figura 25: Vecinos OSPF en PE1, P1 y ASBR1 en el escenario A con tres routers.

En la Figura 25 se muestran las salidas de show ip ospf neighbor en PE1, P1 y ASBR1. En PE1 aparece como vecino P1, identificado por 11.11.11.11, a través de la dirección 11.11.11.2 por la interfaz Ethernet0/1. En P1 se observan dos vecinos, PE1 con router ID 1.1.1.1 y ASBR1 con router ID 2.2.2.2, lo que confirma que P1 actúa como router intermedio dentro del núcleo del proveedor. Por último, en ASBR1 aparece la vecindad con P1 mediante la dirección 12.12.12.1.

El estado FULL confirma que las adyacencias OSPF se han establecido correctamente. Las marcas DR y BDR son normales en enlaces Ethernet, donde OSPF puede elegir un router designado y un router designado de respaldo. Esta comprobación es importante porque OSPF proporciona la conectividad interna entre las loopbacks y enlaces del proveedor, lo que permite que después funcionen correctamente las sesiones de control asociadas al núcleo MPLS (Moy, 1998).

3.2.4. Verificación de vecinos LDP



The image displays three terminal windows from the GNS3 application, each showing the output of the 'show mpls ldp neighbor' command for a different router. The windows are titled 'PE1', 'P1', and 'ASBR1'. Each window shows the LDP session details, including the peer LDP ID, local LDP ID, TCP connection, state, and discovery sources.

```
PE1#show mpls ldp neighbor
Peer LDP Ident: 11.11.11.11:0; Local LDP Ident 1.1.1.1:0
TCP connection: 11.11.11.11.47598 - 1.1.1.1.646
State: Oper; Msgs sent/rcvd: 36/35; Downstream
Up time: 00:24:52
LDP discovery sources:
Ethernet0/1, Src IP addr: 11.11.11.2
Addresses bound to peer LDP Ident:
11.11.11.2    11.11.11.11    12.12.12.1
PE1#
```

```
P1#show mpls ldp neighbor
Peer LDP Ident: 1.1.1.1:0; Local LDP Ident 11.11.11.11:0
TCP connection: 1.1.1.1.646 - 11.11.11.11.47598
State: Oper; Msgs sent/rcvd: 39/39; Downstream
Up time: 00:27:58
LDP discovery sources:
Ethernet0/0, Src IP addr: 11.11.11.1
Addresses bound to peer LDP Ident:
11.11.11.1    1.1.1.1
Peer LDP Ident: 2.2.2.2:0; Local LDP Ident 11.11.11.11:0
TCP connection: 2.2.2.2.646 - 11.11.11.11.47370
State: Oper; Msgs sent/rcvd: 40/39; Downstream
Up time: 00:27:58
LDP discovery sources:
Ethernet0/1, Src IP addr: 12.12.12.2
Addresses bound to peer LDP Ident:
12.12.12.2    2.2.2.2
P1#
```

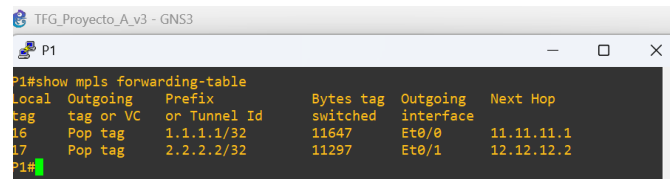
```
ASBR1#show mpls ldp neighbor
Peer LDP Ident: 11.11.11.11:0; Local LDP Ident 2.2.2.2:0
TCP connection: 11.11.11.11.47370 - 2.2.2.2.646
State: Oper; Msgs sent/rcvd: 35/36; Downstream
Up time: 00:24:19
LDP discovery sources:
Ethernet0/0, Src IP addr: 12.12.12.1
Addresses bound to peer LDP Ident:
11.11.11.2    11.11.11.11    12.12.12.1
ASBR1#
```

Figura 26: Vecinos LDP en PE1, P1 y ASBR1 en el escenario A con tres routers.

La Figura 26 permite comprobar el estado de LDP mediante show mpls ldp neighbor en PE1, P1 y ASBR1. En PE1 se observa una sesión LDP operativa con P1, cuyo identificador LDP es 11.11.11.11:0. En P1 aparecen dos vecinos LDP, uno hacia PE1 con identificador 1.1.1.1:0 y otro hacia ASBR1 con identificador 2.2.2.2:0. En ASBR1 se confirma también la sesión LDP con P1.

El estado Oper indica que las sesiones LDP están activas y funcionando. Además, las direcciones asociadas a cada vecino muestran qué direcciones están vinculadas al identificador LDP de cada router. Esta verificación es fundamental en este escenario, ya que la presencia de routers P implica que el núcleo interno del proveedor debe poder intercambiar etiquetas MPLS para transportar correctamente el tráfico entre PE1 y ASBR1 (Andersson et al., 2007).

3.2.5. Verificación de la tabla de reenvío MPLS



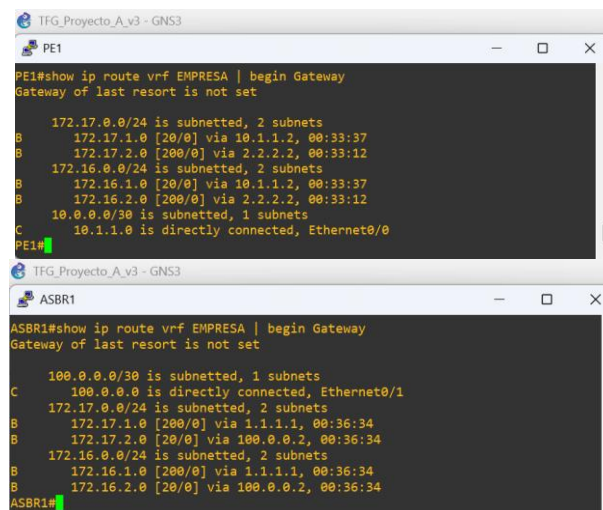
```
TFG_Proyecto_A_v3 - GNS3
P1
P1#show mpls forwarding-table
Local   Outgoing   Prefix      Bytes tag  Outgoing     Next Hop
tag     tag or VC  or Tunnel Id switched   interface
16      Pop tag    1.1.1.1/32  11647      Et0/0        11.11.11.1
17      Pop tag    2.2.2.2/32  11297      Et0/1        12.12.12.2
P1#
```

Figura 27: Tabla de reenvío MPLS en P1 en el escenario A con tres routers.

En la Figura 27 aparece el comando `show mpls forwarding-table` ejecutado en P1. La salida muestra entradas MPLS hacia las loopbacks 1.1.1.1/32 y 2.2.2.2/32, asociadas respectivamente a PE1 y ASBR1. Para cada prefijo se indica la etiqueta local, la acción de salida, la interfaz de salida y el siguiente salto.

La aparición de Pop tag indica que P1 realiza retirada de etiqueta en determinados trayectos, lo que resulta coherente cuando el siguiente salto está próximo al destino del tramo MPLS. Esta tabla confirma que P1 no solo tiene vecindades OSPF y LDP activas, sino que además dispone de entradas de reenvío MPLS reales. Por tanto, el router P cumple su función como nodo de tránsito dentro del núcleo del proveedor (Rosen et al., 2001).

3.2.6. Verificación de rutas en la VRF



```
TFG_Proyecto_A_v3 - GNS3
PE1
PE1#show ip route vrf EMPRESA | begin Gateway
Gateway of last resort is not set

 172.17.0.0/24 is subnetted, 2 subnets
B   172.17.1.0 [200/0] via 10.1.1.2, 00:33:37
B   172.17.2.0 [200/0] via 2.2.2.2, 00:33:12
 172.16.0.0/24 is subnetted, 2 subnets
B   172.16.1.0 [200/0] via 10.1.1.2, 00:33:37
B   172.16.2.0 [200/0] via 2.2.2.2, 00:33:12
C   10.0.0.0/30 is subnetted, 1 subnets
C   10.1.1.0 is directly connected, Ethernet0/0
PE1#

TFG_Proyecto_A_v3 - GNS3
ASBR1
ASBR1#show ip route vrf EMPRESA | begin Gateway
Gateway of last resort is not set

 100.0.0.0/30 is subnetted, 1 subnets
C   100.0.0.0 is directly connected, Ethernet0/1
C   172.17.0.0/24 is subnetted, 2 subnets
B   172.17.1.0 [200/0] via 1.1.1.1, 00:36:34
B   172.17.2.0 [200/0] via 100.0.0.2, 00:36:34
 172.16.0.0/24 is subnetted, 2 subnets
B   172.16.1.0 [200/0] via 1.1.1.1, 00:36:34
B   172.16.2.0 [200/0] via 100.0.0.2, 00:36:34
ASBR1#
```

Figura 28: Tablas de rutas de la VRF EMPRESA en PE1 y ASBR1 en el escenario A con tres routers.

En la Figura 28 se muestran las tablas de rutas de la VRF EMPRESA en PE1 y ASBR1 mediante `show ip route vrf EMPRESA`. En PE1 aparecen las redes locales 172.17.1.0/24 y 172.16.1.0/24 aprendidas a través de 10.1.1.2, que corresponde a CE1. Las redes remotas 172.17.2.0/24 y 172.16.2.0/24 aparecen a través de 2.2.2.2, que corresponde a ASBR1. Esto demuestra que PE1 conoce tanto las redes conectadas a su cliente local como las redes que se alcanzan a través del proveedor remoto.

En ASBR1 se aprecia una visión coherente con la función que realiza este router. Las redes del lado izquierdo se alcanzan mediante 1.1.1.1, es decir, hacia PE1, mientras que las redes del lado derecho se alcanzan mediante 100.0.0.2, que corresponde al ASBR del otro sistema autónomo. También aparece el enlace 100.0.0.0/30 como red directamente conectada. De esta forma, la figura 28 confirma que ASBR1 mantiene la continuidad de la VRF EMPRESA entre el proveedor local y el proveedor remoto.

3.2.7. Verificación de BGP en la VRF

```

TFG_Proyecto_A_v3 - GNS3
PE1
PE1#show ip bgp vpnv4 vrf EMPRESA
BGP table version is 7, local router ID is 1.1.1.1
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network          Next Hop          Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*> 172.16.1.0/24     10.1.1.2             0         0 100 i
*>i172.16.2.0/24     2.2.2.2              0      100   0 65002 200 i
*> 172.17.1.0/24     10.1.1.2             0         0 100 i
*>i172.17.2.0/24     2.2.2.2              0      100   0 65002 200 i
PE1#

TFG_Proyecto_A_v3 - GNS3
ASBR1
ASBR1#show ip bgp vpnv4 vrf EMPRESA
BGP table version is 7, local router ID is 2.2.2.2
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network          Next Hop          Metric LocPrf Weight Path
Route Distinguisher: 65001:10 (default for vrf EMPRESA)
*>i172.16.1.0/24     1.1.1.1             0      100   0 100 i
*> 172.16.2.0/24     100.0.0.2           0      100   0 65002 200 i
*>i172.17.1.0/24     1.1.1.1             0      100   0 100 i
*> 172.17.2.0/24     100.0.0.2           0      100   0 65002 200 i
ASBR1#

```

Figura 29: Tabla BGP VPNv4 de la VRF EMPRESA en PE1 y ASBR1 en el escenario A con tres routers.

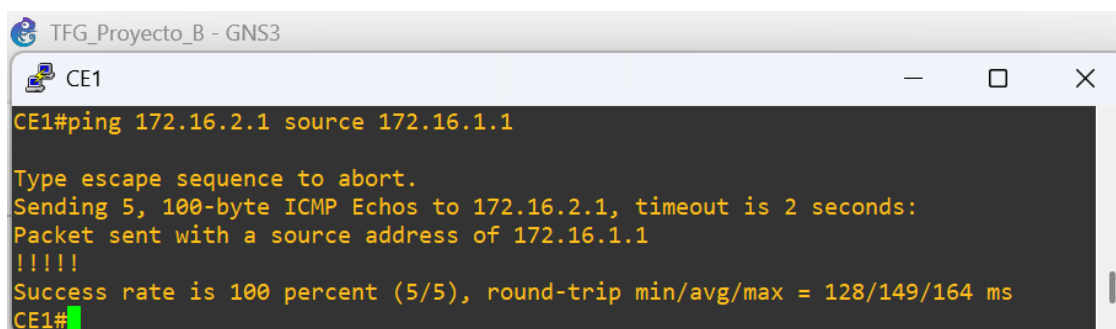
En la Figura 29 se observa el comando `show ip bgp vpnv4 vrf EMPRESA` ejecutado en PE1 y ASBR1. En PE1 aparecen los prefijos 172.16.1.0/24, 172.16.2.0/24, 172.17.1.0/24 y 172.17.2.0/24 dentro de la VRF EMPRESA. Las rutas locales se alcanzan mediante 10.1.1.2, mientras que las rutas remotas se alcanzan mediante 2.2.2.2. Además, los prefijos remotos muestran un camino AS que incluye 65002 200, lo que confirma que proceden del sistema autónomo remoto y del cliente situado al otro extremo.

En ASBR1 aparecen los mismos prefijos de cliente, pero con siguientes saltos adaptados a su posición dentro de la topología. Las redes del lado izquierdo se alcanzan mediante 1.1.1.1, mientras que las redes del lado derecho se alcanzan mediante 100.0.0.2. También se aprecia que ASBR1 utiliza su propio Route Distinguisher para la VRF EMPRESA, lo cual no impide que las rutas del servicio sean visibles dentro de la tabla BGP correspondiente. Esta salida confirma que BGP mantiene la información de encaminamiento necesaria para que las redes del cliente puedan comunicarse a través del escenario (Rosen & Rekhter, 2006).

Las figuras 23 a 29 validan el funcionamiento del escenario A con tres routers. Las pruebas de ping y traceroute demuestran la conectividad extremo a extremo, mientras que los comandos referidos a OSPF, LDP, MPLS, VRF y BGP permiten comprobar el funcionamiento interno del plano de control y del transporte MPLS. Por tanto, queda justificado que la incorporación de routers P no rompe la comunicación entre sedes, sino que añade un núcleo interno de proveedor que transporta el tráfico de forma correcta.

3.3 Escenario B con dos routers

3.3.1. Pruebas de conectividad mediante ping



```
TFG_Proyecto_B - GNS3
CE1
CE1#ping 172.16.2.1 source 172.16.1.1

Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
Packet sent with a source address of 172.16.1.1
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 128/149/164 ms
CE1#
```

The figure displays four terminal windows from a GNS3 project titled 'TFG_Proyecto_B - GNS3'. Each window shows the results of network connectivity tests performed from different hosts in the network.

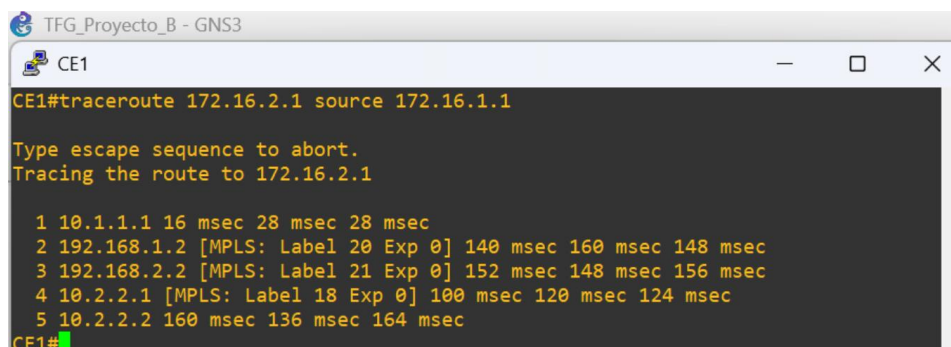
- LinuxCore-1 - PuTTY:** Shows ping results to 172.17.2.10 (16% packet loss) and 172.17.2.20 (16% packet loss). It also shows a successful ping to 172.17.1.20 (0% packet loss).
- LinuxCore-2 - PuTTY:** Shows ping results to 172.17.1.10 (0% packet loss) and 172.17.1.20 (16% packet loss). It also shows a successful ping to 172.17.2.20 (0% packet loss).
- PC1 - PuTTY:** Shows successful ping results to 172.17.2.10, 172.17.2.20, and 172.17.1.10.
- PC2 - PuTTY:** Shows successful ping results to 172.17.1.10, 172.17.1.20, and 172.17.2.10.

The output includes detailed ping statistics such as packets transmitted, packets received, packet loss percentage, and round-trip times. Some initial attempts show timeouts, which are followed by successful connections.

Figura 30: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario B con dos routers.

En la Figura 30 se agrupan las pruebas de conectividad del escenario B con dos routers por sistema autónomo. El ping realizado desde CE1 hacia la loopback remota de CE2 mediante ping 172.16.2.1 source 172.16.1.1 obtiene una tasa de éxito del 100%. Además, las pruebas desde LinuxCore1, LinuxCore2, PC1 y PC2 confirman la comunicación entre los equipos finales de ambas sedes. Al igual que en los apartados anteriores, los tiempos hacia destinos remotos son superiores a los tiempos locales debido al recorrido por la red de proveedor. Las pequeñas pérdidas parciales o timeouts iniciales responden al mismo comportamiento ya explicado en el apartado 3.2.1, por lo que no se consideran un fallo del escenario.

3.3.2. Trazado de ruta mediante traceroute

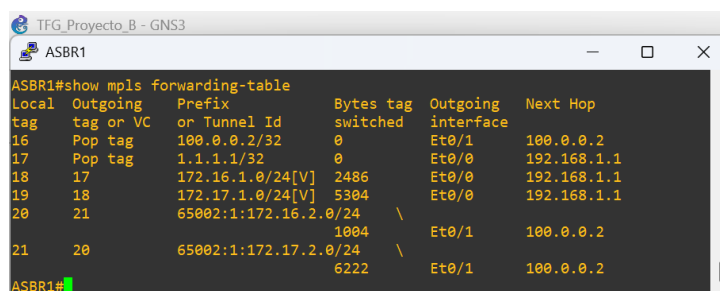


```
TFG_Proyecto_B - GNS3
CE1
CE1#traceroute 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Tracing the route to 172.16.2.1
 0 10.1.1.1 16 msec 28 msec 28 msec
 1 192.168.1.2 [MPLS: Label 20 Exp 0] 140 msec 160 msec 148 msec
 2 192.168.2.2 [MPLS: Label 21 Exp 0] 152 msec 148 msec 156 msec
 3 10.2.2.1 [MPLS: Label 18 Exp 0] 100 msec 120 msec 124 msec
 4 10.2.2.2 160 msec 136 msec 164 msec
CE1#
```

Figura 31: Trazado de ruta desde CE1 hacia la loopback remota de CE2 utilizando origen 172.16.1.1 en el escenario B con dos routers.

En la Figura 31 aparece el traceroute ejecutado desde CE1 hacia 172.16.2.1 utilizando como origen 172.16.1.1. El recorrido pasa por PE1, ASBR1, ASBR2, PE2 y finalmente CE2, por lo que coincide con la estructura del escenario B con dos routers por sistema autónomo. La diferencia es la aparición de una etiqueta MPLS entre los routers ASBR de cada sistema autónomo, además de en el interior de cada uno de ellos. Esto se aprecia en la aparición de etiquetas MPLS en los saltos intermedios, como Label 20, Label 21 y Label 18. Esto indica que el tráfico atraviesa el camino entre los routers PE origen y destino etiquetado entre los routers de proveedor, algo coherente con el funcionamiento del escenario B.

3.3.3. Verificación de la tabla de reenvío MPLS



```
TFG_Proyecto_B - GNS3
ASBR1
ASBR1#show mpls forwarding-table
Local tag Outgoing tag or VC Prefix or Tunnel Id Bytes tag switched Outgoing interface Next Hop
16 Pop tag 100.0.0.2/32 0 Et0/1 100.0.0.2
17 Pop tag 1.1.1.1/32 0 Et0/0 192.168.1.1
18 17 172.16.1.0/24[V] 2486 Et0/0 192.168.1.1
19 18 172.17.1.0/24[V] 5304 Et0/0 192.168.1.1
20 21 65002:1:172.16.2.0/24 \ 1004 Et0/1 100.0.0.2
21 20 65002:1:172.17.2.0/24 \ 6222 Et0/1 100.0.0.2
ASBR1#
```

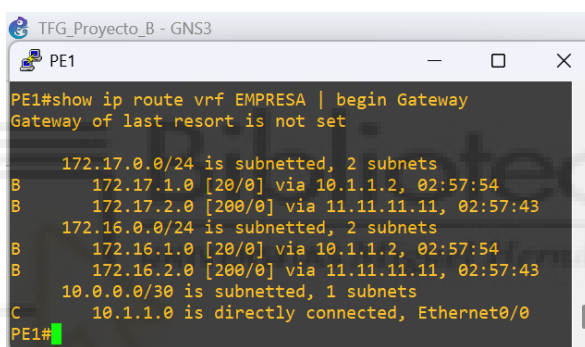
Figura 32: Tabla de reenvío MPLS en ASBR1 en el escenario B con dos routers.

En la Figura 32 se recoge el comando `show mpls forwarding-table` ejecutado en ASBR1. La tabla muestra entradas asociadas a prefijos de transporte, como 100.0.0.2/32 y 1.1.1.1/32, y también entradas vinculadas a prefijos de la VPN. Esto permite comprobar

que ASBR1 no solo participa en el intercambio BGP VPNv4, sino que además dispone de información de reenvío MPLS para transportar el tráfico etiquetado.

La aparición de acciones como Pop tag indica que en determinados casos ASBR1 retira la etiqueta antes de reenviar el paquete hacia el siguiente salto, lo cual sucede cuando el tráfico va a salir del sistema autónomo. También se observan etiquetas de salida asociadas a redes del cliente, lo que confirma que el escenario B utiliza un transporte MPLS para permitir la comunicación entre las sedes. Esta captura refuerza la diferencia con el escenario A, donde el intercambio entre proveedores se apoyaba directamente en la continuidad de la VRF.

3.3.4. Verificación de rutas en la VRF



```
TFG_Proyecto_B - GNS3
PE1
PE1#show ip route vrf EMPRESA | begin Gateway
Gateway of last resort is not set

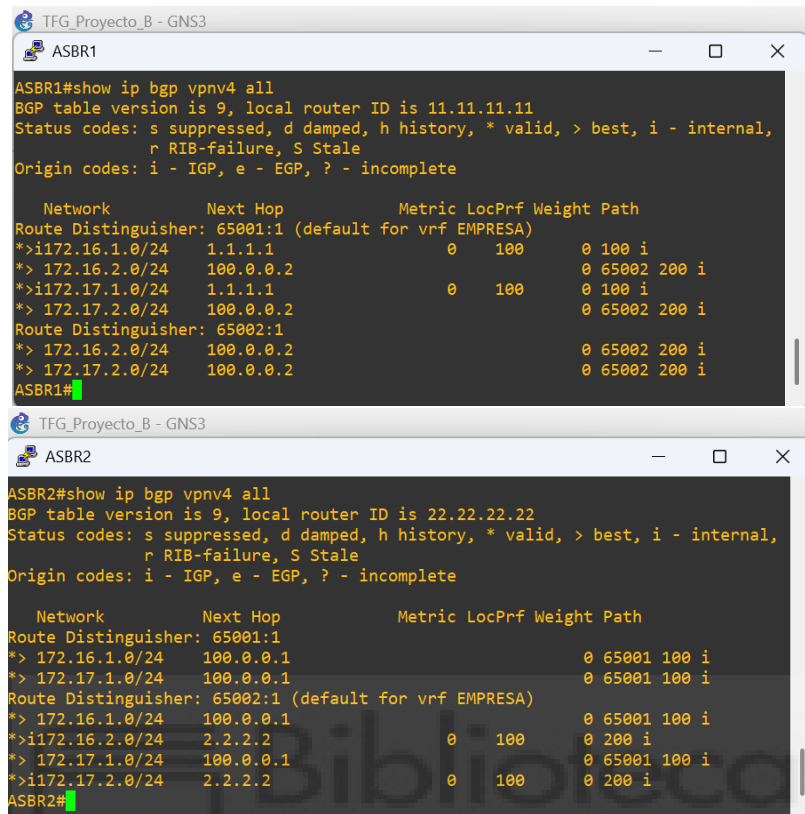
    172.17.0.0/24 is subnetted, 2 subnets
B       172.17.1.0 [20/0] via 10.1.1.2, 02:57:54
B       172.17.2.0 [200/0] via 11.11.11.11, 02:57:43
    172.16.0.0/24 is subnetted, 2 subnets
B       172.16.1.0 [20/0] via 10.1.1.2, 02:57:54
B       172.16.2.0 [200/0] via 11.11.11.11, 02:57:43
    10.0.0.0/30 is subnetted, 1 subnets
C       10.1.1.0 is directly connected, Ethernet0/0
PE1#
```

Figura 33: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario B con dos routers.

En la Figura 33 se muestra el comando `show ip route vrf EMPRESA` ejecutado en PE1. La tabla confirma que PE1 mantiene rutas hacia las redes locales del cliente, como 172.17.1.0/24 y 172.16.1.0/24, aprendidas a través de 10.1.1.2, que corresponde a CE1. También aparecen las redes remotas 172.17.2.0/24 y 172.16.2.0/24, aprendidas mediante 11.11.11.11, que corresponde a ASBR1.

Esta salida permite comprobar que PE1 dispone de las rutas necesarias dentro de la VRF EMPRESA para alcanzar tanto las redes de su sede local como las redes situadas en el otro extremo. Al igual que en escenarios anteriores, no es necesario repetir la misma comprobación en PE2 si se quiere reducir el número de capturas, ya que su comportamiento sería simétrico con las direcciones del lado derecho.

3.3.5. Verificación de rutas VPNv4 en los ASBR



The image shows two terminal windows from a GNS3 simulation. The top window is for ASBR1, showing the output of the command 'show ip bgp vpnv4 all'. It displays BGP table information for local router ID 11.11.11.11. The table lists routes for two Route Distinguishers: 65001:1 (default for vrf EMPRESA) and 65002:1. For 65001:1, routes 172.16.1.0/24 and 172.17.1.0/24 are learned via 1.1.1.1, while 172.16.2.0/24 and 172.17.2.0/24 are learned via 100.0.0.2. For 65002:1, routes 172.16.2.0/24 and 172.17.2.0/24 are learned via 100.0.0.2. The bottom window is for ASBR2, showing the output of 'show ip bgp vpnv4 all' for local router ID 22.22.22.22. It lists routes for Route Distinguishers 65001:1 and 65002:1. For 65001:1, routes 172.16.1.0/24 and 172.17.1.0/24 are learned via 100.0.0.1, while 172.16.2.0/24 and 172.17.2.0/24 are learned via 2.2.2.2. For 65002:1, routes 172.16.1.0/24 and 172.17.1.0/24 are learned via 100.0.0.1, while 172.16.2.0/24 and 172.17.2.0/24 are learned via 2.2.2.2.

```
ASBR1#show ip bgp vpnv4 all
BGP table version is 9, local router ID is 11.11.11.11
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop           Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*>i172.16.1.0/24    1.1.1.1              0      100      0 100 i
*> 172.16.2.0/24    100.0.0.2            0      100      0 65002 200 i
*>i172.17.1.0/24    1.1.1.1              0      100      0 100 i
*> 172.17.2.0/24    100.0.0.2            0      100      0 65002 200 i
Route Distinguisher: 65002:1
*> 172.16.2.0/24    100.0.0.2            0      100      0 65002 200 i
*> 172.17.2.0/24    100.0.0.2            0      100      0 65002 200 i
ASBR1#
```

```
ASBR2#show ip bgp vpnv4 all
BGP table version is 9, local router ID is 22.22.22.22
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

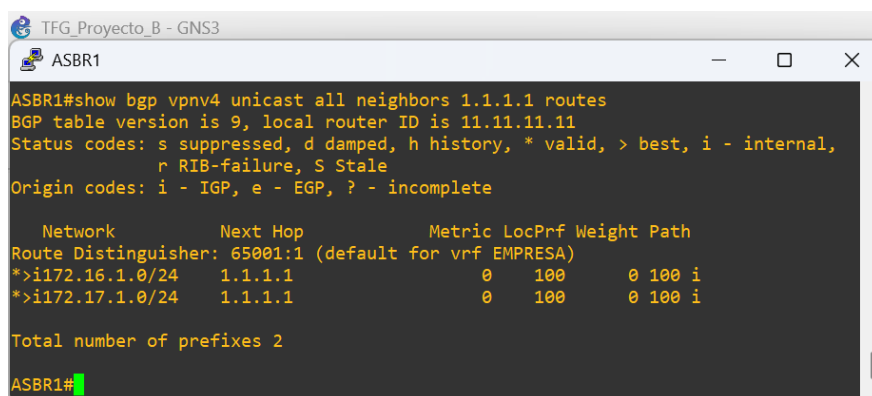
   Network        Next Hop           Metric LocPrf Weight Path
Route Distinguisher: 65001:1
*> 172.16.1.0/24    100.0.0.1            0      100      0 65001 100 i
*> 172.17.1.0/24    100.0.0.1            0      100      0 65001 100 i
Route Distinguisher: 65002:1 (default for vrf EMPRESA)
*> 172.16.1.0/24    100.0.0.1            0      100      0 65001 100 i
*>i172.16.2.0/24    2.2.2.2              0      100      0 200 i
*> 172.17.1.0/24    100.0.0.1            0      100      0 65001 100 i
*>i172.17.2.0/24    2.2.2.2              0      100      0 200 i
ASBR2#
```

Figura 34: Tabla BGP VPNv4 en ASBR1 y ASBR2 en el escenario B con dos routers.

En la Figura 34 se muestran las tablas BGP VPNv4 de ASBR1 y ASBR2 mediante `show ip bgp vpnv4 all`. Esta captura es especialmente importante en el escenario B, ya que los ASBR son los equipos encargados de intercambiar las rutas VPNv4 entre sistemas autónomos. En ASBR1 se observan rutas asociadas al Route Distinguisher 65001:1, correspondiente al lado izquierdo, y también rutas asociadas al Route Distinguisher 65002:1, correspondiente al lado derecho. Las redes del cliente local aparecen con siguiente salto 1.1.1.1, mientras que las redes remotas se alcanzan mediante 100.0.0.2.

En ASBR2 se aprecia el comportamiento equivalente desde el lado contrario. Las redes del lado izquierdo se alcanzan mediante 100.0.0.1, mientras que las redes propias del extremo derecho aparecen a través de 2.2.2.2. Esta comparación permite comprobar que ambos ASBR tienen visibilidad de las rutas VPNv4 necesarias y que el intercambio entre los dos sistemas autónomos se está realizando correctamente.

3.3.6. Verificación de rutas recibidas y anunciadas por BGP VPNv4



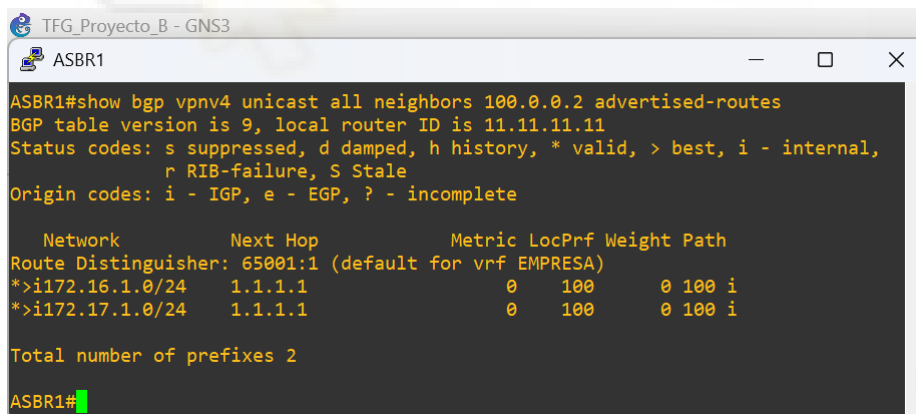
```
TFG_Proyecto_B - GNS3
ASBR1
ASBR1#show bgp vpnv4 unicast all neighbors 1.1.1.1 routes
BGP table version is 9, local router ID is 11.11.11.11
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network                Next Hop                Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*>i172.16.1.0/24          1.1.1.1                      0     100      0 100 i
*>i172.17.1.0/24          1.1.1.1                      0     100      0 100 i

Total number of prefixes 2
ASBR1#
```

Figura 35: Rutas VPNv4 recibidas por ASBR1 desde PE1 en el escenario B con dos routers.

En la Figura 35 se utiliza el comando `show bgp vpnv4 unicast all neighbors 1.1.1.1 routes` en ASBR1. La salida muestra las rutas recibidas desde PE1, concretamente los prefijos 172.16.1.0/24 y 172.17.1.0/24, ambos asociados al lado izquierdo del cliente. Esta comprobación permite ver de forma más precisa qué rutas está aprendiendo ASBR1 desde su vecino interno antes de propagarlas hacia el sistema autónomo remoto.



```
TFG_Proyecto_B - GNS3
ASBR1
ASBR1#show bgp vpnv4 unicast all neighbors 100.0.0.2 advertised-routes
BGP table version is 9, local router ID is 11.11.11.11
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network                Next Hop                Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*>i172.16.1.0/24          1.1.1.1                      0     100      0 100 i
*>i172.17.1.0/24          1.1.1.1                      0     100      0 100 i

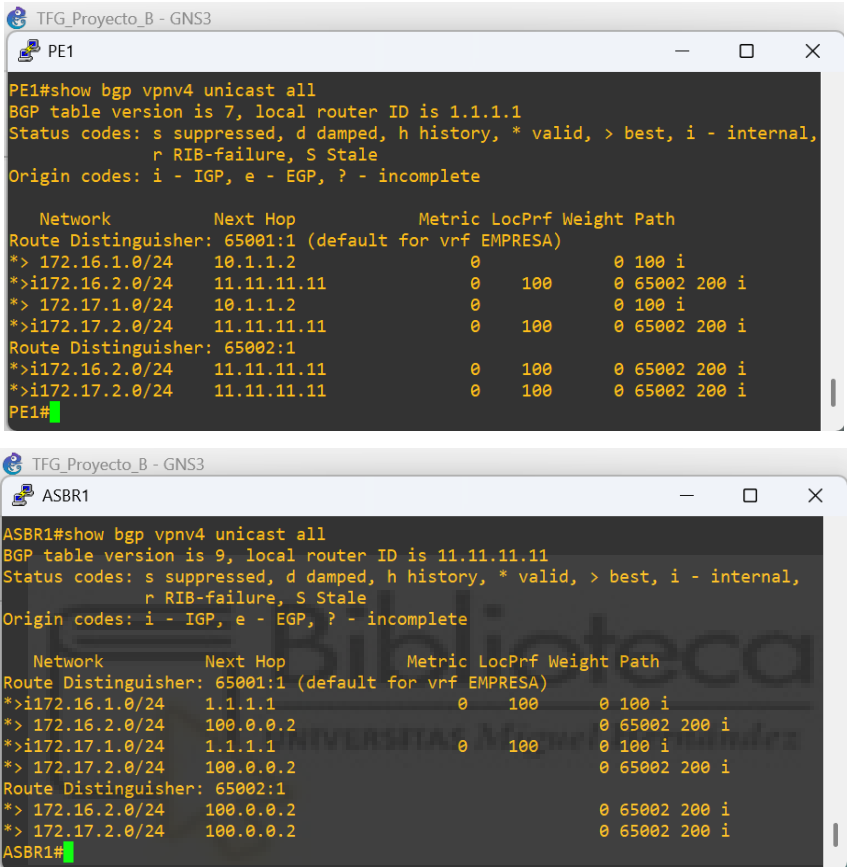
Total number of prefixes 2
ASBR1#
```

Figura 36: Rutas VPNv4 anunciadas por ASBR1 hacia ASBR2 en el escenario B con dos routers.

En la Figura 36 se muestra el comando `show bgp vpnv4 unicast all neighbors 100.0.0.2 advertised-routes`. En este caso se comprueba qué rutas anuncia ASBR1 hacia ASBR2. Los prefijos anunciados son 172.16.1.0/24 y 172.17.1.0/24, que corresponden a las redes del cliente situadas en el lado izquierdo. Esta salida complementa a la anterior,

ya que permite comprobar el flujo de rutas desde PE1 hacia ASBR1 y desde ASBR1 hacia el ASBR remoto.

3.3.7. Comparación de tablas BGP VPNv4 en PE1 y ASBR1



The figure consists of two screenshots of GNS3 terminal windows. The top window is titled 'PE1' and shows the output of the command 'show bgp vpnv4 unicast all'. It displays BGP table version 7, local router ID 1.1.1.1, and a list of routes for VRF EMPRESA. The bottom window is titled 'ASBR1' and shows the output of the same command. It displays BGP table version 9, local router ID 11.11.11.11, and a list of routes for VRF EMPRESA. Both windows show a table with columns: Network, Next Hop, Metric, LocPrf, Weight, and Path.

```
PE1#show bgp vpnv4 unicast all
BGP table version is 7, local router ID is 1.1.1.1
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop        Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*> 172.16.1.0/24   10.1.1.2          0      100     0 100 i
*>i172.16.2.0/24   11.11.11.11       0      100     0 65002 200 i
*> 172.17.1.0/24   10.1.1.2          0      100     0 100 i
*>i172.17.2.0/24   11.11.11.11       0      100     0 65002 200 i
Route Distinguisher: 65002:1
*>i172.16.2.0/24   11.11.11.11       0      100     0 65002 200 i
*>i172.17.2.0/24   11.11.11.11       0      100     0 65002 200 i
PE1#
```

```
ASBR1#show bgp vpnv4 unicast all
BGP table version is 9, local router ID is 11.11.11.11
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop        Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*>i172.16.1.0/24   1.1.1.1          0      100     0 100 i
*> 172.16.2.0/24   100.0.0.2        0      100     0 65002 200 i
*>i172.17.1.0/24   1.1.1.1          0      100     0 100 i
*> 172.17.2.0/24   100.0.0.2        0      100     0 65002 200 i
Route Distinguisher: 65002:1
*> 172.16.2.0/24   100.0.0.2        0      100     0 65002 200 i
*> 172.17.2.0/24   100.0.0.2        0      100     0 65002 200 i
ASBR1#
```

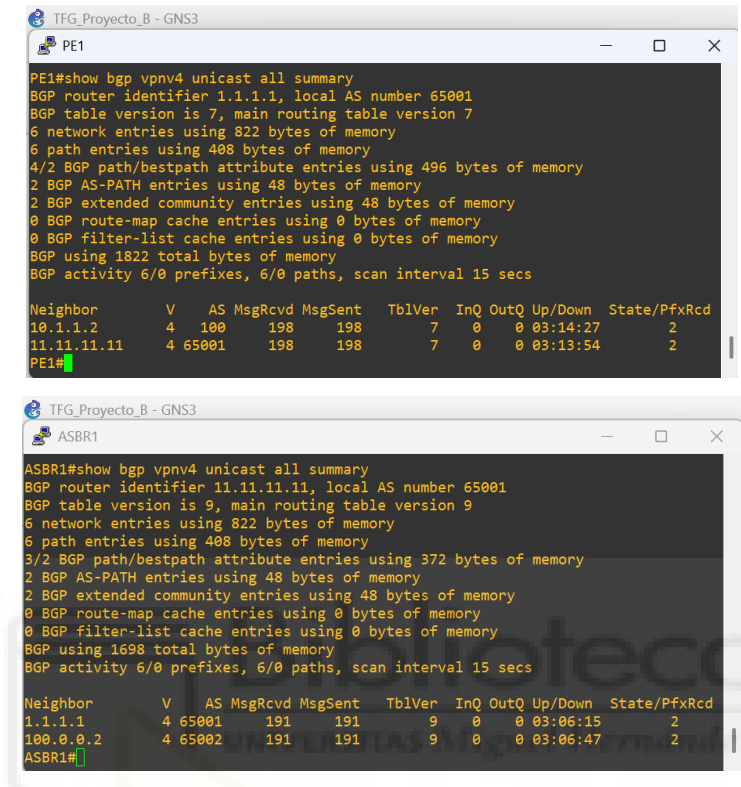
Figura 37: Tabla BGP VPNv4 en PE1 y ASBR1 en el escenario B con dos routers.

En la Figura 37 se comparan las salidas de show bgp vpnv4 unicast all en PE1 y ASBR1. En PE1 aparecen las rutas de la VRF EMPRESA desde la perspectiva del router de borde del cliente. Las redes locales se alcanzan mediante 10.1.1.2, mientras que las redes remotas se alcanzan mediante 11.11.11.11, que corresponde a ASBR1. Esto confirma que PE1 recibe del ASBR la información necesaria para llegar al extremo remoto.

En ASBR1 la tabla muestra una visión más amplia del intercambio VPNv4. Las rutas del lado izquierdo aparecen mediante 1.1.1.1, mientras que las rutas del lado derecho se alcanzan mediante 100.0.0.2. Esta diferencia es relevante porque PE1 actúa como router de entrada del cliente a la VPN, mientras que ASBR1 actúa como punto de intercambio entre sistemas autónomos. Por tanto, aunque ambos comandos muestran

información VPNv4, la función que representa cada salida dentro del escenario no es la misma.

3.3.8. Resumen de vecinos BGP VPNv4



The figure consists of two screenshots of GNS3 terminal windows. The top window is titled 'PE1' and shows the output of the command 'show bgp vpnv4 unicast all summary'. The output includes BGP router identifier 1.1.1.1, local AS number 65001, BGP table version 7, and a summary of network, path, and attribute entries. It also shows a table of neighbors: 10.1.1.2 and 11.11.11.11, both with AS 65001, 198 messages received and sent, and a state of 2. The bottom window is titled 'ASBR1' and shows the output of the same command. It includes BGP router identifier 11.11.11.11, local AS number 65001, BGP table version 9, and a summary of network, path, and attribute entries. It also shows a table of neighbors: 1.1.1.1 and 100.0.0.2, both with AS 65001, 191 messages received and sent, and a state of 2.

```
PE1#show bgp vpnv4 unicast all summary
BGP router identifier 1.1.1.1, local AS number 65001
BGP table version is 7, main routing table version 7
6 network entries using 822 bytes of memory
6 path entries using 408 bytes of memory
4/2 BGP path/bestpath attribute entries using 496 bytes of memory
2 BGP AS-PATH entries using 48 bytes of memory
2 BGP extended community entries using 48 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
BGP using 1822 total bytes of memory
BGP activity 6/0 prefixes, 6/0 paths, scan interval 15 secs

Neighbor      V    AS MsgRcvd MsgSent  TblVer  InQ OutQ Up/Down  State/PfxRcd
10.1.1.2      4   100    198    198      7    0    0 03:14:27      2
11.11.11.11   4  65001    198    198      7    0    0 03:13:54      2
PE1#
```

```
ASBR1#show bgp vpnv4 unicast all summary
BGP router identifier 11.11.11.11, local AS number 65001
BGP table version is 9, main routing table version 9
6 network entries using 822 bytes of memory
6 path entries using 408 bytes of memory
3/2 BGP path/bestpath attribute entries using 372 bytes of memory
2 BGP AS-PATH entries using 48 bytes of memory
2 BGP extended community entries using 48 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
BGP using 1698 total bytes of memory
BGP activity 6/0 prefixes, 6/0 paths, scan interval 15 secs

Neighbor      V    AS MsgRcvd MsgSent  TblVer  InQ OutQ Up/Down  State/PfxRcd
1.1.1.1       4  65001    191    191      9    0    0 03:06:15      2
100.0.0.2     4  65002    191    191      9    0    0 03:06:47      2
ASBR1#
```

Figura 38: Resumen de vecinos BGP VPNv4 en PE1 y ASBR1 en el escenario B con dos routers.

En la Figura 38 se muestra el comando `show bgp vpnv4 unicast all summary` en PE1 y ASBR1. Este comando ofrece una vista resumida de los vecinos BGP asociados a la familia VPNv4 unicast. En PE1 aparecen dos vecinos principales: 10.1.1.2, correspondiente a CE1 dentro de la VRF, y 11.11.11.11, correspondiente a ASBR1. En ambos casos se reciben dos prefijos, lo que confirma que las sesiones están establecidas y que PE1 dispone de las rutas necesarias del cliente.

En ASBR1 también aparecen dos vecinos, pero con una función diferente. El vecino 1.1.1.1 corresponde a PE1 dentro del mismo sistema autónomo, mientras que 100.0.0.2 corresponde a ASBR2 en el sistema autónomo remoto. Esta diferencia es clave en el escenario B, ya que ASBR1 actúa como punto intermedio entre el PE local y el ASBR remoto. Además, el comando utilizado en la captura, `show bgp vpnv4 unicast all summary`, proporciona una verificación equivalente a `show ip bgp vpnv4 all summary`,

por lo que no se ha incluido una captura adicional con esa otra sintaxis para evitar repetir información prácticamente idéntica.

Las figuras utilizadas en este apartado validan el escenario B con dos routers por sistema autónomo. Las pruebas de ping y traceroute demuestran la conectividad extremo a extremo, mientras que las salidas referidas a MPLS, VRF y BGP VPNv4 permiten justificar técnicamente cómo se consigue esa conectividad. La diferencia principal frente al escenario A es que ahora los ASBR participan directamente en el intercambio de rutas VPNv4 entre sistemas autónomos, mientras que el transporte se apoya en MPLS.

3.4 Escenario B con tres routers

3.4.1. Pruebas de conectividad mediante ping

```
TFG_Proyecto_B_v3 - GNS3
CE1
CE1#ping 172.16.2.1 source 172.16.1.1

Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
Packet sent with a source address of 172.16.1.1
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 180/208/240 ms
CE1#

TFG_Proyecto_B_v3 - GNS3
LinuxCore-1 - PuTTY
tc@box:~$ ping 172.17.2.10
PING 172.17.2.10 (172.17.2.10): 56 data bytes
64 bytes from 172.17.2.10: seq=1 ttl=56 time=242.390 ms
64 bytes from 172.17.2.10: seq=2 ttl=56 time=236.856 ms
64 bytes from 172.17.2.10: seq=3 ttl=56 time=235.356 ms
64 bytes from 172.17.2.10: seq=4 ttl=56 time=238.813 ms
64 bytes from 172.17.2.10: seq=5 ttl=56 time=247.673 ms

--- 172.17.2.10 ping statistics ---
6 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 235.356/240.217/247.673 ms
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=1 ttl=56 time=240.834 ms
64 bytes from 172.17.2.20: seq=2 ttl=56 time=238.022 ms
64 bytes from 172.17.2.20: seq=3 ttl=56 time=249.121 ms
64 bytes from 172.17.2.20: seq=4 ttl=56 time=216.329 ms
64 bytes from 172.17.2.20: seq=5 ttl=56 time=240.904 ms

--- 172.17.2.20 ping statistics ---
6 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 216.329/237.042/249.121 ms
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=0 ttl=64 time=3.994 ms
64 bytes from 172.17.1.20: seq=1 ttl=64 time=5.560 ms
64 bytes from 172.17.1.20: seq=2 ttl=64 time=0.712 ms
64 bytes from 172.17.1.20: seq=3 ttl=64 time=20.321 ms

--- 172.17.1.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 0.712/6.687/20.321 ms
tc@box:~$

TFG_Proyecto_B_v3 - GNS3
LinuxCore-2 - PuTTY
tc@box:~$ ping 172.17.1.10
PING 172.17.1.10 (172.17.1.10): 56 data bytes
64 bytes from 172.17.1.10: seq=0 ttl=56 time=242.969 ms
64 bytes from 172.17.1.10: seq=1 ttl=56 time=228.557 ms
64 bytes from 172.17.1.10: seq=2 ttl=56 time=246.596 ms
64 bytes from 172.17.1.10: seq=3 ttl=56 time=241.255 ms
64 bytes from 172.17.1.10: seq=4 ttl=56 time=219.880 ms

--- 172.17.1.10 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 219.880/235.851/246.596 ms
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=1 ttl=56 time=235.912 ms
64 bytes from 172.17.1.20: seq=2 ttl=56 time=242.275 ms
64 bytes from 172.17.1.20: seq=3 ttl=56 time=235.928 ms
64 bytes from 172.17.1.20: seq=4 ttl=56 time=231.191 ms
64 bytes from 172.17.1.20: seq=5 ttl=56 time=257.844 ms

--- 172.17.1.20 ping statistics ---
6 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 231.191/240.630/257.844 ms
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=0 ttl=64 time=2.379 ms
64 bytes from 172.17.2.20: seq=1 ttl=64 time=5.295 ms
64 bytes from 172.17.2.20: seq=2 ttl=64 time=8.143 ms
64 bytes from 172.17.2.20: seq=3 ttl=64 time=1.187 ms
64 bytes from 172.17.2.20: seq=4 ttl=64 time=5.604 ms

--- 172.17.2.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 1.187/4.521/8.143 ms
tc@box:~$

TFG_Proyecto_B_v3 - GNS3
PC1 - PuTTY
PC1> ping 172.17.2.10
84 bytes from 172.17.2.10: icmp_seq=1 ttl=56 time=240.856 ms
84 bytes from 172.17.2.10: icmp_seq=2 ttl=56 time=240.316 ms
84 bytes from 172.17.2.10: icmp_seq=3 ttl=56 time=240.003 ms
84 bytes from 172.17.2.10: icmp_seq=4 ttl=56 time=240.601 ms
84 bytes from 172.17.2.10: icmp_seq=5 ttl=56 time=240.520 ms

PC1> ping 172.17.2.20
84 bytes from 172.17.2.20: icmp_seq=1 ttl=56 time=240.500 ms
84 bytes from 172.17.2.20: icmp_seq=2 ttl=56 time=241.201 ms
84 bytes from 172.17.2.20: icmp_seq=3 ttl=56 time=240.331 ms
84 bytes from 172.17.2.20: icmp_seq=4 ttl=56 time=240.333 ms
84 bytes from 172.17.2.20: icmp_seq=5 ttl=56 time=240.104 ms

PC1> ping 172.17.1.10
84 bytes from 172.17.1.10: icmp_seq=1 ttl=64 time=7.366 ms
84 bytes from 172.17.1.10: icmp_seq=2 ttl=64 time=2.813 ms
84 bytes from 172.17.1.10: icmp_seq=3 ttl=64 time=2.288 ms
84 bytes from 172.17.1.10: icmp_seq=4 ttl=64 time=0.775 ms
84 bytes from 172.17.1.10: icmp_seq=5 ttl=64 time=1.217 ms

PC1>

TFG_Proyecto_B_v3 - GNS3
PC2 - PuTTY
PC2> ping 172.17.1.10
84 bytes from 172.17.1.10: icmp_seq=1 ttl=56 time=240.889 ms
84 bytes from 172.17.1.10: icmp_seq=2 ttl=56 time=240.213 ms
84 bytes from 172.17.1.10: icmp_seq=3 ttl=56 time=240.665 ms
84 bytes from 172.17.1.10: icmp_seq=4 ttl=56 time=240.361 ms
84 bytes from 172.17.1.10: icmp_seq=5 ttl=56 time=240.678 ms

PC2> ping 172.17.1.20
172.17.1.20: icmp: seq=1 timeout
172.17.1.20: icmp: seq=2 timeout
84 bytes from 172.17.1.20: icmp_seq=3 ttl=56 time=225.306 ms
84 bytes from 172.17.1.20: icmp_seq=4 ttl=56 time=240.866 ms
84 bytes from 172.17.1.20: icmp_seq=5 ttl=56 time=240.345 ms

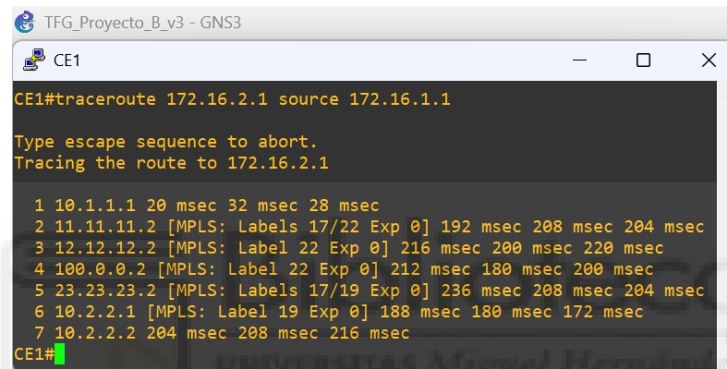
PC2> ping 172.17.2.10
84 bytes from 172.17.2.10: icmp_seq=1 ttl=64 time=1.117 ms
84 bytes from 172.17.2.10: icmp_seq=2 ttl=64 time=2.337 ms
84 bytes from 172.17.2.10: icmp_seq=3 ttl=64 time=1.223 ms
84 bytes from 172.17.2.10: icmp_seq=4 ttl=64 time=1.225 ms
84 bytes from 172.17.2.10: icmp_seq=5 ttl=64 time=1.899 ms

PC2>
```

Figura 39: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario B con tres routers.

La Figura 39 confirma la conectividad entre las dos sedes del cliente en el escenario B con tres routers por sistema autónomo. El ping desde CE1 hacia la loopback remota de CE2 obtiene una tasa de éxito del 100%, y las pruebas desde LinuxCore1, LinuxCore2, PC1 y PC2 también muestran comunicación entre los equipos finales. Los tiempos remotos son superiores a los locales y los posibles timeouts iniciales responden al comportamiento ya explicado en los apartados anteriores.

3.4.2. Trazado de ruta mediante traceroute



```

TFG_Proyecto_B_v3 - GNS3
CE1
CE1#traceroute 172.16.2.1 source 172.16.1.1

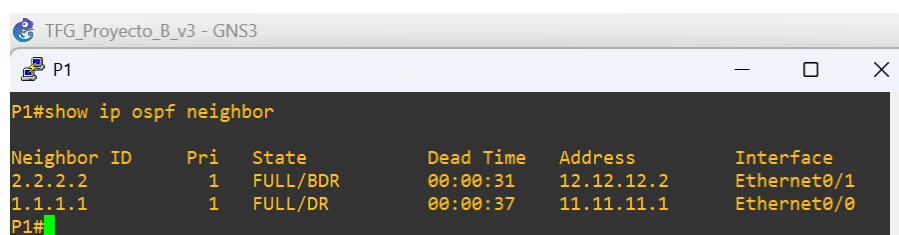
Type escape sequence to abort.
Tracing the route to 172.16.2.1

 0 10.1.1.1 20 msec 32 msec 28 msec
 1 11.11.11.2 [MPLS: Label 17/22 Exp 0] 192 msec 208 msec 204 msec
 2 12.12.12.2 [MPLS: Label 22 Exp 0] 216 msec 200 msec 220 msec
 3 100.0.0.2 [MPLS: Label 22 Exp 0] 212 msec 180 msec 200 msec
 4 23.23.23.2 [MPLS: Label 17/19 Exp 0] 236 msec 208 msec 204 msec
 5 10.2.2.1 [MPLS: Label 19 Exp 0] 188 msec 180 msec 172 msec
 6 10.2.2.2 204 msec 208 msec 216 msec
CE1#
  
```

Figura 40: Trazado de ruta desde CE1 hacia la loopback remota de CE2 en el escenario B con tres routers.

En la Figura 40 se observa el traceroute desde CE1 hacia 172.16.2.1 utilizando como origen 172.16.1.1. A diferencia del escenario B con dos routers, el recorrido incluye más saltos, ya que el tráfico atraviesa PE1, P1, ASBR1, ASBR2, P2, PE2 y finalmente CE2. También aparecen etiquetas MPLS en varios puntos del camino, lo que confirma que el tráfico cruza un núcleo MPLS interno en cada proveedor antes de alcanzar el extremo remoto.

3.4.3. Verificación de vecinos OSPF



```

TFG_Proyecto_B_v3 - GNS3
P1
P1#show ip ospf neighbor

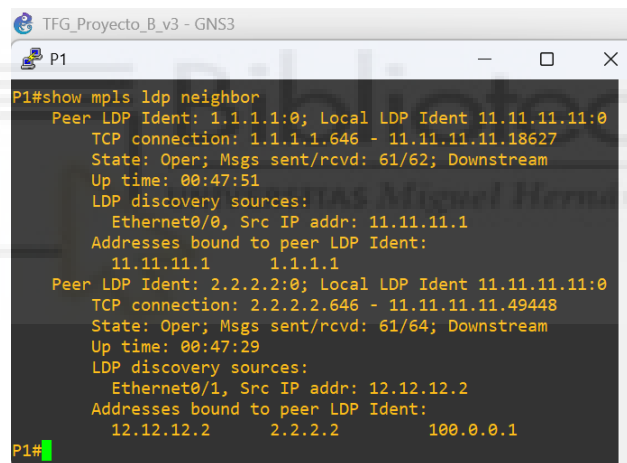
Neighbor ID    Pri   State           Dead Time   Address        Interface
2.2.2.2        1     FULL/BDR        00:00:31    12.12.12.2     Ethernet0/1
1.1.1.1        1     FULL/DR         00:00:37    11.11.11.1     Ethernet0/0
P1#
  
```

Figura 41: Vecinos OSPF en P1 en el escenario B con tres routers.

En la Figura 41 se muestra el comando `show ip ospf neighbor` ejecutado en P1. La salida confirma que P1 mantiene dos vecindades OSPF, una con PE1, identificado por 1.1.1.1, y otra con ASBR1, identificado por 2.2.2.2. Ambas aparecen en estado FULL, por lo que la conectividad interna del proveedor izquierdo está correctamente establecida.

Esta comprobación es importante porque P1 actúa como router de tránsito dentro del núcleo del proveedor. Al tener vecindad con PE1 y ASBR1, permite que ambos extremos internos conozcan sus loopbacks y enlaces de transporte, lo que es necesario para que funcionen correctamente las sesiones MPLS y BGP asociadas al escenario.

3.4.4. Verificación de vecinos LDP



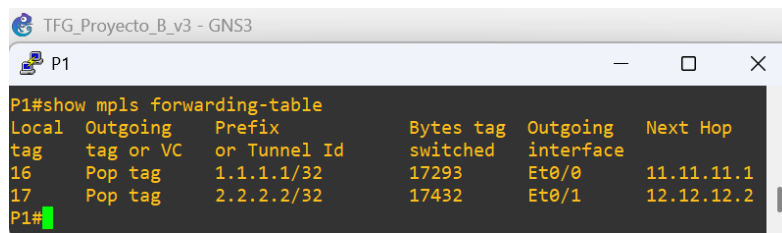
```
P1#show mpls ldp neighbor
Peer LDP Ident: 1.1.1.1:0; Local LDP Ident 11.11.11.11:0
TCP connection: 1.1.1.1.646 - 11.11.11.11.18627
State: Oper; Msgs sent/rcvd: 61/62; Downstream
Up time: 00:47:51
LDP discovery sources:
Ethernet0/0, Src IP addr: 11.11.11.1
Addresses bound to peer LDP Ident:
11.11.11.1 1.1.1.1
Peer LDP Ident: 2.2.2.2:0; Local LDP Ident 11.11.11.11:0
TCP connection: 2.2.2.2.646 - 11.11.11.11.49448
State: Oper; Msgs sent/rcvd: 61/64; Downstream
Up time: 00:47:29
LDP discovery sources:
Ethernet0/1, Src IP addr: 12.12.12.2
Addresses bound to peer LDP Ident:
12.12.12.2 2.2.2.2 100.0.0.1
P1#
```

Figura 42: Vecinos LDP en P1 en el escenario B con tres routers.

La Figura 42 recoge el comando `show mpls ldp neighbor` en P1. En la salida aparecen dos vecinos LDP, uno hacia PE1 con identificador 1.1.1.1:0 y otro hacia ASBR1 con identificador 2.2.2.2:0. En ambos casos el estado aparece como Oper, lo que indica que las sesiones LDP están operativas.

Esta verificación demuestra que el núcleo interno no solo tiene conectividad IP mediante OSPF, sino que también puede intercambiar etiquetas MPLS entre PE1, P1 y ASBR1. Por tanto, P1 queda validado como router de tránsito dentro del transporte MPLS del proveedor.

3.4.5. Verificación de la tabla de reenvío MPLS



The screenshot shows a terminal window titled 'TFG_Proyecto_B_v3 - GNS3' with a tab for 'P1'. The command 'P1#show mpls forwarding-table' has been executed, displaying the following table:

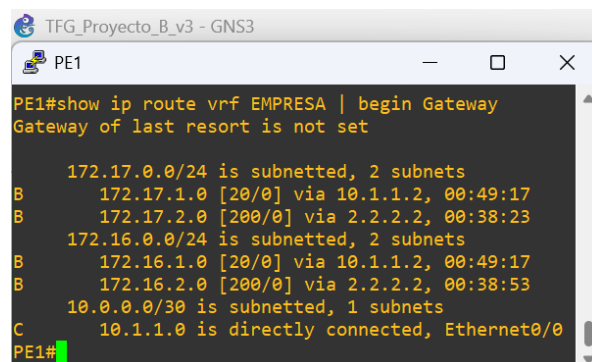
Local tag	Outgoing tag or VC	Prefix or Tunnel Id	Bytes tag switched	Outgoing interface	Next Hop
16	Pop tag	1.1.1.1/32	17293	Et0/0	11.11.11.1
17	Pop tag	2.2.2.2/32	17432	Et0/1	12.12.12.2

Figura 43: Tabla de reenvío MPLS en P1 en el escenario B con tres routers.

En la Figura 43 se muestra `show mpls forwarding-table` en P1. La tabla contiene entradas hacia las loopbacks 1.1.1.1/32 y 2.2.2.2/32, asociadas a PE1 y ASBR1 respectivamente. Para cada destino se indica una etiqueta local, la acción de salida, la interfaz utilizada y el siguiente salto.

La acción `Pop tag` indica que P1 retira la etiqueta antes de entregar el tráfico al siguiente router del tramo MPLS, en el caso que dicho router permite salir del sistema autónomo. Esta salida complementa las verificaciones de OSPF y LDP, ya que confirma que el router P no solo forma vecindades, sino que también dispone de información real de reenvío MPLS.

3.4.6. Verificación de rutas en la VRF



The screenshot shows a terminal window titled 'TFG_Proyecto_B_v3 - GNS3' with a tab for 'PE1'. The command 'PE1#show ip route vrf EMPRESA | begin Gateway' has been executed, displaying the following routing information:

```
Gateway of last resort is not set

 172.17.0.0/24 is subnetted, 2 subnets
B       172.17.1.0 [20/0] via 10.1.1.2, 00:49:17
B       172.17.2.0 [200/0] via 2.2.2.2, 00:38:23
 172.16.0.0/24 is subnetted, 2 subnets
B       172.16.1.0 [20/0] via 10.1.1.2, 00:49:17
B       172.16.2.0 [200/0] via 2.2.2.2, 00:38:53
 10.0.0.0/30 is subnetted, 1 subnets
C       10.1.1.0 is directly connected, Ethernet0/0
PE1#
```

Figura 44: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario B con tres routers.

En la Figura 44 se muestra `show ip route vrf EMPRESA` en PE1. La tabla confirma que las redes locales del cliente, 172.17.1.0/24 y 172.16.1.0/24, se alcanzan mediante

10.1.1.2, que corresponde a CE1. Las redes remotas, 172.17.2.0/24 y 172.16.2.0/24, aparecen a través de 2.2.2.2, que corresponde a ASBR1.

Esta salida verifica que PE1 dispone de las rutas necesarias dentro de la VRF EMPRESA para comunicar la sede local con la sede remota. En este caso, no es necesario repetir la captura en PE2, ya que el lado derecho mantiene un comportamiento equivalente con sus propias direcciones.

3.4.7. Verificación de rutas VPNv4 en los ASBR

```

TFG_Proyecto_B_v3 - GNS3
ASBR1
ASBR1#show ip bgp vpnv4 all
BGP table version is 5, local router ID is 2.2.2.2
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop           Metric LocPrf Weight Path
Route Distinguisher: 65001:1
*>i172.16.1.0/24   1.1.1.1             0      100      0 100 i
*>i172.17.1.0/24   1.1.1.1             0      100      0 100 i
Route Distinguisher: 65002:1
*> 172.16.2.0/24   100.0.0.2           0      100      0 65002 200 i
*> 172.17.2.0/24   100.0.0.2           0      100      0 65002 200 i
ASBR1#

TFG_Proyecto_B_v3 - GNS3
ASBR2
ASBR2#show ip bgp vpnv4 all
BGP table version is 5, local router ID is 3.3.3.3
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop           Metric LocPrf Weight Path
Route Distinguisher: 65001:1
*> 172.16.1.0/24   100.0.0.1           0      100      0 65001 100 i
*> 172.17.1.0/24   100.0.0.1           0      100      0 65001 100 i
Route Distinguisher: 65002:1
*>i172.16.2.0/24   4.4.4.4             0      100      0 200 i
*>i172.17.2.0/24   4.4.4.4             0      100      0 200 i
ASBR2#

```

Figura 45: Tabla BGP VPNv4 en ASBR1 y ASBR2 en el escenario B con tres routers.

En la Figura 45 se comparan las salidas de show ip bgp vpnv4 all en ASBR1 y ASBR2. En ASBR1 aparecen rutas asociadas al Route Distinguisher 65001:1, correspondientes al lado izquierdo, y rutas asociadas al Route Distinguisher 65002:1, correspondientes al lado derecho. Las rutas locales se alcanzan mediante 1.1.1.1, mientras que las rutas remotas se alcanzan mediante 100.0.0.2. En ASBR2 se observa la misma lógica desde el extremo contrario. Las rutas del lado izquierdo se alcanzan mediante 100.0.0.1, mientras que las rutas del lado derecho se alcanzan mediante 4.4.4.4. Esta comparación es clave en el escenario B, ya que confirma que los ASBR intercambian rutas VPNv4 entre sistemas autónomos y no se limitan únicamente al transporte interno.

Los resultados obtenidos son coherentes con el diseño del escenario B con tres routers. La conectividad entre sedes funciona correctamente, el núcleo interno se valida mediante OSPF, LDP y MPLS, y los ASBR muestran las rutas VPNv4 necesarias para el intercambio entre proveedores. La diferencia principal respecto al escenario A con tres routers es que, en este caso, los ASBR participan directamente en el intercambio VPNv4 entre sistemas autónomos, mientras que el escenario A se apoya en la continuidad de la VRF en la frontera entre proveedores.

3.5 Escenario C con dos routers

3.5.1. Pruebas de conectividad mediante ping

```
CE1#ping 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
Packet sent with a source address of 172.16.1.1
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 140/146/152 ms
CE1#

LinuxCore-1 - PuTTY
tc@box:~$ ping 172.17.2.10
PING 172.17.2.10 (172.17.2.10): 56 data bytes
64 bytes from 172.17.2.10: seq=0 ttl=58 time=207.767 ms
64 bytes from 172.17.2.10: seq=1 ttl=58 time=168.571 ms
64 bytes from 172.17.2.10: seq=2 ttl=58 time=183.071 ms
64 bytes from 172.17.2.10: seq=3 ttl=58 time=176.576 ms
64 bytes from 172.17.2.10: seq=4 ttl=58 time=174.847 ms
--- 172.17.2.10 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 168.571/182.166/207.767 ms
tc@box:~$
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=1 ttl=58 time=180.442 ms
64 bytes from 172.17.2.20: seq=2 ttl=58 time=181.549 ms
64 bytes from 172.17.2.20: seq=3 ttl=58 time=169.355 ms
64 bytes from 172.17.2.20: seq=4 ttl=58 time=176.127 ms
64 bytes from 172.17.2.20: seq=5 ttl=58 time=178.129 ms
--- 172.17.2.20 ping statistics ---
6 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 169.355/175.520/181.549 ms
tc@box:~$
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=0 ttl=64 time=19.693 ms
64 bytes from 172.17.1.20: seq=1 ttl=64 time=9.684 ms
64 bytes from 172.17.1.20: seq=2 ttl=64 time=9.917 ms
64 bytes from 172.17.1.20: seq=3 ttl=64 time=1.451 ms
64 bytes from 172.17.1.20: seq=4 ttl=64 time=9.396 ms
--- 172.17.1.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 0.684/7.857/19.693 ms
tc@box:~$

LinuxCore-2 - PuTTY
tc@box:~$ ping 172.17.1.10
PING 172.17.1.10 (172.17.1.10): 56 data bytes
64 bytes from 172.17.1.10: seq=0 ttl=58 time=172.660 ms
64 bytes from 172.17.1.10: seq=1 ttl=58 time=178.138 ms
64 bytes from 172.17.1.10: seq=2 ttl=58 time=182.407 ms
64 bytes from 172.17.1.10: seq=3 ttl=58 time=181.671 ms
64 bytes from 172.17.1.10: seq=4 ttl=58 time=175.899 ms
--- 172.17.1.10 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 172.660/178.155/182.407 ms
tc@box:~$
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=1 ttl=58 time=172.315 ms
64 bytes from 172.17.1.20: seq=2 ttl=58 time=176.174 ms
64 bytes from 172.17.1.20: seq=3 ttl=58 time=178.898 ms
64 bytes from 172.17.1.20: seq=4 ttl=58 time=174.410 ms
64 bytes from 172.17.1.20: seq=5 ttl=58 time=177.605 ms
--- 172.17.1.20 ping statistics ---
6 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 172.315/175.880/178.898 ms
tc@box:~$
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=0 ttl=64 time=9.238 ms
64 bytes from 172.17.2.20: seq=1 ttl=64 time=0.814 ms
64 bytes from 172.17.2.20: seq=2 ttl=64 time=0.917 ms
64 bytes from 172.17.2.20: seq=3 ttl=64 time=0.772 ms
64 bytes from 172.17.2.20: seq=4 ttl=64 time=2.017 ms
--- 172.17.2.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 0.772/2.753/9.238 ms
tc@box:~$

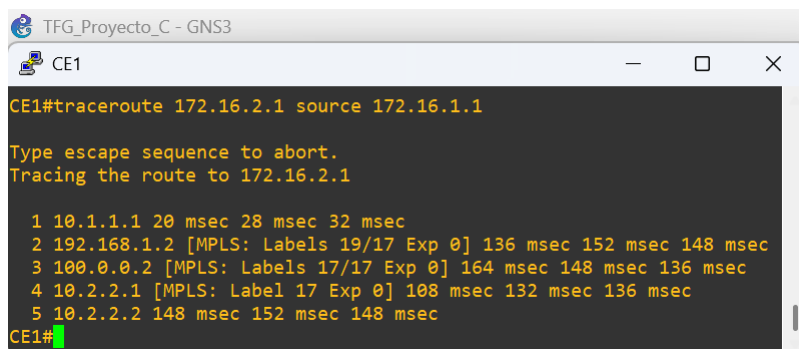
PC1 - PuTTY
PC1> ping 172.17.2.10
64 bytes from 172.17.2.10 icmp_seq=1 ttl=58 time=180.307 ms
64 bytes from 172.17.2.10 icmp_seq=2 ttl=58 time=180.765 ms
64 bytes from 172.17.2.10 icmp_seq=3 ttl=58 time=180.325 ms
64 bytes from 172.17.2.10 icmp_seq=4 ttl=58 time=180.757 ms
64 bytes from 172.17.2.10 icmp_seq=5 ttl=58 time=180.195 ms
PC1> ping 172.17.2.20
172.17.2.20 icmp_seq=2 timeout
64 bytes from 172.17.2.20 icmp_seq=3 ttl=58 time=180.476 ms
64 bytes from 172.17.2.20 icmp_seq=4 ttl=58 time=180.694 ms
64 bytes from 172.17.2.20 icmp_seq=5 ttl=58 time=180.756 ms
PC1> ping 172.17.1.10
64 bytes from 172.17.1.10 icmp_seq=1 ttl=64 time=2.337 ms
64 bytes from 172.17.1.10 icmp_seq=2 ttl=64 time=0.898 ms
64 bytes from 172.17.1.10 icmp_seq=3 ttl=64 time=2.177 ms
64 bytes from 172.17.1.10 icmp_seq=4 ttl=64 time=2.298 ms
64 bytes from 172.17.1.10 icmp_seq=5 ttl=64 time=10.550 ms
PC1>

PC2 - PuTTY
PC2> ping 172.17.1.10
64 bytes from 172.17.1.10 icmp_seq=1 ttl=58 time=180.134 ms
64 bytes from 172.17.1.10 icmp_seq=2 ttl=58 time=180.683 ms
64 bytes from 172.17.1.10 icmp_seq=3 ttl=58 time=180.285 ms
64 bytes from 172.17.1.10 icmp_seq=4 ttl=58 time=180.634 ms
64 bytes from 172.17.1.10 icmp_seq=5 ttl=58 time=180.366 ms
PC2> ping 172.17.1.20
64 bytes from 172.17.1.20 icmp_seq=1 ttl=58 time=180.770 ms
64 bytes from 172.17.1.20 icmp_seq=2 ttl=58 time=165.284 ms
64 bytes from 172.17.1.20 icmp_seq=3 ttl=58 time=180.546 ms
64 bytes from 172.17.1.20 icmp_seq=4 ttl=58 time=180.077 ms
64 bytes from 172.17.1.20 icmp_seq=5 ttl=58 time=165.181 ms
PC2> ping 172.17.2.10
64 bytes from 172.17.2.10 icmp_seq=1 ttl=64 time=1.781 ms
64 bytes from 172.17.2.10 icmp_seq=2 ttl=64 time=1.123 ms
64 bytes from 172.17.2.10 icmp_seq=3 ttl=64 time=2.088 ms
64 bytes from 172.17.2.10 icmp_seq=4 ttl=64 time=1.172 ms
64 bytes from 172.17.2.10 icmp_seq=5 ttl=64 time=3.198 ms
PC2>
```

Figura 46: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario C con dos routers.

La Figura 46 confirma que el escenario C con dos routers permite la comunicación entre las dos sedes del cliente. El ping desde CE1 hacia la loopback remota de CE2 alcanza una tasa de éxito del 100%, y las pruebas desde LinuxCore1, LinuxCore2, PC1 y PC2 también muestran conectividad entre equipos locales y remotos.

3.5.2. Trazado de ruta mediante traceroute



```

TFG_Proyecto_C - GNS3
CE1
CE1#traceroute 172.16.2.1 source 172.16.1.1

Type escape sequence to abort.
Tracing the route to 172.16.2.1

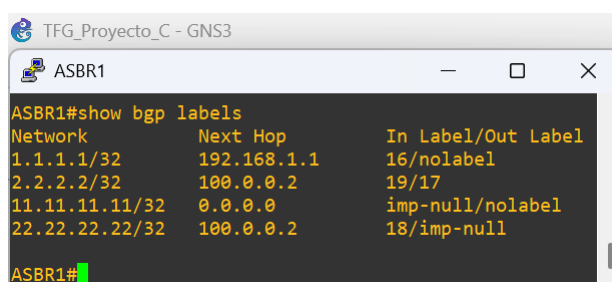
 0 10.1.1.1 20 msec 28 msec 32 msec
 1 192.168.1.2 [MPLS: Labels 19/17 Exp 0] 136 msec 152 msec 148 msec
 2 100.0.0.2 [MPLS: Labels 17/17 Exp 0] 164 msec 148 msec 136 msec
 3 10.2.2.1 [MPLS: Label 17 Exp 0] 108 msec 132 msec 136 msec
 4 10.2.2.2 148 msec 152 msec 148 msec
CE1#

```

Figura 47: Trazado de ruta desde CE1 hacia la loopback remota de CE2 en el escenario C con dos routers.

En la Figura 47 se observa el traceroute desde CE1 hacia 172.16.2.1 utilizando como origen 172.16.1.1. El recorrido pasa por PE1, ASBR1, ASBR2, PE2 y finalmente CE2. Durante el trayecto aparecen etiquetas MPLS en los saltos intermedios, lo que confirma que el tráfico atraviesa un transporte etiquetado. La diferencia respecto al escenario B no está tanto en el camino físico, que puede parecer el mismo, sino en el plano de control, ya que en este escenario los ASBR no intercambian rutas VPNv4, sino rutas IPv4 etiquetadas para permitir la comunicación directa VPNv4 entre los PE.

3.5.3. Verificación de rutas IPv4 etiquetadas en BGP



```

TFG_Proyecto_C - GNS3
ASBR1
ASBR1#show bgp labels
Network          Next Hop          In Label/Out Label
1.1.1.1/32       192.168.1.1       16/nolabel
2.2.2.2/32       100.0.0.2         19/17
11.11.11.11/32   0.0.0.0           imp-null/nolabel
22.22.22.22/32   100.0.0.2         18/imp-null
ASBR1#

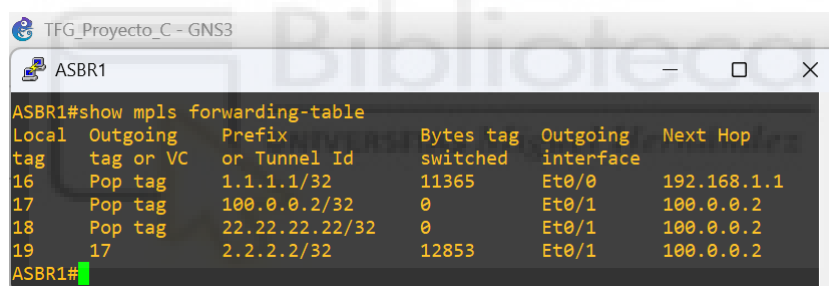
```

Figura 48: Rutas IPv4 etiquetadas en ASBR1 mediante BGP en el escenario C con dos routers.

En la Figura 48 se muestra el comando `show bgp labels` ejecutado en ASBR1. Esta salida permite comprobar las rutas IPv4 etiquetadas que utiliza el escenario C para construir el transporte entre los PE. Aparecen prefijos como 1.1.1.1/32, 2.2.2.2/32, 11.11.11.11/32 y 22.22.22.22/32, junto con su siguiente salto y las etiquetas de entrada y salida asociadas.

Este punto es clave en el escenario C, porque ASBR1 no aparece aquí como intercambiador de rutas VPNv4, sino como router encargado de transportar rutas IPv4 etiquetadas. Aunque también se puede utilizar el comando `show ip bgp labels`, en esta captura se utiliza `show bgp labels`, que cumple la misma función de verificación en esta sintaxis de IOS. Por tanto, también podría usarse la variante con `ip` si la versión del router la admite.

3.5.4. Verificación de la tabla de reenvío MPLS



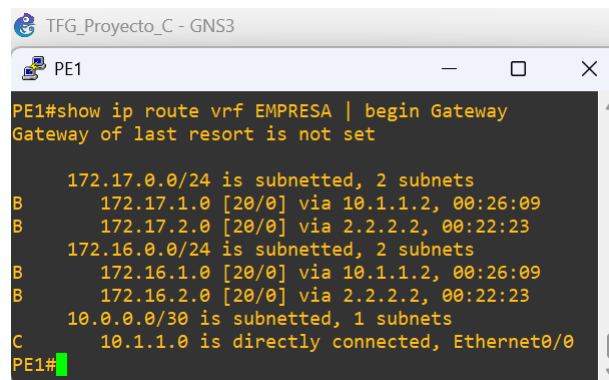
Local tag	Outgoing tag or VC	Prefix or Tunnel Id	Bytes switched	Outgoing interface	Next Hop
16	Pop tag	1.1.1.1/32	11365	Eth0/0	192.168.1.1
17	Pop tag	100.0.0.2/32	0	Eth0/1	100.0.0.2
18	Pop tag	22.22.22.22/32	0	Eth0/1	100.0.0.2
19	17	2.2.2.2/32	12853	Eth0/1	100.0.0.2

Figura 49: Tabla de reenvío MPLS en ASBR1 en el escenario C con dos routers.

La Figura 49 recoge el comando `show mpls forwarding-table` en ASBR1. La tabla muestra entradas MPLS hacia prefijos de transporte como 1.1.1.1/32, 2.2.2.2/32, 22.22.22.22/32 y 100.0.0.2/32. En cada caso se indica la etiqueta local, la acción de salida, la interfaz utilizada y el siguiente salto.

Esta salida complementa la figura anterior, ya que no solo se comprueba que BGP conoce rutas etiquetadas, sino que ASBR1 también dispone de información real de reenvío MPLS. De esta forma, queda validado el transporte etiquetado que permite que los PE puedan intercambiar información VPNv4 directamente entre sistemas autónomos.

3.5.5. Verificación de rutas en la VRF



```
TFG_Proyecto_C - GNS3
PE1
PE1#show ip route vrf EMPRESA | begin Gateway
Gateway of last resort is not set

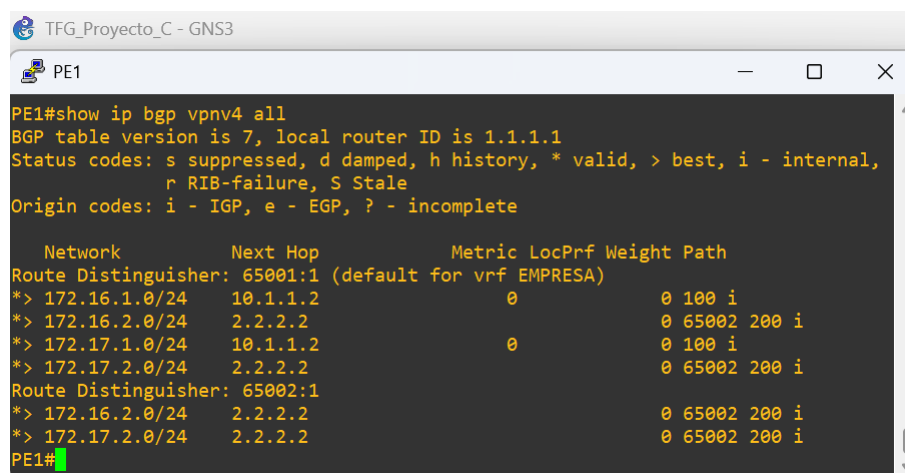
    172.17.0.0/24 is subnetted, 2 subnets
B       172.17.1.0 [20/0] via 10.1.1.2, 00:26:09
B       172.17.2.0 [20/0] via 2.2.2.2, 00:22:23
    172.16.0.0/24 is subnetted, 2 subnets
B       172.16.1.0 [20/0] via 10.1.1.2, 00:26:09
B       172.16.2.0 [20/0] via 2.2.2.2, 00:22:23
    10.0.0.0/30 is subnetted, 1 subnets
C       10.1.1.0 is directly connected, Ethernet0/0
PE1#
```

Figura 50: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario C con dos routers.

En la Figura 50 se muestra `show ip route vrf EMPRESA` en PE1. La tabla confirma que las redes locales 172.17.1.0/24 y 172.16.1.0/24 se alcanzan mediante 10.1.1.2, correspondiente a CE1. Las redes remotas 172.17.2.0/24 y 172.16.2.0/24 aparecen a través de 2.2.2.2, que corresponde al PE remoto.

Este resultado es importante porque muestra una diferencia clara respecto al escenario B. En el escenario B, PE1 alcanzaba las rutas remotas a través de ASBR1, mientras que aquí las rutas remotas apuntan directamente hacia el PE del otro sistema autónomo. Esto encaja con el funcionamiento del escenario C, donde el intercambio VPNv4 se realiza entre PE mediante eBGP multihop.

3.5.6. Verificación de rutas VPNv4 en PE1



```
TFG_Proyecto_C - GNS3
PE1
PE1#show ip bgp vpnv4 all
BGP table version is 7, local router ID is 1.1.1.1
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network          Next Hop          Metric LocPrf Weight Path
Route Distinguisher: 65001:1 (default for vrf EMPRESA)
*> 172.16.1.0/24    10.1.1.2              0           0 100 i
*> 172.16.2.0/24    2.2.2.2              0           0 65002 200 i
*> 172.17.1.0/24    10.1.1.2              0           0 100 i
*> 172.17.2.0/24    2.2.2.2              0           0 65002 200 i
Route Distinguisher: 65002:1
*> 172.16.2.0/24    2.2.2.2              0           0 65002 200 i
*> 172.17.2.0/24    2.2.2.2              0           0 65002 200 i
PE1#
```

Figura 51: Tabla BGP VPNv4 en PE1 en el escenario C con dos routers.

En la Figura 51 se presenta la tabla BGP VPNv4 de PE1 mediante show ip bgp vpnv4 all. En ella aparecen los prefijos del cliente asociados a la VRF EMPRESA. Las rutas locales, como 172.16.1.0/24 y 172.17.1.0/24, se alcanzan mediante 10.1.1.2, mientras que las rutas remotas, como 172.16.2.0/24 y 172.17.2.0/24, se alcanzan mediante 2.2.2.2.

También aparecen rutas asociadas al Route Distinguisher del extremo remoto, lo que confirma que PE1 está recibiendo información VPNv4 desde el otro sistema autónomo. Esta captura refuerza la idea principal del escenario C, ya que la información VPNv4 no se intercambia en los ASBR, sino directamente entre PE1 y PE2.

3.5.7. Resumen de vecinos BGP VPNv4

```

PE1#show ip bgp vpnv4 all summary
BGP router identifier 1.1.1.1, local AS number 65001
BGP table version is 7, main routing table version 7
6 network entries using 822 bytes of memory
6 path entries using 408 bytes of memory
8/2 BGP path/bestpath attribute entries using 992 bytes of memory
3 BGP AS-PATH entries using 72 bytes of memory
2 BGP extended community entries using 48 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
BGP using 2342 total bytes of memory
BGP activity 10/0 prefixes, 10/0 paths, scan interval 15 secs

Neighbor      V      AS MsgRcvd MsgSent  TblVer  InQ OutQ Up/Down  State/PfxRcd
2.2.2.2        4 65002    23     23      7    0    0 00:18:38      2
10.1.1.2       4   100    26     26      7    0    0 00:22:36      2
PE1#
  
```

Figura 52: Resumen de vecinos BGP VPNv4 en PE1 en el escenario C con dos routers.

En la Figura 52 se muestra show ip bgp vpnv4 all summary en PE1. La salida confirma que PE1 mantiene una sesión BGP VPNv4 con 2.2.2.2, correspondiente al PE remoto en el AS 65002, y otra relación asociada al cliente mediante 10.1.1.2 en el AS 100. En ambos casos se reciben dos prefijos, lo que indica que las sesiones están establecidas y que PE1 dispone de las rutas necesarias.

Los resultados obtenidos son correctos y muestran la diferencia principal frente a los escenarios anteriores. El escenario A se apoyaba en la continuidad de la VRF en la frontera entre proveedores, el escenario B hacía que los ASBR intercambiaran rutas VPNv4, y en este escenario C los ASBR se limitan al transporte IPv4 etiquetado mientras los PE intercambian directamente la información VPNv4. Por tanto, las capturas validan tanto la conectividad como el comportamiento específico del escenario C con dos routers.

3.6 Escenario C con tres routers

3.6.1. Pruebas de conectividad mediante ping

The image displays five terminal windows from a GNS3 simulation, each showing the results of a ping command. The windows are titled 'CE1', 'LinuxCore-1 - PuTTY', 'LinuxCore-2 - PuTTY', 'PC1 - PuTTY', and 'PC2 - PuTTY'. Each window shows the command being executed and the resulting output, including packet statistics and round-trip times.

```
CE1#ping 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
Packet sent with a source address of 172.16.1.1
!!!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 184/202/224 ms
CE1#
```

```
LinuxCore-1 - PuTTY
tc@box:~$ ping 172.17.2.10
PING 172.17.2.10 (172.17.2.10): 56 data bytes
64 bytes from 172.17.2.10: seq=1 ttl=56 time=221.104 ms
64 bytes from 172.17.2.10: seq=2 ttl=56 time=238.322 ms
64 bytes from 172.17.2.10: seq=3 ttl=56 time=255.644 ms
64 bytes from 172.17.2.10: seq=4 ttl=56 time=242.128 ms
64 bytes from 172.17.2.10: seq=5 ttl=56 time=232.192 ms
--- 172.17.2.10 ping statistics ---
5 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 221.104/237.878/255.644 ms
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=1 ttl=56 time=243.434 ms
64 bytes from 172.17.2.20: seq=2 ttl=56 time=237.977 ms
64 bytes from 172.17.2.20: seq=3 ttl=56 time=230.945 ms
64 bytes from 172.17.2.20: seq=4 ttl=56 time=242.052 ms
64 bytes from 172.17.2.20: seq=5 ttl=56 time=240.481 ms
--- 172.17.2.20 ping statistics ---
5 packets transmitted, 5 packets received, 16% packet loss
round-trip min/avg/max = 230.945/238.977/243.434 ms
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=0 ttl=64 time=0.312 ms
64 bytes from 172.17.1.20: seq=1 ttl=64 time=0.766 ms
64 bytes from 172.17.1.20: seq=2 ttl=64 time=0.822 ms
64 bytes from 172.17.1.20: seq=3 ttl=64 time=0.766 ms
64 bytes from 172.17.1.20: seq=4 ttl=64 time=0.706 ms
--- 172.17.1.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 0.706/0.814/0.822 ms
tc@box:~$
```

```
LinuxCore-2 - PuTTY
tc@box:~$ ping 172.17.1.10
PING 172.17.1.10 (172.17.1.10): 56 data bytes
64 bytes from 172.17.1.10: seq=0 ttl=56 time=233.375 ms
64 bytes from 172.17.1.10: seq=1 ttl=56 time=232.595 ms
64 bytes from 172.17.1.10: seq=2 ttl=56 time=233.505 ms
64 bytes from 172.17.1.10: seq=3 ttl=56 time=241.023 ms
64 bytes from 172.17.1.10: seq=4 ttl=56 time=232.399 ms
--- 172.17.1.10 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 232.399/234.579/241.023 ms
tc@box:~$ ping 172.17.1.20
PING 172.17.1.20 (172.17.1.20): 56 data bytes
64 bytes from 172.17.1.20: seq=0 ttl=56 time=3244.982 ms
64 bytes from 172.17.1.20: seq=1 ttl=56 time=2253.471 ms
64 bytes from 172.17.1.20: seq=2 ttl=56 time=1255.463 ms
64 bytes from 172.17.1.20: seq=3 ttl=56 time=246.132 ms
64 bytes from 172.17.1.20: seq=4 ttl=56 time=235.078 ms
--- 172.17.1.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 235.078/1447.025/3244.982 ms
tc@box:~$ ping 172.17.2.20
PING 172.17.2.20 (172.17.2.20): 56 data bytes
64 bytes from 172.17.2.20: seq=0 ttl=64 time=0.879 ms
64 bytes from 172.17.2.20: seq=1 ttl=64 time=0.813 ms
64 bytes from 172.17.2.20: seq=2 ttl=64 time=1.127 ms
64 bytes from 172.17.2.20: seq=3 ttl=64 time=0.642 ms
64 bytes from 172.17.2.20: seq=4 ttl=64 time=0.682 ms
--- 172.17.2.20 ping statistics ---
5 packets transmitted, 5 packets received, 0% packet loss
round-trip min/avg/max = 0.642/1.028/2.127 ms
tc@box:~$
```

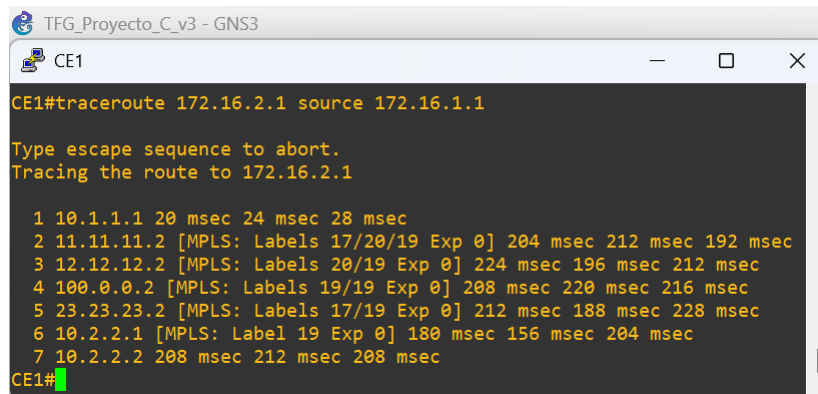
```
PC1 - PuTTY
PC1> ping 172.17.2.10
64 bytes from 172.17.2.10: icmp_seq=1 ttl=56 time=240.292 ms
64 bytes from 172.17.2.10: icmp_seq=2 ttl=56 time=240.631 ms
64 bytes from 172.17.2.10: icmp_seq=3 ttl=56 time=239.980 ms
64 bytes from 172.17.2.10: icmp_seq=4 ttl=56 time=240.252 ms
64 bytes from 172.17.2.10: icmp_seq=5 ttl=56 time=240.427 ms
PC1> ping 172.17.2.20
172.17.2.20: icmp_seq=1 timeout
172.17.2.20: icmp_seq=2 timeout
64 bytes from 172.17.2.20: icmp_seq=3 ttl=56 time=240.221 ms
64 bytes from 172.17.2.20: icmp_seq=4 ttl=56 time=240.380 ms
64 bytes from 172.17.2.20: icmp_seq=5 ttl=56 time=240.872 ms
PC1> ping 172.17.1.10
64 bytes from 172.17.1.10: icmp_seq=1 ttl=64 time=2.965 ms
64 bytes from 172.17.1.10: icmp_seq=2 ttl=64 time=1.239 ms
64 bytes from 172.17.1.10: icmp_seq=3 ttl=64 time=1.650 ms
64 bytes from 172.17.1.10: icmp_seq=4 ttl=64 time=1.748 ms
64 bytes from 172.17.1.10: icmp_seq=5 ttl=64 time=1.022 ms
PC1>
```

```
PC2 - PuTTY
PC2> ping 172.17.1.10
64 bytes from 172.17.1.10: icmp_seq=1 ttl=56 time=240.410 ms
64 bytes from 172.17.1.10: icmp_seq=2 ttl=56 time=240.333 ms
64 bytes from 172.17.1.10: icmp_seq=3 ttl=56 time=240.195 ms
64 bytes from 172.17.1.10: icmp_seq=4 ttl=56 time=240.134 ms
64 bytes from 172.17.1.10: icmp_seq=5 ttl=56 time=240.786 ms
PC2> ping 172.17.1.20
64 bytes from 172.17.1.20: icmp_seq=1 ttl=56 time=240.255 ms
64 bytes from 172.17.1.20: icmp_seq=2 ttl=56 time=240.522 ms
64 bytes from 172.17.1.20: icmp_seq=3 ttl=56 time=240.726 ms
64 bytes from 172.17.1.20: icmp_seq=4 ttl=56 time=240.963 ms
64 bytes from 172.17.1.20: icmp_seq=5 ttl=56 time=240.423 ms
PC2> ping 172.17.2.10
64 bytes from 172.17.2.10: icmp_seq=1 ttl=64 time=7.105 ms
64 bytes from 172.17.2.10: icmp_seq=2 ttl=64 time=2.884 ms
64 bytes from 172.17.2.10: icmp_seq=3 ttl=64 time=3.121 ms
64 bytes from 172.17.2.10: icmp_seq=4 ttl=64 time=0.961 ms
64 bytes from 172.17.2.10: icmp_seq=5 ttl=64 time=5.330 ms
PC2>
```

Figura 53: Pruebas de conectividad desde CE1, LinuxCore1, LinuxCore2, PC1 y PC2 en el escenario C con tres routers.

La Figura 53 confirma la conectividad entre las dos sedes del cliente en el escenario C con tres routers por sistema autónomo. El ping desde CE1 hacia la loopback remota de CE2 obtiene una tasa de éxito del 100%, y las pruebas desde LinuxCore1, LinuxCore2, PC1 y PC2 muestran comunicación entre equipos locales y remotos. Los tiempos y pequeños comportamientos iniciales siguen la misma lógica ya comentada en los apartados anteriores, por lo que no se vuelven a desarrollar en detalle.

3.6.2. Trazado de ruta mediante traceroute



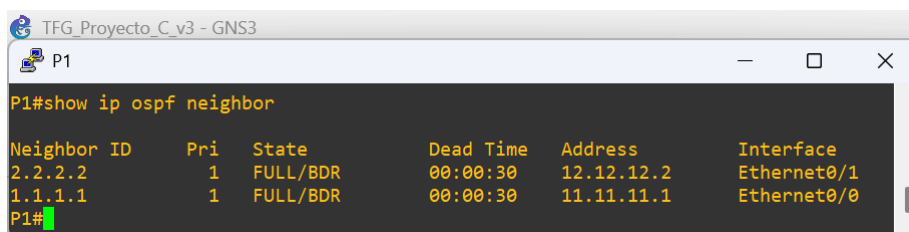
```
TFG_Proyecto_C_v3 - GNS3
CE1
CE1#traceroute 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Tracing the route to 172.16.2.1

 0 10.1.1.1 20 msec 24 msec 28 msec
 1 11.11.11.2 [MPLS: Labels 17/20/19 Exp 0] 204 msec 212 msec 192 msec
 2 12.12.12.2 [MPLS: Labels 20/19 Exp 0] 224 msec 196 msec 212 msec
 3 100.0.0.2 [MPLS: Labels 19/19 Exp 0] 208 msec 220 msec 216 msec
 4 23.23.23.2 [MPLS: Labels 17/19 Exp 0] 212 msec 188 msec 228 msec
 5 10.2.2.1 [MPLS: Label 19 Exp 0] 180 msec 156 msec 204 msec
 6 10.2.2.2 208 msec 212 msec 208 msec
CE1#
```

Figura 54: Trazado de ruta desde CE1 hacia la loopback remota de CE2 en el escenario C con tres routers.

En la Figura 54 se muestra el traceroute desde CE1 hacia 172.16.2.1 utilizando como origen 172.16.1.1. El recorrido atraviesa PE1, P1, ASBR1, ASBR2, P2, PE2 y finalmente CE2. La aparición de etiquetas MPLS en varios saltos confirma que el tráfico cruza un núcleo etiquetado dentro de cada proveedor. La diferencia principal respecto al escenario C con dos routers es el aumento de saltos intermedios debido a la incorporación de P1 y P2.

3.6.3. Verificación de vecinos OSPF



```
TFG_Proyecto_C_v3 - GNS3
P1
P1#show ip ospf neighbor
Neighbor ID    Pri   State           Dead Time   Address        Interface
2.2.2.2        1     FULL/BDR        00:00:30    12.12.12.2     Ethernet0/1
1.1.1.1        1     FULL/BDR        00:00:30    11.11.11.1     Ethernet0/0
P1#
```

Figura 55: Vecinos OSPF en P1 en el escenario C con tres routers.

En la Figura 55 se recoge el comando `show ip ospf neighbor` ejecutado en P1. La salida muestra dos vecinos OSPF, PE1 con router ID 1.1.1.1 y ASBR1 con router ID 2.2.2.2. Ambos aparecen en estado FULL, lo que confirma que P1 mantiene correctamente la conectividad interna con los dos routers del proveedor izquierdo.

Esta comprobación valida el papel de P1 como router de tránsito dentro del núcleo. Al igual que en los escenarios anteriores con tres routers, OSPF permite que las loopbacks y enlaces internos sean alcanzables, lo que resulta necesario para el funcionamiento posterior de LDP, MPLS y BGP.

3.6.4. Verificación de vecinos LDP

```

P1#show mpls ldp neighbor
Peer LDP Ident: 1.1.1.1:0; Local LDP Ident 11.11.11.11:0
TCP connection: 1.1.1.1.646 - 11.11.11.11.59028
State: Oper; Msgs sent/rcvd: 39/39; Downstream
Up time: 00:27:38
LDP discovery sources:
 Ethernet0/0, Src IP addr: 11.11.11.1
Addresses bound to peer LDP Ident:
 11.11.11.1 1.1.1.1
Peer LDP Ident: 2.2.2.2:0; Local LDP Ident 11.11.11.11:0
TCP connection: 2.2.2.2.646 - 11.11.11.11.38794
State: Oper; Msgs sent/rcvd: 39/41; Downstream
Up time: 00:27:20
LDP discovery sources:
 Ethernet0/1, Src IP addr: 12.12.12.2
Addresses bound to peer LDP Ident:
 12.12.12.2 2.2.2.2 100.0.0.1
P1#
  
```

Figura 56: Vecinos LDP en P1 en el escenario C con tres routers.

La Figura 56 muestra show mpls ldp neighbor en P1. En la salida aparecen dos sesiones LDP operativas, una con PE1 mediante el identificador 1.1.1.1:0 y otra con ASBR1 mediante 2.2.2.2:0. En ambos casos el estado es Oper, por lo que el intercambio de etiquetas dentro del núcleo izquierdo está activo.

Esta salida confirma que el núcleo interno no solo tiene conectividad IP, sino que también dispone de las sesiones LDP necesarias para construir el transporte MPLS entre PE1, P1 y ASBR1.

3.6.5. Verificación de la tabla de reenvío MPLS en P1

```

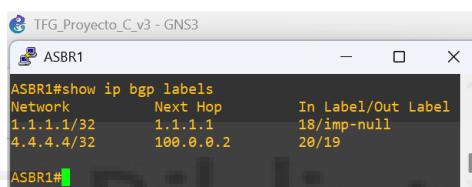
ASBR1#show mpls forwarding-table
Local  Outgoing  Prefix          Bytes tag  Outgoing     Next Hop
tag    tag or VC  or Tunnel Id    switched   interface
16     Pop tag    100.0.0.2/32    0          Et0/1        100.0.0.2
17     Pop tag    11.11.11.0/30   0          Et0/0        12.12.12.1
18     16         1.1.1.1/32      13265      Et0/0        12.12.12.1
19     Pop tag    11.11.11.11/32  0          Et0/0        12.12.12.1
20     19         4.4.4.4/32      13169      Et0/1        100.0.0.2
ASBR1#
  
```

Figura 57: Tabla de reenvío MPLS en P1 en el escenario C con tres routers.

En la Figura 57 se observa show mpls forwarding-table en P1. La tabla contiene entradas hacia las loopbacks 1.1.1.1/32 y 2.2.2.2/32, asociadas a PE1 y ASBR1. Para cada una se muestra la etiqueta local, la acción de salida, la interfaz utilizada y el siguiente salto.

La acción Pop tag indica que P1 retira la etiqueta en esos trayectos antes de reenviar el tráfico al siguiente router, si dicho router permite salir dentro del sistema autónomo. Esta comprobación completa la validación del núcleo izquierdo, ya que P1 tiene vecindades OSPF, sesiones LDP y entradas reales de reenvío MPLS.

3.6.6. Verificación de rutas IPv4 etiquetadas en BGP



```

ASBR1#show ip bgp labels
Network      Next Hop      In Label/Out Label
1.1.1.1/32   1.1.1.1       18/imp-null
4.4.4.4/32   100.0.0.2     20/19
ASBR1#

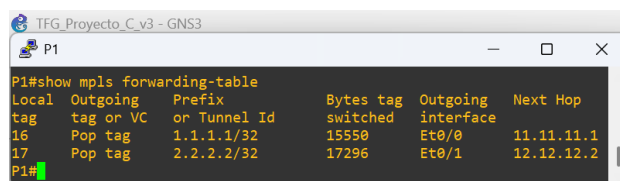
```

Figura 58: Rutas IPv4 etiquetadas en ASBR1 mediante BGP en el escenario C con tres routers.

En la Figura 58 se muestra el comando show ip bgp labels en ASBR1. La salida contiene las rutas IPv4 etiquetadas necesarias para el transporte entre PE1 y PE2, destacando las loopbacks 1.1.1.1/32 y 4.4.4.4/32. Esto es especialmente importante en el escenario C, ya que los ASBR no intercambian las rutas VPNv4 del cliente, sino rutas IPv4 etiquetadas que permiten alcanzar las loopbacks de los PE.

Esta salida cumple la misma función que show bgp labels, utilizado en el apartado 3.5. Lo he querido implementar para que se pueda observar que no hay diferencias de utilizar un comando respecto al otro.

3.6.7. Verificación de la tabla de reenvío MPLS en ASBR1



```

P1#show mpls forwarding-table
Local  Outgoing  Prefix      Bytes tag  Outgoing     Next Hop
tag    tag or VC  or Tunnel Id  switched   interface
16     Pop tag    1.1.1.1/32   15550      Et0/0        11.11.11.1
17     Pop tag    2.2.2.2/32   17296      Et0/1        12.12.12.2
P1#

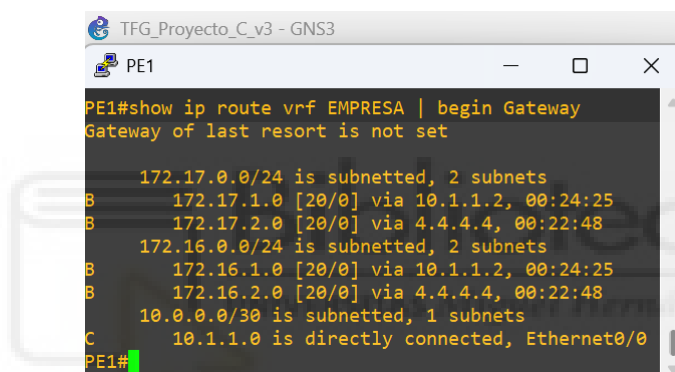
```

Figura 59: Tabla de reenvío MPLS en ASBR1 en el escenario C con tres routers.

En la Figura 59 aparece show mpls forwarding-table en ASBR1. La tabla muestra entradas hacia prefijos de transporte como 100.0.0.2/32, 1.1.1.1/32, 11.11.11.11/32 y 4.4.4.4/32. Para cada destino se indica la etiqueta local, la etiqueta de salida, la interfaz de salida y el siguiente salto.

Esta comprobación complementa la anterior, porque no solo se observa que BGP conoce rutas etiquetadas, sino que ASBR1 también tiene información MPLS instalada para reenviar el tráfico. De esta forma se valida la función principal de ASBR1 en el escenario C, que consiste en sostener el transporte etiquetado entre sistemas autónomos.

3.6.8. Verificación de rutas en la VRF



```
TFG_Proyecto_C_v3 - GNS3
PE1
PE1#show ip route vrf EMPRESA | begin Gateway
Gateway of last resort is not set

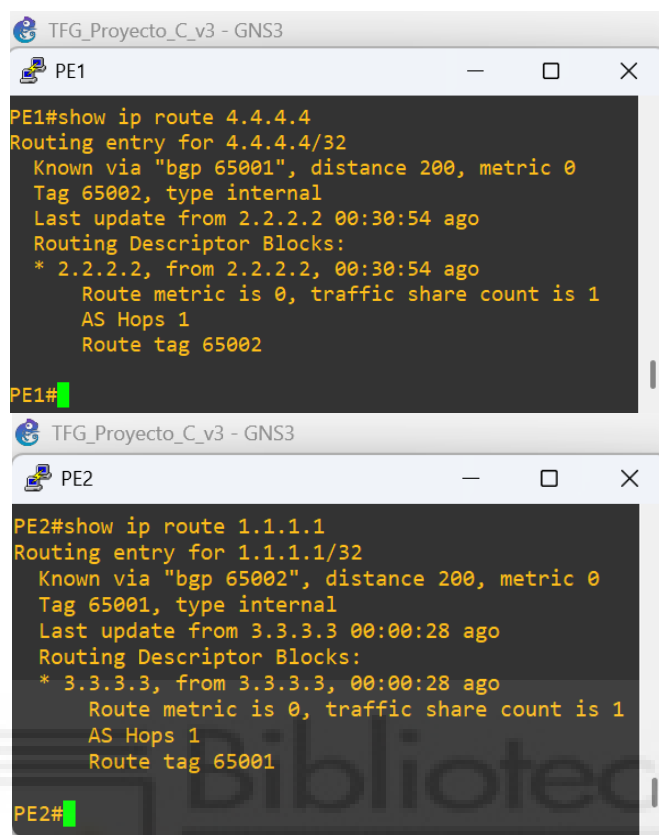
    172.17.0.0/24 is subnetted, 2 subnets
B       172.17.1.0 [20/0] via 10.1.1.2, 00:24:25
B       172.17.2.0 [20/0] via 4.4.4.4, 00:22:48
    172.16.0.0/24 is subnetted, 2 subnets
B       172.16.1.0 [20/0] via 10.1.1.2, 00:24:25
B       172.16.2.0 [20/0] via 4.4.4.4, 00:22:48
    10.0.0.0/30 is subnetted, 1 subnets
C       10.1.1.0 is directly connected, Ethernet0/0
PE1#
```

Figura 60: Tabla de rutas de la VRF EMPRESA en PE1 en el escenario C con tres routers.

En la Figura 60 se muestra show ip route vrf EMPRESA en PE1. Las redes locales 172.17.1.0/24 y 172.16.1.0/24 se alcanzan mediante 10.1.1.2, que corresponde a CE1. Las redes remotas 172.17.2.0/24 y 172.16.2.0/24 aparecen a través de 4.4.4.4, que corresponde a PE2.

Este resultado refleja la característica principal del escenario C. Aunque el tráfico atraviesa ASBR y routers P, las rutas VPNv4 remotas dentro de la VRF apuntan al PE remoto y no al ASBR. Por tanto, PE1 recibe la información VPN directamente desde PE2, mientras que el núcleo se encarga únicamente de proporcionar transporte etiquetado.

3.6.9. Verificación de alcanzabilidad entre loopbacks de PE



The image shows two terminal windows from a GNS3 simulation. The top window is titled 'PE1' and shows the output of the command 'show ip route 4.4.4.4'. The output indicates a routing entry for 4.4.4.4/32, known via BGP 65001, with a distance of 200 and metric 0. It also shows the last update from 2.2.2.2 and the routing descriptor blocks, including the route metric, traffic share count, AS hops, and route tag. The bottom window is titled 'PE2' and shows the output of the command 'show ip route 1.1.1.1'. The output indicates a routing entry for 1.1.1.1/32, known via BGP 65002, with a distance of 200 and metric 0. It also shows the last update from 3.3.3.3 and the routing descriptor blocks, including the route metric, traffic share count, AS hops, and route tag.

```
PE1#show ip route 4.4.4.4
Routing entry for 4.4.4.4/32
  Known via "bgp 65001", distance 200, metric 0
  Tag 65002, type internal
  Last update from 2.2.2.2 00:30:54 ago
  Routing Descriptor Blocks:
    * 2.2.2.2, from 2.2.2.2, 00:30:54 ago
      Route metric is 0, traffic share count is 1
      AS Hops 1
      Route tag 65002
PE1#
```

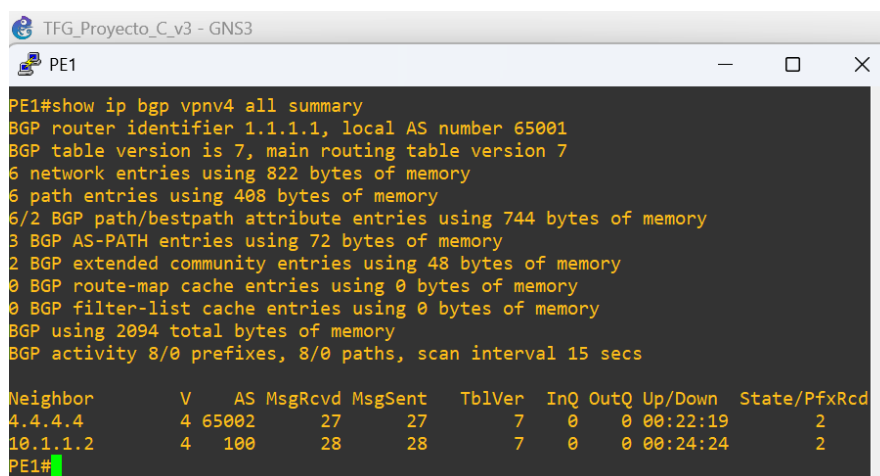
```
PE2#show ip route 1.1.1.1
Routing entry for 1.1.1.1/32
  Known via "bgp 65002", distance 200, metric 0
  Tag 65001, type internal
  Last update from 3.3.3.3 00:00:28 ago
  Routing Descriptor Blocks:
    * 3.3.3.3, from 3.3.3.3, 00:00:28 ago
      Route metric is 0, traffic share count is 1
      AS Hops 1
      Route tag 65001
PE2#
```

Figura 61: Rutas globales hacia las loopbacks remotas de los PE en el escenario C con tres routers.

En la Figura 61 se incorporan las comprobaciones show ip route 4.4.4.4 en PE1 y show ip route 1.1.1.1 en PE2. En PE1 se observa que la loopback de PE2 se alcanza mediante BGP, con actualización desde 2.2.2.2, que corresponde a ASBR1. En PE2 ocurre lo equivalente hacia la loopback de PE1, aprendida desde 3.3.3.3, correspondiente a ASBR2.

Esta comprobación se incluye antes del resumen de vecinos porque permite justificar la base necesaria para que la sesión eBGP multihop VPNv4 entre PE1 y PE2 funcione correctamente. En el escenario C, los PE deben poder alcanzarse a través de sus loopbacks en la tabla global, ya que sobre esa conectividad se establece el intercambio directo de rutas VPNv4.

3.6.10. Resumen de vecinos BGP VPNv4



The screenshot shows a terminal window titled 'TFG_Proyecto_C_v3 - GNS3' with a sub-header 'PE1'. The terminal displays the output of the command 'show ip bgp vpnv4 all summary'. The output includes BGP statistics and a table of neighbors.

```
PE1#show ip bgp vpnv4 all summary
BGP router identifier 1.1.1.1, local AS number 65001
BGP table version is 7, main routing table version 7
6 network entries using 822 bytes of memory
6 path entries using 408 bytes of memory
6/2 BGP path/bestpath attribute entries using 744 bytes of memory
3 BGP AS-PATH entries using 72 bytes of memory
2 BGP extended community entries using 48 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
BGP using 2094 total bytes of memory
BGP activity 8/0 prefixes, 8/0 paths, scan interval 15 secs

Neighbor      V    AS MsgRcvd MsgSent  TblVer  InQ OutQ Up/Down  State/PfxRcd
4.4.4.4        4 65002    27     27      7    0   0 00:22:19      2
10.1.1.2       4   100    28     28      7    0   0 00:24:24      2
PE1#
```

Figura 62: Resumen de vecinos BGP VPNv4 en PE1 en el escenario C con tres routers.

En la Figura 62 se muestra `show ip bgp vpnv4 all summary` en PE1. La salida confirma que PE1 mantiene una sesión VPNv4 con 4.4.4.4, correspondiente a PE2 en el AS 65002, y otra relación asociada al cliente mediante 10.1.1.2. En ambos casos se reciben dos prefijos, lo que indica que las sesiones están establecidas y que las rutas necesarias están siendo intercambiadas.

Los resultados obtenidos encajan con el diseño del escenario C con tres routers. El núcleo interno queda validado demostrando el transporte IPv4 etiquetado, y PE1 confirma tanto la presencia de rutas en la VRF como la sesión VPNv4 directa con PE2. La diferencia respecto al escenario B con tres routers es que aquí los ASBR no son los encargados de intercambiar rutas VPNv4, sino de permitir el transporte etiquetado entre los PE de ambos sistemas autónomos.

3.7. Resumen comparativo de ventajas y desventajas de los escenarios

3.7.1. Escenario A

El escenario A es el más sencillo de entender y de comprobar, ya que la comunicación entre proveedores se realiza manteniendo la continuidad de la VRF entre los ASBR. Su principal ventaja es que la configuración resulta bastante directa y permite ver de forma clara cómo las rutas del cliente se mantienen separadas dentro de la VRF

EMPRESA. Además, las pruebas de ping, traceroute y rutas VRF son fáciles de interpretar. Como desventaja, este escenario puede ser menos escalable si se quisiera aplicar a una red más grande, porque obliga a mantener configuración de VRF en la frontera entre proveedores y hace que los ASBR tengan una relación más directa con el servicio del cliente.

3.7.2. Escenario B

El escenario B supone un paso intermedio entre sencillez y escalabilidad. Su principal ventaja es que los ASBR intercambian rutas VPNv4 entre sistemas autónomos, por lo que ya no se depende de una continuidad VRF tan directa como en el escenario A. Esto permite separar mejor el funcionamiento interno de cada proveedor y hace que el modelo sea más adecuado para redes más grandes. Como desventaja, la configuración y la verificación son algo más complejas, ya que hay que comprobar no solo la VRF y la conectividad, sino también las rutas VPNv4 en los ASBR y el funcionamiento del transporte MPLS.

3.7.3. Escenario C

El escenario C es el más completo y también el más complejo de los tres. Su ventaja principal es que las rutas VPNv4 se intercambian directamente entre los routers PE mediante eBGP multihop, mientras que los ASBR se encargan únicamente del transporte IPv4 etiquetado. Esto permite separar mejor el servicio VPN del transporte entre proveedores. Además, desde el punto de vista del diseño, es un escenario más avanzado y limpio en cuanto a funciones. Como desventaja, requiere más comprobaciones para verificar que todo funciona correctamente, ya que no basta con ver la VRF o las rutas VPNv4, sino que también hay que comprobar la alcanzabilidad entre loopbacks de PE, las rutas IPv4 etiquetadas y el reenvío MPLS en los ASBR.

3.7.4. Conclusión comparativa de los escenarios

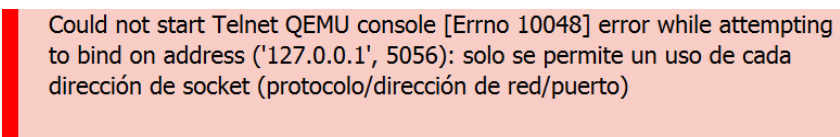
Criterio	Escenario A	Escenario B	Escenario C
Interconexión inter-AS	VRF entre ASBR	VPNv4 entre ASBR	VPNv4 entre PE
Papel de ASBR	Continuidad de VRF	Intercambio VPNv4	Transporte etiquetado
Escalabilidad	Baja	Media	Alta
Complejidad	Baja	Media	Alta

Tabla 1: Comparación general de los escenarios A, B y C.

En la tabla 1 se comparan de forma resumida los tres escenarios desarrollados en el trabajo. El escenario A es el más sencillo, ya que se basa en mantener la VRF entre los ASBR, aunque por ello también resulta menos escalable. El escenario B representa un punto intermedio, porque los ASBR intercambian rutas VPNv4 y permiten una separación mayor entre proveedores. Por último, el escenario C es el más avanzado, ya que los PE intercambian directamente la información VPNv4 y los ASBR se centran en el transporte etiquetado, aunque esto hace que su configuración y verificación sean más complejas.

3.8. Observaciones complementarias durante la validación

3.8.1. Error de consola Telnet QEMU



```
Could not start Telnet QEMU console [Errno 10048] error while attempting  
to bind on address ('127.0.0.1', 5056): solo se permite un uso de cada  
dirección de socket (protocolo/dirección de red/puerto)
```

Figura 63: Error de consola Telnet QEMU durante la apertura de un equipo en GNS3.

En la Figura 63 se muestra un error de GNS3 relacionado con la apertura de la consola Telnet de QEMU. El mensaje indica que no se puede enlazar la dirección 127.0.0.1 con el puerto 5056, ya que solo se permite un uso de cada dirección de socket. Este error no está relacionado con la configuración de red del escenario, sino con el entorno de emulación y con el uso de puertos locales en el sistema operativo.

El código Errno 10048 corresponde al error WSAEADDRINUSE, que indica que una aplicación intenta utilizar una dirección IP y un puerto que ya están siendo usados por otro socket, o que no se han cerrado correctamente. En este caso, GNS3 intenta abrir una consola local mediante Telnet, pero el puerto asignado ya está ocupado o bloqueado temporalmente. Por tanto, la solución habitual sería cerrar la consola o proceso que está usando ese puerto, reiniciar el nodo afectado, cambiar el puerto de consola o reiniciar GNS3 si el puerto ha quedado reservado por una sesión anterior (Microsoft, 2021).

3.8.2. Pérdida inicial de paquetes en la primera prueba de ping

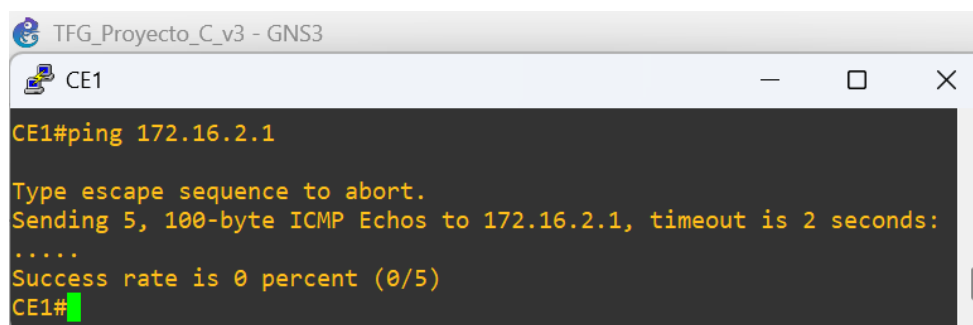
```
CE1#ping 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
Packet sent with a source address of 172.16.1.1
...!!
Success rate is 40 percent (2/5), round-trip min/avg/max = 144/146/148 ms
CE1#ping 172.16.2.1 source 172.16.1.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
Packet sent with a source address of 172.16.1.1
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 140/150/164 ms
CE1#
```

Figura 64: Pérdida inicial de paquetes y posterior conectividad correcta en una prueba de ping.

En la Figura 64 se observa que el primer ping desde CE1 hacia 172.16.2.1, usando como origen 172.16.1.1, obtiene inicialmente una tasa de éxito del 40%, mientras que al repetir la misma prueba el resultado pasa a ser del 100%. Este comportamiento puede aparecer justo después de iniciar o modificar el escenario, cuando todavía se están completando procesos internos como la resolución ARP, la instalación de rutas, la estabilización de BGP o la disponibilidad del transporte MPLS.

La causa más habitual en los primeros paquetes es que el equipo todavía necesite resolver direcciones de enlace antes de poder reenviar el tráfico correctamente. ARP se encarga precisamente de obtener la dirección física asociada a una dirección IP dentro de una red Ethernet, y si esa información aún no está disponible, los primeros paquetes pueden perderse mientras se completa la resolución. Por eso, al repetir el ping unos segundos después, la conectividad ya aparece estabilizada y la tasa de éxito pasa al 100% (Plummer, 1982).

3.8.3. Ping sin origen específico

A screenshot of a GNS3 terminal window titled 'TFG_Proyecto_C_v3 - GNS3'. The terminal shows a command prompt 'CE1' followed by the command 'ping 172.16.2.1'. The output of the command is: 'Type escape sequence to abort. Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds: Success rate is 0 percent (0/5) CE1#'. The text is displayed in a monospaced font with a dark background and light-colored text.

```
TFG_Proyecto_C_v3 - GNS3
CE1
CE1#ping 172.16.2.1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 172.16.2.1, timeout is 2 seconds:
.....
Success rate is 0 percent (0/5)
CE1#
```

Figura 65: Prueba de ping desde CE1 hacia la loopback remota de CE2 sin utilizar origen específico.

En la Figura 65 se muestra un ping desde CE1 hacia 172.16.2.1 sin indicar el parámetro source. El resultado es una tasa de éxito del 0%, a diferencia de las pruebas anteriores donde sí se utilizaba source 172.16.1.1. Esto ocurre porque, al no forzar el origen, el router no representa necesariamente el tráfico desde la loopback del cliente, sino desde otra dirección propia del camino de salida, por ejemplo, la dirección del enlace CE-PE, cuya dirección de red no ha sido configurada para salir de su sistema autónomo.

En este trabajo interesa comprobar la comunicación entre las redes del cliente que se anuncian dentro de la VPN, como las loopbacks 172.16.1.1 y 172.16.2.1. En una VPN MPLS, las rutas que el PE aprende del CE se instalan dentro de la VRF correspondiente y se distribuyen como rutas VPN hacia otros PE. Por eso, al usar source 172.16.1.1, la prueba representa correctamente el tráfico entre sedes del cliente. En cambio, sin ese origen, la respuesta puede no encontrar una ruta válida de vuelta hacia la dirección usada como origen, y el ping falla, aunque la VPN esté funcionando correctamente (Rosen & Rekhter, 2006).

4. CONCLUSIONES

Con este trabajo se ha podido comprobar de forma práctica cómo funciona una red MPLS VPN Inter-AS en GNS3 y cómo cambia su comportamiento según el escenario utilizado. Aunque al principio los tres modelos pueden parecer similares, durante la configuración y la validación se ha visto que cada uno resuelve la interconexión entre proveedores de una forma distinta. En todos los casos se ha conseguido conectividad entre las sedes del cliente, lo que confirma que las configuraciones aplicadas cumplen el objetivo principal del trabajo.

El escenario A ha sido el más sencillo de seguir, porque la comunicación entre proveedores se mantiene dentro de la VRF y eso hace que las rutas del cliente sean más fáciles de interpretar. El escenario B añade un paso más, ya que los ASBR empiezan a tener un papel más importante al intercambiar rutas VPNv4 entre sistemas autónomos. Por último, el escenario C ha sido el más avanzado, porque separa mejor el transporte del servicio VPN y permite que los PE intercambien directamente la información VPNv4, mientras los ASBR se encargan de transportar rutas IPv4 etiquetadas.

Las pruebas realizadas también han demostrado que no basta con comprobar que un ping responde. Para entender realmente si el escenario está funcionando, ha sido necesario revisar las tablas de rutas de la VRF, las rutas BGP VPNv4, las vecindades OSPF y LDP, las tablas MPLS y, en el escenario C, las rutas etiquetadas y la alcanzabilidad entre loopbacks de los PE. Gracias a estos comandos se ha podido justificar de forma más completa por qué existe conectividad y qué papel cumple cada router dentro de la topología.

Además, durante la validación han aparecido detalles propios de un entorno emulado, como pérdidas iniciales de paquetes, la necesidad de usar ping con origen específico o algún error puntual de consola en GNS3. Estos aspectos no se consideran fallos de los escenarios, sino situaciones normales que pueden aparecer durante las pruebas y que hay que saber interpretar.

También hay que tener en cuenta que el trabajo se ha realizado sobre una única VRF, lo que ha condicionado algunas decisiones de configuración y simplifica el análisis realizado. Como trabajo futuro, se podría ampliar el laboratorio incorporando múltiples clientes, varias VRF y más sistemas autónomos, con el objetivo de acercar todavía más la topología a un entorno real de proveedor. En general, el trabajo ha permitido comparar los tres escenarios, entender mejor sus ventajas y limitaciones, y ver de forma práctica cómo se puede establecer una VPN entre clientes conectados a sistemas autónomos diferentes.



5. ANEXO

En este apartado se recogen y se comentan los fragmentos de configuración más relevantes empleados en el desarrollo de los distintos escenarios del trabajo. Su finalidad no es repetir la explicación teórica ya presentada en los apartados anteriores, sino complementar esa información mostrando de forma más directa cómo se ha materializado cada escenario en GNS3. De este modo, el anexo actúa como un apoyo técnico a la memoria principal y permite relacionar la descripción metodológica con la configuración realmente utilizada. Además, este apartado no se plantea como una simple acumulación de líneas de código, sino como una guía breve y ordenada que acompaña cada bloque de configuración con una explicación de su función dentro del escenario correspondiente.

La organización del anexo sigue una lógica progresiva. En primer lugar, se presentarán aquellos elementos cuya configuración se repite de manera prácticamente idéntica en todos los escenarios, como ocurre con los equipos finales y otros componentes comunes del entorno. Posteriormente, se abordarán las configuraciones específicas de cada alternativa estudiada, diferenciando entre las versiones más simples y las ampliadas. Esta estructura permite evitar repeticiones innecesarias y facilita que el lector identifique con claridad qué partes del código son comunes y cuáles responden a las particularidades de cada escenario.

5.1. VPCS

5.1.1. PC1

```
ip 172.17.1.20/24 172.17.1.1  
save
```

En el caso de PC1, la configuración se realiza mediante una sintaxis muy simple, en la que se introduce primero el comando `ip` y, a continuación, la dirección IP del equipo, la máscara en formato prefijo y la puerta de enlace predeterminada. En este ejemplo, `172.17.1.20/24` identifica al host dentro de la red local `172.17.1.0/24`, mientras que `172.17.1.1` corresponde a la puerta de enlace o gateway de esa LAN, es decir, a la interfaz del router CE conectada al segmento local y encargada de dar salida al tráfico hacia el

resto de la red. Después de definir estos parámetros, se emplea el comando save, que en este tipo de PC permite guardar la configuración introducida para que no se pierda al reiniciar el nodo.

5.1.2. PC2

```
ip 172.17.2.20/24 172.17.2.1  
save
```

En PC2, la configuración sigue exactamente la misma lógica que en PC1, utilizando el comando ip para indicar la dirección del host, la máscara y la puerta de enlace, y save para guardar los cambios. La única diferencia está en los valores introducidos, ya que este equipo pertenece a la red 172.17.2.0/24, por lo que su dirección IP es 172.17.2.20/24 y su gateway correspondiente es 172.17.2.1, es decir, la interfaz LAN del router CE de ese extremo.

5.2. LinuxCore

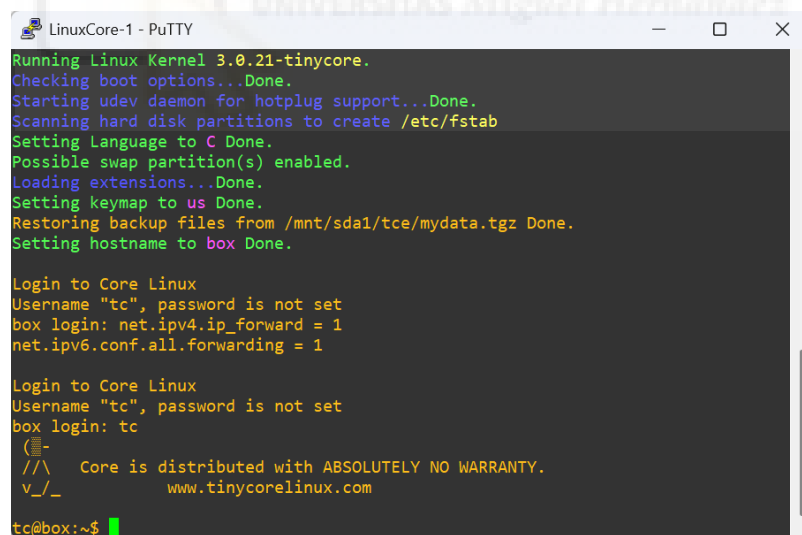
A screenshot of a PuTTY terminal window titled 'LinuxCore-1 - PuTTY'. The terminal shows the boot sequence of a LinuxCore system. It starts with 'Running Linux Kernel 3.0.21-tinycore.', followed by 'Checking boot options...Done.', 'Starting udev daemon for hotplug support...Done.', 'Scanning hard disk partitions to create /etc/fstab', 'Setting Language to C Done.', 'Possible swap partition(s) enabled.', 'Loading extensions...Done.', 'Setting keymap to us Done.', 'Restoring backup files from /mnt/sda1/tce/mydata.tgz Done.', and 'Setting hostname to box Done.'. Then it prompts for login: 'Login to Core Linux', 'Username "tc", password is not set', 'box login: net.ipv4.ip_forward = 1', 'net.ipv6.conf.all.forwarding = 1'. It then prompts for a shell: 'Login to Core Linux', 'Username "tc", password is not set', 'box login: tc'. Finally, it shows a prompt 'tc@box:~\$' with a green cursor. The terminal also displays a copyright notice: 'Core is distributed with ABSOLUTELY NO WARRANTY. www.tinycorelinux.com'.

Figura 66: Arranque e inicio de sesión en Linux Core.

En la

Figura 66 se observa el proceso de arranque de un nodo Linux Core dentro de GNS3. Durante esta secuencia, el sistema carga sus componentes básicos, aplica los elementos restaurados desde la copia de seguridad y deja preparado el entorno para su

utilización. Una vez finalizado el arranque, para poder acceder al sistema es necesario introducir el usuario tc, que es el usuario por defecto en este entorno. Esta fase inicial, aunque sencilla, resulta importante dentro del trabajo porque marca el punto de partida desde el que después se introducen los comandos de configuración o se comprueba que la configuración previamente guardada se ha cargado correctamente.

5.2.1. LinuxCore1

```
sudo ifconfig eth0 172.17.1.10 netmask 255.255.255.0 up  
sudo route add default gw 172.17.1.1 eth0
```

En Linux Core, la dirección IP y la puerta de enlace se configuran desde terminal. En este caso, `sudo ifconfig eth0 172.17.1.10 netmask 255.255.255.0 up` asigna a la interfaz `eth0` la dirección `172.17.1.10`, activa la interfaz y fija de forma explícita la máscara de red `255.255.255.0`, mientras que `sudo route add default gw 172.17.1.1 eth0` añade la ruta por defecto hacia `172.17.1.1`, que corresponde a la interfaz LAN del router CE de ese extremo. El prefijo `sudo` solo se explica una vez, ya que su función es ejecutar el comando con privilegios de administración, algo necesario para modificar la configuración de red del sistema. Además, la propia sintaxis clásica de `ifconfig` contempla el uso explícito de `netmask`, y `route` se utiliza para definir la puerta de enlace por defecto del host (Linux man-pages project, 2025).

En este trabajo se prefirió escribir la máscara completa `255.255.255.0` en lugar de utilizar la notación `/24`, a diferencia de lo que ocurre en los VPCS. La razón es práctica. Durante la configuración del nodo Linux Core, el uso de `ifconfig` con notación CIDR dio lugar a una interpretación incorrecta de la máscara, por lo que se optó por la forma más segura y explícita dentro de este entorno. De hecho, la documentación de `ifconfig` recoge la máscara como un parámetro independiente mediante `netmask MASK`, lo que encaja plenamente con la sintaxis utilizada aquí. Por tanto, aunque en otros contextos Linux pueda usarse notación abreviada, en este trabajo resultó más fiable definir la máscara completa para evitar ambigüedades en la configuración del host (Linux man-pages project, 2025).

A continuación, se muestran los comandos utilizados para guardar la configuración en los nodos Linux, ya que cada vez que se cerraba el proyecto era necesario volver a introducirla manualmente. Dado que este proceso resultaba tedioso, se investigó una forma de conservar dicha configuración entre reinicios.

```
echo '#!/bin/sh' > /opt/eth0.sh
echo 'ifconfig eth0 172.17.1.10 netmask 255.255.255.0 up' >>
/opt/eth0.sh
echo 'route add default gw 172.17.1.1 eth0' >> /opt/eth0.sh
chmod +x /opt/eth0.sh
echo 'sysctl -w net.ipv4.ip_forward=1' >> /opt/bootlocal.sh
echo 'sysctl -w net.ipv6.conf.all.forwarding=1' >>
/opt/bootlocal.sh
echo '/opt/eth0.sh &' >> /opt/bootlocal.sh
echo 'opt/eth0.sh' >> /opt/.filetool.lst
echo 'opt/bootlocal.sh' >> /opt/.filetool.lst
filetool.sh -b
```

Una vez definida correctamente la dirección IP, el siguiente paso consistió en conseguir que esa configuración no se perdiera tras reiniciar el nodo. Para ello, se creó primero el archivo `/opt/eth0.sh`, que reúne en un único script las órdenes necesarias para configurar la interfaz `eth0` y añadir la ruta por defecto. Después, ese script se marcó como ejecutable y se invocó desde `/opt/bootlocal.sh`, que en Tiny Core Linux es el archivo destinado a ejecutar comandos de arranque una vez que el sistema ya ha completado la fase inicial del boot. La propia documentación oficial de Tiny Core indica que `/opt/bootlocal.sh` se usa precisamente para lanzar comandos cada vez que el sistema arranca y que la mayoría de órdenes personalizadas deben colocarse ahí (Tiny Core Linux Wiki, 2011a).

Sin embargo, en Tiny Core no basta con crear o modificar esos archivos. También es necesario indicar qué elementos deben conservarse entre reinicios. Para eso se añadió `opt/eth0.sh` y `opt/bootlocal.sh` a `/opt/.filetool.lst`, que es el fichero donde Tiny Core guarda la lista de archivos y directorios incluidos en la copia de seguridad. La documentación oficial explica que las rutas incluidas en `/opt/.filetool.lst` deben escribirse de forma relativa a la raíz y sin la barra inicial, exactamente como se ha hecho en este caso. Después de actualizar esa lista, se ejecutó `filetool.sh -b`, que es el comando oficial de Tiny Core para generar la copia de seguridad y guardar los archivos persistentes en `mydata.tgz`. Así, el guardado no depende solo del script de arranque, sino de la combinación entre

/opt/bootlocal.sh, /opt/.filetool.lst y el uso de filetool.sh -b como mecanismo de persistencia (Tiny Core Linux Wiki, 2011b).

En consecuencia, el proceso de guardado en Linux Core puede entenderse como una secuencia de tres pasos. Primero se encapsula la configuración de red en un script reutilizable. Después se fuerza su ejecución automática en cada arranque mediante /opt/bootlocal.sh. Finalmente, se incluye ese contenido en la copia de seguridad del sistema para que Tiny Core lo restaure tras reiniciar. Esta forma de trabajo resulta especialmente útil en un entorno de emulación como el de este trabajo, porque evita tener que reconfigurar manualmente la interfaz Linux cada vez que se vuelve a abrir el escenario y asegura que el host mantenga una configuración coherente con la red del cliente en todos los montajes posteriores.

5.2.2. LinuxCore2

```
sudo ifconfig eth0 172.17.2.10 netmask 255.255.255.0 up
sudo route add default gw 172.17.2.1 eth0

echo '#!/bin/sh' > /opt/eth0.sh
echo 'ifconfig eth0 172.17.2.10 netmask 255.255.255.0 up' >>
/opt/eth0.sh
echo 'route add default gw 172.17.2.1 eth0' >> /opt/eth0.sh
chmod +x /opt/eth0.sh
echo 'sysctl -w net.ipv4.ip_forward=1' >> /opt/bootlocal.sh
echo 'sysctl -w net.ipv6.conf.all.forwarding=1' >>
/opt/bootlocal.sh
echo '/opt/eth0.sh &' >> /opt/bootlocal.sh
echo 'opt/eth0.sh' >> /opt/.filetool.lst
echo 'opt/bootlocal.sh' >> /opt/.filetool.lst
filetool.sh -b
```

En LinuxCore2, la configuración sigue exactamente la misma lógica que en el nodo anterior, cambiando únicamente los valores propios de su subred. Primero se asigna a eth0 la dirección 172.17.2.10 con máscara 255.255.255.0 y se establece como puerta de enlace por defecto 172.17.2.1, que corresponde a la interfaz LAN del router CE de ese extremo. Después, esa configuración se guarda creando el script /opt/eth0.sh, que reúne las órdenes necesarias para levantar la interfaz y añadir la ruta por defecto, dándole permisos de ejecución e incorporándolo al arranque mediante /opt/bootlocal.sh. A continuación, tanto este script como el propio bootlocal.sh se añaden a /opt/.filetool.lst

para que Tiny Core los incluya en la copia de seguridad del sistema, y finalmente se ejecuta `filetool.sh -b` para guardar los cambios. Cabe recalcar que para comprobar que se ha guardado satisfactoriamente, podemos usar `sudo reboot` para reiniciar el nodo y comprobar que la configuración queda aplicada de forma persistente.

5.3. Router CE

En este apartado se recogen los comandos correspondientes a los routers CE, es decir, los equipos que representan el borde del cliente en cada escenario. Su función dentro del trabajo consiste en enlazar la red local del cliente con la red del proveedor y anunciar hacia el router PE las redes que pertenecen a esa sede. Por ello, aunque cambien los escenarios A, B o C, la lógica general de configuración de los CE se mantiene prácticamente igual, ya que estos equipos no participan en la complejidad interna del proveedor, sino que siguen actuando como routers de cliente en el extremo de la topología.

Precisamente por esa razón, en todos los escenarios se utilizará una estructura de comandos muy similar en CE1 y CE2. En ambos casos se configura una interfaz orientada hacia el proveedor, otra interfaz asociada a la LAN local y una loopback adicional que representa otra red del cliente. Después, se activa BGP para anunciar esos prefijos hacia el PE correspondiente. Esta lógica encaja con el funcionamiento general de BGP-4, donde los vecinos se establecen manualmente y las rutas de alcanzabilidad se anuncian mediante prefijos asociados al sistema autónomo del router que los origina (Rekhter et al., 2006).

5.3.1. Router CE1

```
conf term
hostname CE1

int e0/0
description CE1-PE1
ip address 10.1.1.2 255.255.255.252
no shut

int e0/1
description LAN_CE1
ip address 172.17.1.1 255.255.255.0
```

```
no shut

int Loopback0
ip address 172.16.1.1 255.255.255.0

router bgp 100
neighbor 10.1.1.1 remote-as 65001
network 172.16.1.0 mask 255.255.255.0
network 172.17.1.0 mask 255.255.255.0

end
```

En CE1, la configuración comienza con `conf term`, que permite entrar en modo de configuración global, y con `hostname CE1`, que asigna un identificador claro al equipo dentro del escenario. A continuación, en `int e0/0` se configura la interfaz que conecta con PE1, añadiendo una descripción para identificar su función, una dirección IP de enlace punto a punto y el comando `no shut` para dejar la interfaz activa. Después, en `int e0/1` se define la interfaz correspondiente a la LAN_CE1, que en este caso actúa como puerta de enlace para los hosts conectados a ese extremo del trabajo.

Tras ello, se crea Loopback0 con la dirección 172.16.1.1 255.255.255.0, que representa una red adicional del cliente y permite que el escenario no se limite únicamente a la LAN conectada físicamente al CE. Después se entra en `router bgp 100`, donde el número 100 identifica el sistema autónomo del cliente en ese extremo. Mediante `neighbor 10.1.1.1 remote-as 65001` se establece la relación BGP con PE1, y con los comandos `network 172.16.1.0 mask 255.255.255.0` y `network 172.17.1.0 mask 255.255.255.0` se anuncian hacia el proveedor tanto la red de la loopback como la red local de la sede. Esta mecánica coincide con el funcionamiento general descrito en BGP-4, donde el router intercambia información de alcanzabilidad con sus vecinos y anuncia prefijos concretos a través de la sesión configurada (Rekhter et al., 2006).

En consecuencia, CE1 debe entenderse como el router que concentra las redes del cliente en el extremo izquierdo de la topología y las entrega al proveedor mediante una configuración sencilla y estable. Lo importante aquí no es tanto la complejidad de los comandos como su función dentro del trabajo: definir correctamente la salida hacia el proveedor, mantener operativa la LAN local y anunciar las redes de esa sede. Como esta lógica no cambia entre escenarios, la configuración de CE1 se repetirá prácticamente

igual en las distintas alternativas, variando solo cuando sea necesario ajustar algún valor concreto del direccionamiento.

5.3.2. Router CE2

```
conf term
hostname CE2

int e0/0
description CE2-PE2
ip address 10.2.2.2 255.255.255.252
no shut

int e0/1
description LAN_CE2
ip address 172.17.2.1 255.255.255.0
no shut

int Loopback0
ip address 172.16.2.1 255.255.255.0

router bgp 200
neighbor 10.2.2.1 remote-as 65002
network 172.16.2.0 mask 255.255.255.0
network 172.17.2.0 mask 255.255.255.0

end
```

La configuración de CE2 sigue exactamente la misma lógica que la de CE1, por lo que no es necesario volver a desarrollar su estructura con el mismo nivel de detalle. En este caso, se repiten los mismos pasos generales, es decir, la definición del nombre del router, la configuración de la interfaz que conecta con el proveedor, la configuración de la interfaz asociada a la LAN local, la creación de la loopback adicional y la activación del proceso BGP para anunciar las redes del cliente. Las diferencias se encuentran únicamente en los valores concretos utilizados, como la dirección IP del enlace hacia PE2, la red local 172.17.2.0/24, la loopback 172.16.2.0/24 y el número de sistema autónomo del cliente en este extremo, que pasa a ser AS 200. Por tanto, CE2 debe entenderse como una configuración equivalente a la de CE1, adaptada únicamente al direccionamiento y al contexto del lado derecho de la topología.

5.4. Escenario A con dos routers

En este apartado se recogen los comandos correspondientes a los routers de proveedor del escenario A, es decir, PE1, ASBR1, ASBR2 y PE2. Aunque en esta topología también intervienen CE1 y CE2, sus configuraciones no se volverán a incluir aquí, ya que su explicación ya ha sido desarrollada en el apartado 5.3 y se repiten en todos los escenarios con cambios mínimos de direccionamiento. Del mismo modo, para evitar una fragmentación excesiva del anexo, en esta sección no se abrirán apartados genéricos por tipo de router, sino que se explicará directamente la configuración de cada equipo de forma individual. Por ello, en PE1 y ASBR1, cuya lógica aparece aquí por primera vez, la explicación será más desarrollada, mientras que en ASBR2 y PE2 bastará con remarcar que se trata de configuraciones equivalentes adaptadas al extremo opuesto de la topología.

5.4.1. Router PE1

```
conf term
hostname PE1

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65002:1

int e0/0
description PE1-CE1
ip vrf forwarding EMPRESA
ip address 10.1.1.1 255.255.255.252
no shut

int e0/1
description PE1-ASBR1
ip vrf forwarding EMPRESA
ip address 192.168.1.1 255.255.255.252
no shut

router bgp 65001
bgp router-id 1.1.1.1
address-family ipv4 vrf EMPRESA
neighbor 10.1.1.2 remote-as 100
neighbor 10.1.1.2 activate
neighbor 192.168.1.2 remote-as 65001
neighbor 192.168.1.2 activate
```

```
neighbor 192.168.1.2 next-hop-self
exit-address-family

end
```

En PE1, la configuración comienza con hostname PE1, que identifica el router dentro del escenario, y continúa con la creación de la VRF EMPRESA. En este bloque se definen el route distinguisher mediante rd 65001:1 y los valores de route-target export y route-target import, que son los parámetros que permiten diferenciar y controlar qué rutas pertenecen al servicio. Dentro de la lógica de una L3VPN, esta parte resulta esencial porque convierte al router PE en el punto donde el tráfico del cliente deja de formar parte de la tabla global y pasa a asociarse a una tabla independiente del servicio (Cisco, 2020a).

Después se configuran las dos interfaces físicas del router. En Ethernet0/0 se establece el enlace hacia CE1, aplicando ip vrf forwarding EMPRESA para indicar que esa interfaz queda vinculada a la VRF del cliente. A continuación, se asigna la dirección 10.1.1.1/30, que corresponde al enlace punto a punto con el CE, y se activa la interfaz con no shutdown. La segunda interfaz, Ethernet0/1, enlaza con ASBR1 y también queda asociada a la misma VRF, con la dirección 192.168.1.1/30. Esto significa que, en el escenario A, tanto el enlace hacia el cliente como el enlace hacia la frontera del proveedor quedan dentro del contexto de la misma VRF, lo que encaja con la idea general de esta opción, donde el servicio se prolonga hasta la interconexión entre ASBR (Cisco, 2022c).

Por último, en router bgp 65001 se activa BGP para el sistema autónomo del proveedor y se entra en address-family ipv4 vrf EMPRESA, que es el contexto donde se manejan las rutas del cliente dentro de esa VRF. Aquí se define como vecino a CE1 mediante neighbor 10.1.1.2 remote-as 100 y también a ASBR1 con neighbor 192.168.1.2 remote-as 65001. El comando activate habilita ambos vecinos dentro de esta familia de direcciones y next-hop-self fuerza a PE1 a anunciarse a sí mismo como siguiente salto hacia ASBR1. Con ello, PE1 actúa como el punto de entrada del servicio en la red del proveedor y como el equipo encargado de vincular la red del cliente con la lógica VRF del escenario A (Cisco, 2020a).

5.4.2. Router ASBR1

```
conf term
hostname ASBR1

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65002:1

int e0/0
description ASBR1-PE1
ip vrf forwarding EMPRESA
ip address 192.168.1.2 255.255.255.252
no shut

int e0/1
description ASBR1-ASBR2
ip vrf forwarding EMPRESA
ip address 100.0.0.1 255.255.255.252
no shut

router bgp 65001
bgp router-id 2.2.2.2
address-family ipv4 vrf EMPRESA
neighbor 192.168.1.1 remote-as 65001
neighbor 192.168.1.1 activate
neighbor 192.168.1.1 next-hop-self
neighbor 100.0.0.2 remote-as 65002
neighbor 100.0.0.2 activate
exit-address-family

end
```

En ASBR1, la configuración comienza igualmente con `hostname ASBR1` y con la definición de la VRF EMPRESA. De nuevo aparecen `rd 65001:1`, `route-target export 65001:1` y `route-target import 65002:1`, lo que indica que este router comparte la misma lógica de servicio que PE1 dentro del AS 65001. Esto es coherente con el escenario A, ya que el ASBR no trabaja con una tabla global para el servicio del cliente, sino que mantiene también la VPN dentro de una VRF específica (Cisco, 2022c).

A continuación, se configuran sus dos interfaces principales. Ethernet0/0 enlaza con PE1, queda asociada a `ip vrf forwarding EMPRESA` y recibe la dirección 192.168.1.2/30. Por su parte, Ethernet0/1 conecta con ASBR2, también dentro de la

misma VRF, y se configura con la dirección 100.0.0.1/30. Esta parte es especialmente importante, porque muestra que en el escenario A la interconexión entre los dos proveedores se mantiene dentro del contexto de la VRF del cliente, en lugar de trasladarse a un intercambio VPNv4 o a un transporte etiquetado como ocurrirá en otros escenarios (Cisco, 2022c).

Finalmente, en router bgp 65001 se configura BGP para este ASBR. Dentro de address-family ipv4 vrf EMPRESA se establece la vecindad con PE1 mediante neighbor 192.168.1.1 remote-as 65001 y con ASBR2 mediante neighbor 100.0.0.2 remote-as 65002. Igual que antes, activate habilita ambas relaciones dentro de la VRF y next-hop-self se aplica sobre el vecino interno para controlar correctamente el siguiente salto dentro del AS 65001. Así, ASBR1 actúa como el equipo que enlaza el proveedor local con el proveedor remoto, pero sin abandonar la lógica VRF propia de la opción A (Cisco, 2022c).

5.4.3. Router ASBR2

```
conf term
hostname ASBR2

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65001:1

int e0/0
description ASBR2-ASBR1
ip vrf forwarding EMPRESA
ip address 100.0.0.2 255.255.255.252
no shut

int e0/1
description ASBR2-PE2
ip vrf forwarding EMPRESA
ip address 192.168.2.2 255.255.255.252
no shut

router bgp 65002
bgp router-id 3.3.3.3
address-family ipv4 vrf EMPRESA
neighbor 100.0.0.1 remote-as 65001
```

```

neighbor 100.0.0.1 activate
neighbor 192.168.2.1 remote-as 65002
neighbor 192.168.2.1 activate
neighbor 192.168.2.1 next-hop-self
exit-address-family

```

```
end
```

La configuración de ASBR2 sigue exactamente la misma idea que la de ASBR1, por lo que no es necesario repetir su explicación con el mismo nivel de detalle. En este caso, vuelven a definirse hostname, la VRF EMPRESA, su rd y los route-target correspondientes, adaptados ahora al AS 65002. También se configuran dos interfaces dentro de la misma VRF, una hacia PE2 y otra hacia ASBR1, y en router bgp 65002 se establecen las vecindades necesarias dentro de address-family ipv4 vrf EMPRESA. Las diferencias se encuentran únicamente en los valores concretos, como las direcciones 192.168.2.2/30 y 100.0.0.2/30, el identificador del router y el hecho de que aquí el ASBR pertenece al lado derecho de la topología. Por tanto, ASBR2 debe entenderse como la réplica funcional de ASBR1 en el otro proveedor.

5.4.4. Router PE2

```

conf term
hostname PE2

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65001:1

int e0/0
description PE2-ASBR2
ip vrf forwarding EMPRESA
ip address 192.168.2.1 255.255.255.252
no shut

int e0/1
description PE2-CE2
ip vrf forwarding EMPRESA
ip address 10.2.2.1 255.255.255.252
no shut

router bgp 65002
bgp router-id 4.4.4.4

```

```
address-family ipv4 vrf EMPRESA
neighbor 10.2.2.2 remote-as 200
neighbor 10.2.2.2 activate
neighbor 192.168.2.2 remote-as 65002
neighbor 192.168.2.2 activate
neighbor 192.168.2.2 next-hop-self
exit-address-family

end
```

En PE2, la lógica vuelve a ser equivalente a la de PE1, ya que este router representa el borde del proveedor en el extremo opuesto del escenario. Se define la VRF EMPRESA con su rd y sus route-target, se configura una interfaz hacia CE2 y otra hacia ASBR2, ambas asociadas a la misma VRF, y después se activa router bgp 65002 con la familia address-family ipv4 vrf EMPRESA. Dentro de ella se establecen los vecinos CE2 y ASBR2, y se aplica next-hop-self hacia el vecino interno del proveedor. Por ello, PE2 debe entenderse como una configuración equivalente a la de PE1, adaptada al direccionamiento, al identificador y al contexto del lado derecho de la topología, manteniendo la misma función de entrada y salida del servicio en la red del proveedor.

5.5. Escenario A con tres routers

En este apartado se recogen las configuraciones correspondientes a la versión ampliada del escenario A, es decir, aquella en la que cada proveedor incorpora un router interno de tránsito dentro de su propia red. Igual que ocurría en el apartado anterior, los routers CE1 y CE2 no se volverán a desarrollar aquí, ya que su configuración ya ha sido explicada en el apartado 5.3 y se mantiene esencialmente igual. La diferencia en esta versión aparece en la parte del proveedor, donde a los routers PE y ASBR se añade ahora un router P dentro de cada AS, lo que obliga a introducir elementos nuevos como OSPF, LDP/MPLS, interfaces de núcleo y sesiones VPNv4 internas. Para no multiplicar innecesariamente los niveles del anexo, se explicará directamente cada router por separado, destacando con más detalle los casos de PE1, P1 y ASBR1, y resumiendo después los equipos equivalentes del extremo derecho cuando su lógica sea prácticamente la misma.

5.5.1. Router PE1

```
conf term
hostname PE1

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65001:1

int Loopback0
ip address 1.1.1.1 255.255.255.255

int e0/0
description PE1-CE1
ip vrf forwarding EMPRESA
ip address 10.1.1.1 255.255.255.252
no shut

int e0/1
description PE1-P1
ip address 11.11.11.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 1.1.1.1
network 1.1.1.1 0.0.0.0 area 0
network 11.11.11.0 0.0.0.3 area 0
exit

router bgp 65001
bgp router-id 1.1.1.1
no bgp default ipv4-unicast
neighbor 2.2.2.2 remote-as 65001
neighbor 2.2.2.2 update-source Loopback0
address-family vpnv4
neighbor 2.2.2.2 activate
neighbor 2.2.2.2 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.1.1.2 remote-as 100
neighbor 10.1.1.2 activate
exit-address-family

end
```

En PE1, una parte importante de la configuración se mantiene próxima a la ya explicada en el apartado 5.4.1, ya que el router sigue actuando como borde del proveedor frente al cliente y continúa utilizando la VRF EMPRESA para separar el servicio. Por eso, comandos como hostname PE1, la creación de ip vrf EMPRESA, el uso de rd 65001:1, los route-target y la asignación de la interfaz hacia CE1 a esa VRF responden a la misma lógica ya comentada anteriormente. La diferencia aparece en que, en esta versión ampliada, el enlace hacia el interior del proveedor deja de estar también dentro de la VRF y pasa a convertirse en una interfaz global de núcleo, lo que explica la presencia de ip cef, mpls label protocol ldp, una Loopback0 propia del router y la configuración de Ethernet0/1 con mpls ip hacia P1 (Cisco, 2022b).

A partir de ahí, PE1 ya no se limita a enlazar el cliente con un ASBR dentro de la misma VRF, sino que debe participar en la red interna del proveedor. Por eso se activa router ospf 1, utilizando la loopback como identificador del router y anunciando tanto esa loopback como el enlace hacia P1 en el área 0. Después, en router bgp 65001, aparece no bgp default ipv4-unicast, se configura a ASBR1 como vecino interno a través de la loopback y se levanta address-family vpnv4, donde se activan la vecindad y el envío de comunidades extendidas. Mientras tanto, la address-family ipv4 vrf EMPRESA sigue utilizándose únicamente para la relación con CE1 (Cisco, 2020a). En consecuencia, PE1 mantiene su papel de borde del proveedor, pero en esta versión añade además las funciones necesarias para integrarse en un núcleo interno con OSPF, LDP y MP-BGP.

5.5.2. Router P1

```
conf term
hostname P1

ip cef
mpls label protocol ldp

int Loopback0
ip address 11.11.11.11 255.255.255.255

int e0/0
description P1-PE1
ip address 11.11.11.2 255.255.255.252
mpls ip
no shut
```

```

int e0/1
description P1-ASBR1
ip address 12.12.12.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 11.11.11.11
network 11.11.11.11 0.0.0.0 area 0
network 11.11.11.0 0.0.0.3 area 0
network 12.12.12.0 0.0.0.3 area 0

end

```

El router P1 es una de las novedades de esta versión del escenario A, ya que no aparecía en la topología reducida. Su papel no es conectar directamente con el cliente ni con la frontera entre proveedores, sino actuar como router de tránsito dentro del núcleo del AS 65001. Por eso, su configuración es más sencilla que la de un PE o un ASBR desde el punto de vista del servicio, pero muy importante desde el punto de vista del transporte. Se activa ip cef, se define mpls label protocol ldp, se configura una Loopback0 propia y se levantan dos interfaces físicas con mpls ip, una hacia PE1 y otra hacia ASBR1. Con ello, P1 se convierte en el nodo encargado de sostener el tráfico interno del proveedor entre el borde y la frontera (Cisco, 2022b).

Después, en router ospf 1, se configura su router-id y se anuncian tanto la loopback como los dos enlaces del core en el área 0. No aparece ninguna VRF ni ningún proceso BGP porque P1 no participa directamente en la construcción del servicio del cliente, sino únicamente en la conectividad y el transporte dentro del núcleo. En otras palabras, su misión consiste en proporcionar tránsito interno entre PE1 y ASBR1, haciendo que la red del proveedor deje de ser una simple conexión directa entre ambos equipos y pase a comportarse como un pequeño core real.

5.5.3. Router ASBR1

```

conf term
hostname ASBR1

ip cef
mpls label protocol ldp

```

```

ip vrf EMPRESA
rd 65001:10
route-target export 65001:1
route-target import 65001:1

int Loopback0
ip address 2.2.2.2 255.255.255.255

int e0/0
description ASBR1-P1
ip address 12.12.12.2 255.255.255.252
mpls ip
no shut

int e0/1
description ASBR1-ASBR2
ip vrf forwarding EMPRESA
ip address 100.0.0.1 255.255.255.252
no shut

router ospf 1
router-id 2.2.2.2
network 2.2.2.2 0.0.0.0 area 0
network 12.12.12.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 2.2.2.2
no bgp default ipv4-unicast
neighbor 1.1.1.1 remote-as 65001
neighbor 1.1.1.1 update-source Loopback0
address-family vpnv4
neighbor 1.1.1.1 activate
neighbor 1.1.1.1 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 100.0.0.2 remote-as 65002
neighbor 100.0.0.2 activate
exit-address-family

end

```

En ASBR1 también se conservan varias ideas ya vistas en el apartado 5.4.2, ya que el router sigue participando en la VRF EMPRESA y continúa siendo el equipo que enlaza su propio sistema autónomo con el proveedor remoto. Sin embargo, en esta versión ampliada ya no conecta directamente con PE1, sino que primero debe integrarse en el núcleo interno del AS 65001. Por eso se activan ip cef, mpls label protocol ldp, una Loopback0 propia y una interfaz global hacia P1 con mpls ip, mientras que la interfaz

hacia ASBR2 sigue asociada a ip vrf forwarding EMPRESA. De este modo, ASBR1 combina dos planos distintos dentro de una misma configuración: por un lado, el transporte interno del proveedor y, por otro, la continuidad del servicio VRF hacia el AS vecino (Cisco, 2022c).

Esta dualidad también se refleja en los protocolos. Por una parte, router ospf 1 anuncia la loopback y el enlace hacia P1, permitiendo que ASBR1 forme parte del core del proveedor (Moy, 1998). Por otra, en router bgp 65001 se configura como vecino interno a PE1 usando la loopback, y en address-family vpnv4 se activa esa sesión y el envío de comunidades. Además, address-family ipv4 vrf EMPRESA se utiliza únicamente para la interconexión con ASBR2, donde se anuncia la continuidad del servicio entre proveedores. En consecuencia, ASBR1 deja de ser un simple punto de frontera unido directamente a PE1 y pasa a actuar como un equipo que enlaza el núcleo interno del proveedor con la frontera inter-AS propia del escenario A (Cisco, 2020a).

5.5.4. Router ASBR2



```
conf term
hostname ASBR2

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65002:10
route-target export 65002:1
route-target import 65002:1
int Loopback0
ip address 3.3.3.3 255.255.255.255

int e0/0
description ASBR2-ASBR1
ip vrf forwarding EMPRESA
ip address 100.0.0.2 255.255.255.252
no shut

int e0/1
description ASBR2-P2
ip address 23.23.23.1 255.255.255.252
mpls ip
no shut
```

```

router ospf 1
router-id 3.3.3.3
network 3.3.3.3 0.0.0.0 area 0
network 23.23.23.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 3.3.3.3
no bgp default ipv4-unicast
neighbor 4.4.4.4 remote-as 65002
neighbor 4.4.4.4 update-source Loopback0
address-family vpnv4
neighbor 4.4.4.4 activate
neighbor 4.4.4.4 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 100.0.0.1 remote-as 65001
neighbor 100.0.0.1 activate
exit-address-family

end

```

La configuración de ASBR2 sigue la misma lógica que la de ASBR1, por lo que no es necesario repetirla con el mismo nivel de detalle. En este caso, vuelven a aparecer `ip cef`, `mpls label protocol ldp`, la `Loopback0`, una interfaz global con `mpls ip` hacia P2 y una interfaz dentro de la VRF `EMPRESA` hacia ASBR1. También se activa OSPF para integrar el router en el núcleo interno del AS 65002, y en BGP se establece la vecindad VPNv4 con PE2 mediante la `loopback`, mientras que la `address-family ipv4 vrf EMPRESA` se reserva para la interconexión con el proveedor remoto. Por tanto, ASBR2 debe entenderse como la réplica funcional de ASBR1 en el lado derecho de la topología, cambiando únicamente el direccionamiento, el `router-id`, los valores del `rd` y el contexto del AS 65002.

5.5.5. Router P2

```

conf term
hostname P2

ip cef
mpls label protocol ldp

int Loopback0
ip address 22.22.22.22 255.255.255.255

```

```

int e0/0
description P2-ASBR2
ip address 23.23.23.2 255.255.255.252
mpls ip
no shut

int e0/1
description P2-PE2
ip address 24.24.24.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 22.22.22.22
network 22.22.22.22 0.0.0.0 area 0
network 23.23.23.0 0.0.0.3 area 0
network 24.24.24.0 0.0.0.3 area 0
end

```

El router P2 desempeña en el AS 65002 la misma función que P1 en el AS 65001, por lo que su lógica de configuración también es equivalente. Se activan ip cef y mpls label protocol ldp, se define una Loopback0 propia y se configuran dos interfaces físicas con mpls ip, una hacia PE2 y otra hacia ASBR2. Después, en router ospf 1, se anuncian la loopback y los enlaces internos del core. Así, P2 no participa en la VRF del servicio ni en BGP, sino que se limita a proporcionar tránsito dentro del núcleo del proveedor, permitiendo que el tráfico avance entre el borde y la frontera de una forma más cercana a una red interna real.

5.5.6. Router PE2

```

conf term
hostname PE2

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65002:1

int Loopback0
ip address 4.4.4.4 255.255.255.255

int e0/0

```

```

description PE2-P2
ip address 24.24.24.2 255.255.255.252
mpls ip
no shut

int e0/1
description PE2-CE2
ip vrf forwarding EMPRESA
ip address 10.2.2.1 255.255.255.252
no shut

router ospf 1
router-id 4.4.4.4
network 4.4.4.4 0.0.0.0 area 0
network 24.24.24.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 4.4.4.4
no bgp default ipv4-unicast
neighbor 3.3.3.3 remote-as 65002
neighbor 3.3.3.3 update-source Loopback0
address-family vpnv4
neighbor 3.3.3.3 activate
neighbor 3.3.3.3 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.2.2.2 remote-as 200
neighbor 10.2.2.2 activate
exit-address-family

end

```

En PE2 se repite esencialmente la misma estructura que en PE1, adaptada al extremo derecho del escenario. Se vuelve a definir ip vrf EMPRESA, se mantiene la interfaz hacia CE2 dentro de la VRF y la interfaz hacia P2 como enlace global con mpls ip, y se incorpora una Loopback0 que actúa como identificador del router y punto de referencia para los protocolos internos. Después, router ospf 1 integra a PE2 en el núcleo del AS 65002, mientras que en router bgp 65002 se levanta la vecindad VPNv4 con ASBR2 mediante la loopback y se mantiene la address-family ipv4 vrf EMPRESA para la relación con CE2. Por ello, PE2 debe entenderse como la réplica funcional de PE1, cambiando únicamente el direccionamiento, el AS y el lado de la topología en el que se ubica.

5.6. Escenario B con dos routers

En este apartado se recogen las configuraciones correspondientes a la versión reducida del escenario B, es decir, aquella en la que cada proveedor se compone únicamente de un router PE y un router ASBR. Igual que en los apartados anteriores, los routers CE1 y CE2 no se volverán a desarrollar aquí, ya que su configuración ya fue explicada en el apartado 5.3 y se mantiene prácticamente igual. La diferencia respecto al escenario A no está en los extremos de cliente, sino en la parte del proveedor, ya que ahora la interconexión entre sistemas autónomos deja de resolverse mediante una prolongación VRF a VRF y pasa a basarse en la distribución de rutas VPNv4 entre ASBR, que es precisamente el rasgo que Cisco identifica como definitorio de la InterAS Option B (Cisco, 2022c).

5.6.1. Router PE1

```
conf term
hostname PE1

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65002:1

int lo0
ip address 1.1.1.1 255.255.255.255

int e0/0
description PE1-CE1
ip vrf forwarding EMPRESA
ip address 10.1.1.1 255.255.255.252
no shut

int e0/1
description PE1-ASBR1
ip address 192.168.1.1 255.255.255.252
mpls ip
no shut

router ospf 10
```

```

router-id 1.1.1.1
network 1.1.1.1 0.0.0.0 area 0
network 192.168.1.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 1.1.1.1
neighbor 11.11.11.11 remote-as 65001
neighbor 11.11.11.11 update-source lo0
address-family vpnv4
neighbor 11.11.11.11 activate
neighbor 11.11.11.11 send-community both
neighbor 11.11.11.11 next-hop-self
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.1.1.2 remote-as 100
neighbor 10.1.1.2 activate
exit-address-family

end

```

En PE1, una parte importante de la configuración sigue una lógica parecida a la ya explicada en los apartados 5.4.1 y 5.5.1, ya que el router continúa actuando como borde del proveedor frente al cliente y mantiene la VRF EMPRESA para separar el servicio. Por eso reaparecen elementos como hostname PE1, ip vrf EMPRESA, el rd 65001:1, los route-target y la interfaz hacia CE1 con ip vrf forwarding EMPRESA. La diferencia aparece en el enlace hacia ASBR1, que aquí ya no se mantiene dentro de la VRF, sino que pasa a configurarse como interfaz global con mpls ip, junto con ip cef, mpls label protocol ldp y una Loopback0 propia (Cisco, 2022b). Esto encaja con la lógica del escenario B, donde el core del proveedor debe transportar tráfico MPLS y sostener el intercambio VPNv4 entre los routers de borde (Cisco, 2022c).

A partir de ahí, PE1 se integra en el núcleo del proveedor mediante router ospf 10, utilizando la loopback como identificador y anunciando tanto esa loopback como el enlace hacia ASBR1 (Moy, 1998). Después, en router bgp 65001, se configura a ASBR1 como vecino interno mediante la loopback y se levanta address-family vpnv4, donde se activan la vecindad, el envío de comunidades extendidas y next-hop-self. Mientras tanto, address-family ipv4 vrf EMPRESA se reserva para la relación con CE1 (Cisco, 2020a). De este modo, PE1 mantiene su función de entrada del servicio en la red del proveedor, pero ya no prolonga la VRF hasta la frontera, sino que participa en una arquitectura donde

el servicio del cliente se sostiene con VPNv4 dentro del AS y con MPLS en el transporte interno.

5.6.2. Router ASBR1

```
conf term
hostname ASBR1

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65002:1
route-target import 65001:1

int lo0
ip address 11.11.11.11 255.255.255.255

int e0/0
description ASBR1-PE1
ip address 192.168.1.2 255.255.255.252
mpls ip
no shut
int e0/1
description ASBR1-ASBR2
ip address 100.0.0.1 255.255.255.252
mpls ip
no shut

router ospf 10
router-id 11.11.11.11
network 11.11.11.11 0.0.0.0 area 0
network 192.168.1.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 11.11.11.11
no bgp default route-target filter
neighbor 1.1.1.1 remote-as 65001
neighbor 1.1.1.1 update-source lo0
neighbor 100.0.0.2 remote-as 65002
address-family vpnv4
neighbor 1.1.1.1 activate
neighbor 1.1.1.1 send-community both
neighbor 1.1.1.1 next-hop-self
```

```
neighbor 100.0.0.2 activate
neighbor 100.0.0.2 send-community both
exit-address-family

end
```

En ASBR1 también se reutilizan varias ideas ya vistas en el apartado 5.5.3, ya que el router combina funciones de borde inter-AS y de tránsito interno del proveedor. Se activan ip cef, mpls label protocol ldp, una Loopback0 propia y dos interfaces globales con mpls ip, una hacia PE1 y otra hacia ASBR2 (Cisco, 2022b). A diferencia del escenario A, aquí la interfaz entre proveedores ya no queda dentro de la VRF EMPRESA, sino que se convierte en un enlace global de transporte MPLS. Esto es coherente con el escenario B, donde la frontera inter-AS no mantiene una relación VRF a VRF, sino que pasa a intercambiar información VPNv4 entre ASBR (Cisco, 2022c).

Esta lógica se refleja con más claridad en BGP. En router bgp 65001 aparece no bgp default route-target filter, que desactiva el filtrado automático de comunidades route-target, y se configuran dos vecinos, uno interno con PE1 usando la loopback y otro externo con ASBR2 (Cisco, 2022c). Después, en address-family vpnv4, se activan ambas relaciones y se habilita send-community both, de manera que las comunidades extendidas y la información del servicio puedan circular correctamente entre vecinos (Cisco, 2020a). Así, ASBR1 deja de limitarse a trasladar una VRF hacia el proveedor remoto y pasa a convertirse en el punto donde se intercambian rutas VPNv4 etiquetadas, que es precisamente la esencia del escenario B (Cisco, 2022c).

5.6.3. Router ASBR2

```
conf term
hostname ASBR2

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65001:1
route-target import 65002:1
```

```

int lo0
ip address 22.22.22.22 255.255.255.255

int e0/0
description ASBR2-ASBR1
ip address 100.0.0.2 255.255.255.252
mpls ip
no shut

int e0/1
description ASBR2-PE2
ip address 192.168.2.2 255.255.255.252
mpls ip
no shut

router ospf 10
router-id 22.22.22.22
network 22.22.22.22 0.0.0.0 area 0
network 192.168.2.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 22.22.22.22
no bgp default route-target filter
neighbor 2.2.2.2 remote-as 65002
neighbor 2.2.2.2 update-source lo0
neighbor 100.0.0.1 remote-as 65001
address-family vpnv4
neighbor 2.2.2.2 activate
neighbor 2.2.2.2 send-community both
neighbor 2.2.2.2 next-hop-self
neighbor 100.0.0.1 activate
neighbor 100.0.0.1 send-community both
exit-address-family

end

```

La configuración de ASBR2 sigue esencialmente la misma lógica que la de ASBR1, por lo que no es necesario repetirla con el mismo nivel de detalle. Vuelven a aparecer ip cef, mpls label protocol ldp, una Loopback0, dos interfaces globales con mpls ip y el proceso router ospf 10 para integrar el router en el pequeño núcleo del AS 65002. En BGP se repite también la misma estructura general, con no bgp default route-target filter, un vecino interno con PE2, un vecino externo con ASBR1 y la activación de ambos dentro de address-family vpnv4. Por tanto, ASBR2 debe entenderse como la réplica funcional de ASBR1 en el extremo derecho de la topología, cambiando únicamente el direccionamiento, el router-id y el contexto del AS 65002.

5.6.4. Router PE2

```
conf term
hostname PE2

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65001:1

int lo0
ip address 2.2.2.2 255.255.255.255

int e0/0
description PE2-ASBR2
ip address 192.168.2.1 255.255.255.252
mpls ip
no shut

int e0/1
description PE2-CE2
ip vrf forwarding EMPRESA
ip address 10.2.2.1 255.255.255.252
no shut

router ospf 20
router-id 2.2.2.2
network 2.2.2.2 0.0.0.0 area 0
network 192.168.2.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 2.2.2.2
neighbor 22.22.22.22 remote-as 65002
neighbor 22.22.22.22 update-source lo0
address-family vpnv4
neighbor 22.22.22.22 activate
neighbor 22.22.22.22 send-community both
neighbor 22.22.22.22 next-hop-self
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.2.2.2 remote-as 200
neighbor 10.2.2.2 activate
exit-address-family

end
```

En PE2 se mantiene la misma lógica ya explicada para PE1, adaptada al lado derecho del escenario. Se vuelve a definir la VRF EMPRESA, se conserva la interfaz hacia CE2 dentro de esa VRF y la interfaz hacia ASBR2 se configura como enlace global con mpls ip, junto con ip cef, mpls label protocol ldp, una Loopback0 y router ospf 20 para la conectividad interna del proveedor. Después, en router bgp 65002, se establece la vecindad interna con ASBR2 usando la loopback, se activa address-family vpnv4 con send-community both y next-hop-self, y se mantiene address-family ipv4 vrf EMPRESA para la relación con CE2. En consecuencia, PE2 debe entenderse como la réplica funcional de PE1, manteniendo la misma estructura y cambiando únicamente el direccionamiento, el número de sistema autónomo y el lado de la topología en el que se ubica.

5.7. Escenario B con tres routers

En este apartado se recoge la configuración de la versión ampliada del escenario B, en la que cada sistema autónomo de proveedor incorpora un router interno de tránsito. Al igual que en los apartados anteriores, CE1 y CE2 no se vuelven a desarrollar, ya que su configuración fue explicada en el apartado 5.3 y se mantiene igual en este escenario. La diferencia respecto al 5.6 está en que la topología ya no queda reducida a PE-ASBR, sino que pasa a tener una estructura PE-P-ASBR en cada proveedor. Por tanto, aparecen de nuevo elementos propios del núcleo, como OSPF, LDP/MPLS, interfaces globales con mpls ip y sesiones VPNv4 internas entre PE y ASBR. Cisco explica precisamente que, en una red MPLS VPN, el core del proveedor debe tener conectividad interna mediante un IGP, anunciar las loopbacks y habilitar MPLS en las interfaces del núcleo (Cisco, 2022b).

La lógica propia del escenario B se mantiene igual que en el apartado anterior: el intercambio entre sistemas autónomos no se realiza mediante una relación VRF a VRF, como ocurría en el escenario A, sino mediante el intercambio de rutas VPNv4 entre los ASBR (Cisco, 2022c). Por eso, aunque visualmente esta versión se parezca bastante al escenario A con núcleo interno, internamente responde a otra lógica. En este caso, los ASBR no están configurados con interfaces dentro de la VRF para la frontera inter-AS, sino que trabajan sobre enlaces globales con MPLS y establecen sesiones BGP en la familia VPNv4. Esta diferencia es la que justifica tratar esta configuración como un apartado propio dentro del anexo.

5.7.1. Router PE1

```
conf term
hostname PE1

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65001:1
route-target import 65002:1

int Loopback0
ip address 1.1.1.1 255.255.255.255

int e0/0
description PE1-CE1
ip vrf forwarding EMPRESA
ip address 10.1.1.1 255.255.255.252
no shut

int e0/1
description PE1-P1
ip address 11.11.11.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 1.1.1.1
network 1.1.1.1 0.0.0.0 area 0
network 11.11.11.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 1.1.1.1
no bgp default ipv4-unicast
neighbor 2.2.2.2 remote-as 65001
neighbor 2.2.2.2 update-source Loopback0
address-family vpnv4
neighbor 2.2.2.2 activate
neighbor 2.2.2.2 send-community both
neighbor 2.2.2.2 next-hop-self
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.1.1.2 remote-as 100
neighbor 10.1.1.2 activate
exit-address-family
```

end

En PE1, parte de la configuración se mantiene muy próxima a la explicada en los apartados anteriores. El router sigue actuando como borde del proveedor frente a CE1, por lo que se crea la VRF EMPRESA, se define el rd 65001:1 y se establecen los route-target correspondientes. La interfaz e0/0, conectada hacia el cliente, se asocia a esa VRF mediante `ip vrf forwarding EMPRESA` y recibe la dirección 10.1.1.1/30. La diferencia respecto a una topología más simple aparece en la interfaz e0/1, que ya no enlaza directamente con el ASBR, sino con P1. Por eso se configura como interfaz global del núcleo, con la dirección 11.11.11.1/30 y el comando `mpls ip`, permitiendo que PE1 participe en el transporte MPLS interno del proveedor. Cisco describe esta separación entre interfaces hacia CE dentro de la VRF e interfaces de core con MPLS como parte habitual de una MPLS L3VPN (Cisco, 2022b).

Después, PE1 se integra en el núcleo del AS 65001 mediante `router ospf 1`, anunciando su Loopback0 y el enlace hacia P1 (Moy, 1998). Esa loopback se utiliza también como origen de la sesión BGP interna con ASBR1, configurado como vecino 2.2.2.2. En `address-family vpnv4` se activa esa vecindad, se habilita `send-community both` y se aplica `next-hop-self`, mientras que la familia `address-family ipv4 vrf EMPRESA` queda reservada para la relación con CE1 (Cisco, 2020a). De esta forma, PE1 mantiene su función de entrada del servicio del cliente, pero el intercambio de rutas VPN dentro del proveedor se realiza mediante MP-BGP y no mediante una extensión directa de la VRF hasta la frontera.

5.7.2. Router P1

```
conf term
hostname P1

ip cef
mpls label protocol ldp

int Loopback0
ip address 11.11.11.11 255.255.255.255

int e0/0
description P1-PE1
```

```

ip address 11.11.11.2 255.255.255.252
mpls ip
no shut

int e0/1
description P1-ASBR1
ip address 12.12.12.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 11.11.11.11
network 11.11.11.11 0.0.0.0 area 0
network 11.11.11.0 0.0.0.3 area 0
network 12.12.12.0 0.0.0.3 area 0

end

```

El router P1 aparece como nodo interno del AS 65001 y su función es proporcionar tránsito entre PE1 y ASBR1. A diferencia de los PE y ASBR, no se configura ninguna VRF ni ningún proceso BGP, porque este equipo no participa directamente en el servicio del cliente. Su configuración se centra en el transporte interno, se activa ip cef, se define mpls label protocol ldp, se crea una Loopback0 con la dirección 11.11.11.11/32 y se habilita mpls ip en las interfaces hacia PE1 y ASBR1 (Cisco, 2022b). Posteriormente, mediante router ospf 1, se anuncian la loopback y los enlaces del core (Moy, 1998). De este modo, P1 permite que la red del proveedor deje de ser una conexión directa entre borde y frontera y pase a comportarse como un pequeño núcleo MPLS.

5.7.3. Router ASBR1

```

conf term
hostname ASBR1

ip cef
mpls label protocol ldp

int Loopback0
ip address 2.2.2.2 255.255.255.255

int e0/0
description ASBR1-P1
ip address 12.12.12.2 255.255.255.252
mpls ip

```

```

no shut

int e0/1
description ASBR1-ASBR2
ip address 100.0.0.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 2.2.2.2
network 2.2.2.2 0.0.0.0 area 0
network 12.12.12.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 2.2.2.2
no bgp default route-target filter
no bgp default ipv4-unicast
neighbor 1.1.1.1 remote-as 65001
neighbor 1.1.1.1 update-source Loopback0
neighbor 100.0.0.2 remote-as 65002
address-family vpnv4
neighbor 1.1.1.1 activate
neighbor 1.1.1.1 send-community both
neighbor 1.1.1.1 next-hop-self
neighbor 100.0.0.2 activate
neighbor 100.0.0.2 send-community both
exit-address-family

end

```

En ASBR1, la configuración cambia de forma importante respecto al escenario A. Aquí el ASBR ya no mantiene una interfaz inter-AS dentro de la VRF del cliente, sino que se integra en el núcleo del proveedor y participa en el intercambio VPNv4 del escenario B (Cisco, 2022c). Por eso se activan ip cef, mpls label protocol ldp, una Loopback0 con la dirección 2.2.2.2/32 y dos interfaces globales con mpls ip, una hacia P1 y otra hacia ASBR2. Además, router ospf 1 anuncia la loopback y el enlace hacia P1, permitiendo que ASBR1 sea alcanzable dentro del AS 65001 a través del núcleo interno. Esta parte mantiene la lógica de un core MPLS, donde el IGP proporciona alcance interno y MPLS se habilita en los enlaces de transporte (Cisco, 2022b).

La parte más característica de ASBR1 aparece en BGP. Se configura router bgp 65001, se desactiva el filtrado automático de route-target mediante no bgp default route-target filter y también se utiliza no bgp default ipv4-unicast para evitar activar por defecto

la familia IPv4 unicast (Cisco, 2022c). Después se definen dos vecinos, PE1, mediante la loopback 1.1.1.1, y ASBR2, mediante el enlace inter-AS 100.0.0.2. Ambos vecinos se activan dentro de address-family vpnv4, con send-community both, de forma que la información VPN y sus comunidades extendidas puedan intercambiarse correctamente (Cisco, 2020a). Así, ASBR1 actúa como el punto que conecta el núcleo interno del AS 65001 con el proveedor remoto, pero lo hace mediante rutas VPNv4 y no mediante una VRF extendida hasta la frontera.

5.7.4. Router ASBR2

```
conf term
hostname ASBR2

ip cef
mpls label protocol ldp

int Loopback0
ip address 3.3.3.3 255.255.255.255

int e0/0
description ASBR2-ASBR1
ip address 100.0.0.2 255.255.255.252
mpls ip
no shut

int e0/1
description ASBR2-P2
ip address 23.23.23.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 3.3.3.3
network 3.3.3.3 0.0.0.0 area 0
network 23.23.23.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 3.3.3.3
no bgp default route-target filter
no bgp default ipv4-unicast
neighbor 4.4.4.4 remote-as 65002
neighbor 4.4.4.4 update-source Loopback0
neighbor 100.0.0.1 remote-as 65001
address-family vpnv4
neighbor 4.4.4.4 activate
```

```

neighbor 4.4.4.4 send-community both
neighbor 4.4.4.4 next-hop-self
neighbor 100.0.0.1 activate
neighbor 100.0.0.1 send-community both
exit-address-family

end

```

La configuración de ASBR2 sigue la misma lógica que la de ASBR1, adaptada al AS 65002. De nuevo se activan ip cef, mpls label protocol ldp, una Loopback0 propia y dos interfaces globales con mpls ip, una hacia ASBR1 y otra hacia P2. También se configura OSPF para anunciar la loopback y el enlace interno del proveedor, y en BGP se establecen dos vecinos dentro de la familia VPNv4: PE2, usando la loopback 4.4.4.4, y ASBR1, mediante el enlace 100.0.0.1. Por tanto, ASBR2 debe entenderse como la réplica funcional de ASBR1 en el extremo derecho de la topología, cambiando únicamente el direccionamiento, el router-id y el sistema autónomo al que pertenece.

5.7.5. Router P2

```

conf term
hostname P2

ip cef
mpls label protocol ldp

int Loopback0
ip address 22.22.22.22 255.255.255.255

int e0/0
description P2-ASBR2
ip address 23.23.23.2 255.255.255.252
mpls ip
no shut

int e0/1
description P2-PE2
ip address 24.24.24.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 22.22.22.22
network 22.22.22.22 0.0.0.0 area 0
network 23.23.23.0 0.0.0.3 area 0

```

```
network 24.24.24.0 0.0.0.3 area 0

end
```

El router P2 cumple en el AS 65002 el mismo papel que P1 en el AS 65001. Su configuración se limita al plano de transporte, por lo que no aparecen ni VRF ni BGP. Se activan ip cef y mpls label protocol ldp, se define una Loopback0 con la dirección 22.22.22.22/32 y se configuran las interfaces hacia ASBR2 y PE2 con mpls ip. Después, mediante router ospf 1, se anuncian la loopback y los enlaces internos del core. Con esta configuración, P2 permite que el lado derecho del proveedor tenga también un núcleo interno entre el borde y la frontera, manteniendo una estructura simétrica respecto al AS 65001.

5.7.6. Router PE2

```
conf term
hostname PE2

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65002:1
route-target import 65001:1

int Loopback0
ip address 4.4.4.4 255.255.255.255

int e0/0
description PE2-P2
ip address 24.24.24.2 255.255.255.252
mpls ip
no shut

int e0/1
description PE2-CE2
ip vrf forwarding EMPRESA
ip address 10.2.2.1 255.255.255.252
no shut

router ospf 1
```

```

router-id 4.4.4.4
network 4.4.4.4 0.0.0.0 area 0
network 24.24.24.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 4.4.4.4
no bgp default ipv4-unicast
neighbor 3.3.3.3 remote-as 65002
neighbor 3.3.3.3 update-source Loopback0
address-family vpnv4
neighbor 3.3.3.3 activate
neighbor 3.3.3.3 send-community both
neighbor 3.3.3.3 next-hop-self
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.2.2.2 remote-as 200
neighbor 10.2.2.2 activate
exit-address-family

end

```

En PE2 se repite la misma lógica que en PE1, pero adaptada al lado derecho del escenario. Se crea la VRF EMPRESA, se configuran sus route-target, se mantiene la interfaz hacia CE2 dentro de la VRF y se utiliza la interfaz hacia P2 como enlace global de núcleo con mpls ip. También se define una Loopback0 con la dirección 4.4.4.4/32 y se activa OSPF para integrar el equipo dentro del AS 65002. En BGP, PE2 establece una sesión VPNv4 con ASBR2 usando la loopback 3.3.3.3, mientras que la familia address-family ipv4 vrf EMPRESA se reserva para la relación con CE2. Por tanto, PE2 funciona como la réplica de PE1 en el otro extremo de la topología, manteniendo el mismo papel de borde del proveedor y cambiando únicamente el direccionamiento y el sistema autónomo.

5.8. Escenario C con dos routers

En este apartado se recoge la configuración de la versión reducida del escenario C, formada por un router PE y un router ASBR en cada sistema autónomo de proveedor. Como en los apartados anteriores, CE1 y CE2 no se vuelven a desarrollar, ya que su configuración fue explicada en el apartado 5.3 y se mantiene prácticamente igual. La diferencia principal respecto a los escenarios A y B está en la forma de separar el servicio VPN y el transporte. En este caso, las rutas VPNv4 no se intercambian entre ASBR, sino

directamente entre los routers PE mediante una sesión eBGP multihop, mientras que los ASBR se encargan de transportar rutas IPv4 etiquetadas hacia las loopbacks de los PE. Esta lógica coincide con el procedimiento descrito en RFC 4364, donde los ASBR no mantienen ni distribuyen rutas VPN-IPv4, sino rutas IPv4 etiquetadas que permiten construir un camino conmutado por etiquetas entre el PE de entrada y el PE de salida (Rosen & Rekhter, 2006).

Además, en esta implementación se utiliza la lógica de RFC 3107 mediante send-label, ya que las rutas IPv4 de transporte deben ir asociadas a una etiqueta MPLS. RFC 3107 describe precisamente cómo BGP puede transportar información de etiquetas junto con las rutas que anuncia, lo que encaja con el uso de neighbor ... send-label en las sesiones PE-ASBR y ASBR-ASBR (Rekhter & Rosen, 2001). Por ello, aunque la topología física de este escenario pueda parecer parecida a la del escenario B con dos routers, su funcionamiento interno es distinto: en el escenario B los ASBR intercambian VPNv4, mientras que en el escenario C los ASBR solo ayudan a construir el transporte etiquetado entre los PE.

5.8.1. Router PE1

```
conf term
hostname PE1

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65001:1
route-target import 65002:1

int Loopback0
ip address 1.1.1.1 255.255.255.255

int e0/0
description PE1-CE1
ip vrf forwarding EMPRESA
ip address 10.1.1.1 255.255.255.252
no shut

int e0/1
```

```

description PE1-ASBR1
ip address 192.168.1.1 255.255.255.252
mpls ip
no shut

router ospf 10
router-id 1.1.1.1
network 1.1.1.1 0.0.0.0 area 0
network 192.168.1.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 1.1.1.1
! PE-to-ASBR iBGP transport using RFC3107 send-label
neighbor 11.11.11.11 remote-as 65001
neighbor 11.11.11.11 update-source Loopback0
! Direct PE-to-PE VPNv4 multihop
neighbor 2.2.2.2 remote-as 65002
neighbor 2.2.2.2 update-source Loopback0
neighbor 2.2.2.2 ebgp-multihop 5
address-family ipv4
neighbor 11.11.11.11 activate
neighbor 11.11.11.11 send-label
exit-address-family
address-family vpnv4
neighbor 2.2.2.2 activate
neighbor 2.2.2.2 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.1.1.2 remote-as 100
neighbor 10.1.1.2 activate
exit-address-family

router bgp 65001
address-family ipv4
no neighbor 2.2.2.2 activate

end

clear ip bgp 2.2.2.2
clear ip bgp 11.11.11.11

```

En PE1, una parte de la configuración se mantiene parecida a la ya vista en otros escenarios, ya que el router sigue actuando como borde del proveedor frente a CE1. Por eso se define la VRF EMPRESA, se configuran sus route-target, se asocia la interfaz e0/0 a esa VRF y se mantiene la relación BGP con CE1 dentro de address-family ipv4 vrf EMPRESA. La interfaz e0/1, en cambio, se configura como enlace global hacia ASBR1, con mpls ip, y se activa OSPF para anunciar tanto la loopback de PE1 como el enlace

interno hacia el ASBR. Esta parte conserva la lógica básica de una red MPLS VPN, donde el PE mantiene la VRF del cliente, pero el enlace hacia el núcleo del proveedor pertenece a la tabla global y participa en el transporte MPLS (Cisco, 2022b).

La diferencia más importante aparece en la configuración BGP. Por un lado, PE1 establece una sesión con ASBR1 mediante la loopback 11.11.11.11 y activa esa relación dentro de address-family ipv4 con send-label, de forma que pueda intercambiar rutas IPv4 etiquetadas para el transporte (Rekhter & Rosen, 2001). Por otro lado, PE1 establece una sesión directa VPNv4 multihop con PE2, usando neighbor 2.2.2.2 remote-as 65002, update-source Loopback0 y ebgp-multihop 5. Después, esa vecindad se activa dentro de address-family vpnv4 con send-community both (Rosen & Rekhter, 2006).

Además, se incluye el comando no neighbor 2.2.2.2 activate dentro de la familia IPv4 para evitar que el vecino PE2 quede activo en la familia de direcciones incorrecta, ya que su función en este escenario es intercambiar rutas VPNv4 y no rutas IPv4 unicast. Finalmente, se ejecutan clear ip bgp 2.2.2.2 y clear ip bgp 11.11.11.11 para reiniciar las sesiones BGP afectadas y forzar que los cambios aplicados en las familias de direcciones se negocien de nuevo. Así, PE1 combina dos funciones separadas, mantiene la relación con el cliente dentro de la VRF y, al mismo tiempo, participa en una arquitectura donde el servicio VPN se intercambia directamente entre PE, mientras que el transporte se apoya en rutas etiquetadas distribuidas por BGP.

5.8.2. Router ASBR1

```
conf term
hostname ASBR1

ip cef
mpls label protocol ldp

int Loopback0
ip address 11.11.11.11 255.255.255.255

int e0/0
description ASBR1-PE1
ip address 192.168.1.2 255.255.255.252
mpls ip
no shut
```

```

int e0/1
description ASBR1-ASBR2
ip address 100.0.0.1 255.255.255.252
mpls ip
no shut

router ospf 10
router-id 11.11.11.11
network 11.11.11.11 0.0.0.0 area 0
network 192.168.1.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 11.11.11.11
no bgp default route-target filter
! iBGP-LU toward local PE
neighbor 1.1.1.1 remote-as 65001
neighbor 1.1.1.1 update-source Loopback0
! eBGP-LU toward remote ASBR
neighbor 100.0.0.2 remote-as 65002
address-family ipv4
neighbor 1.1.1.1 activate
neighbor 1.1.1.1 next-hop-self
neighbor 1.1.1.1 send-label
neighbor 100.0.0.2 activate
neighbor 100.0.0.2 send-label
network 1.1.1.1 mask 255.255.255.255
network 11.11.11.11 mask 255.255.255.255
exit-address-family

end

```

En ASBR1, la configuración es bastante distinta a la de los escenarios A y B, porque aquí no aparece ninguna VRF EMPRESA ni una familia VPNv4 destinada a transportar rutas del cliente. Su función se centra en actuar como router frontera de transporte entre el AS 65001 y el AS 65002 (Rosen & Rekhter, 2006). Para ello, se activan ip cef, mpls label protocol ldp, una Loopback0 con la dirección 11.11.11.11/32 y dos interfaces globales con mpls ip: una hacia PE1 y otra hacia ASBR2. También se configura OSPF para anunciar la loopback y el enlace interno hacia PE1, permitiendo que el ASBR sea alcanzable dentro de su propio sistema autónomo (Cisco, 2022b).

La parte clave se encuentra en router bgp 65001. ASBR1 establece una sesión interna con PE1 usando la loopback y una sesión externa con ASBR2 a través del enlace 100.0.0.2. Ambas relaciones se activan dentro de address-family ipv4 y utilizan send-

label, lo que permite asociar etiquetas MPLS a las rutas IPv4 anunciadas (Rekhter & Rosen, 2001). Además, se anuncian las loopbacks 1.1.1.1/32 y 11.11.11.11/32, que son necesarias para construir el alcance etiquetado hacia los extremos del transporte. En consecuencia, ASBR1 no actúa aquí como distribuidor de rutas VPNv4, sino como pieza intermedia que permite que los PE remotos puedan alcanzarse mediante rutas IPv4 etiquetadas, siguiendo la lógica de RFC 3107 y del procedimiento inter-AS descrito en RFC 4364 (Rosen & Rekhter, 2006).

5.8.3. Router ASBR2

```
conf term
hostname ASBR2

ip cef
mpls label protocol ldp

int Loopback0
ip address 22.22.22.22 255.255.255.255

int e0/0
description ASBR2-ASBR1
ip address 100.0.0.2 255.255.255.252
mpls ip
no shut

int e0/1
description ASBR2-PE2
ip address 192.168.2.2 255.255.255.252
mpls ip
no shut

router ospf 20
router-id 22.22.22.22
network 22.22.22.22 0.0.0.0 area 0
network 192.168.2.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 22.22.22.22
no bgp default route-target filter
! iBGP-LU toward local PE
neighbor 2.2.2.2 remote-as 65002
neighbor 2.2.2.2 update-source Loopback0
! eBGP-LU toward remote ASBR
neighbor 100.0.0.1 remote-as 65001
```

```

address-family ipv4
neighbor 2.2.2.2 activate
neighbor 2.2.2.2 next-hop-self
neighbor 2.2.2.2 send-label
neighbor 100.0.0.1 activate
neighbor 100.0.0.1 send-label
network 2.2.2.2 mask 255.255.255.255
network 22.22.22.22 mask 255.255.255.255
exit-address-family

end

```

La configuración de ASBR2 sigue la misma lógica que la de ASBR1, pero adaptada al AS 65002. De nuevo no se configura ninguna VRF ni ninguna familia VPNv4, ya que este router no mantiene las rutas del servicio del cliente. Su función consiste en participar en el transporte etiquetado entre sistemas autónomos, utilizando interfaces globales con mpls ip, OSPF para alcanzar su propia loopback y el enlace interno hacia PE2, y BGP IPv4 con send-label tanto hacia PE2 como hacia ASBR1. Las diferencias se encuentran únicamente en el direccionamiento, el router-id, la loopback 22.22.22.22/32 y las redes anunciadas dentro de BGP. Por tanto, ASBR2 debe entenderse como la réplica funcional de ASBR1 en el extremo derecho de la topología, manteniendo la misma misión de transporte y no de distribución directa del servicio VPN.

5.8.4. Router PE2

```

conf term
hostname PE2

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65002:1
route-target import 65001:1

int Loopback0
ip address 2.2.2.2 255.255.255.255

int e0/0
description PE2-ASBR2
ip address 192.168.2.1 255.255.255.252

```

```

mpls ip
no shut

int e0/1
description PE2-CE2
ip vrf forwarding EMPRESA
ip address 10.2.2.1 255.255.255.252
no shut

router ospf 20
router-id 2.2.2.2
network 2.2.2.2 0.0.0.0 area 0
network 192.168.2.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 2.2.2.2
! PE-to-ASBR iBGP transport using RFC3107 send-label
neighbor 22.22.22.22 remote-as 65002
neighbor 22.22.22.22 update-source Loopback0
! Direct PE-to-PE VPNv4 multihop
neighbor 1.1.1.1 remote-as 65001
neighbor 1.1.1.1 update-source Loopback0
neighbor 1.1.1.1 ebgp-multihop 5
address-family ipv4
neighbor 22.22.22.22 activate
neighbor 22.22.22.22 send-label
exit-address-family
address-family vpnv4
neighbor 1.1.1.1 activate
neighbor 1.1.1.1 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.2.2.2 remote-as 200
neighbor 10.2.2.2 activate
exit-address-family

router bgp 65002
address-family ipv4
no neighbor 1.1.1.1 activate

end

clear ip bgp 1.1.1.1
clear ip bgp 22.22.22.22

```

En PE2 se repite la misma estructura general que en PE1, adaptada al lado derecho del escenario. Se define la VRF EMPRESA, se configura la interfaz hacia CE2 dentro de esa VRF y se mantiene la relación BGP con el cliente mediante address-family ipv4 vrf

EMPRESA. La interfaz hacia ASBR2 queda como enlace global con mpls ip, y OSPF anuncia tanto la loopback 2.2.2.2/32 como el enlace interno hacia el ASBR. En BGP, PE2 establece una relación IPv4 etiquetada con ASBR2 mediante send-label y, al mismo tiempo, una sesión VPNv4 multihop directa con PE1 usando la loopback remota 1.1.1.1. Por ello, PE2 debe entenderse como la réplica funcional de PE1, cambiando únicamente el direccionamiento, el sistema autónomo y el lado de la topología en el que se ubica.

5.9. Escenario C con tres routers

En este apartado se recoge la configuración de la versión ampliada del escenario C, en la que cada proveedor incorpora un router interno de tránsito entre el PE y el ASBR. Como en los apartados anteriores, CE1 y CE2 no se vuelven a desarrollar, ya que su configuración fue explicada en el apartado 5.3 y se mantiene igual en este escenario. La diferencia respecto al apartado 5.8 está en que la topología deja de ser PE-ASBR y pasa a organizarse como PE-P-ASBR en cada sistema autónomo. Por tanto, aparecen de nuevo los routers P1 y P2, que no participan directamente en la VPN del cliente, sino en el transporte interno del proveedor. RFC 4364 indica que los routers P no instalan rutas VPN-IPv4, lo que encaja con el papel asignado a estos nodos dentro del escenario (Rosen & Rekhter, 2006).

La lógica del escenario C se mantiene igual que en la versión de dos routers, las rutas VPNv4 se intercambian directamente entre los routers PE mediante una sesión eBGP multihop, mientras que los ASBR se encargan de distribuir rutas IPv4 etiquetadas para construir el transporte entre sistemas autónomos. RFC 4364 describe esta alternativa indicando que los ASBR no mantienen ni distribuyen rutas VPN-IPv4, sino rutas IPv4 etiquetadas hacia los PE, permitiendo que los PE de distintos AS establezcan sesiones eBGP multihop para intercambiar las rutas VPN-IPv4 (Rosen & Rekhter, 2006).

5.9.1. Router PE1

```
conf term
hostname PE1

ip cef
mpls label protocol ldp
```

```

ip vrf EMPRESA
rd 65001:1
route-target export 65001:1
route-target import 65001:1
route-target import 65002:1

int Loopback0
ip address 1.1.1.1 255.255.255.255

int e0/0
description PE1-CE1
ip vrf forwarding EMPRESA
ip address 10.1.1.1 255.255.255.252
no shut

int e0/1
description PE1-P1
ip address 11.11.11.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 1.1.1.1
network 1.1.1.1 0.0.0.0 area 0
network 11.11.11.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 1.1.1.1
no bgp default ipv4-unicast
neighbor 2.2.2.2 remote-as 65001
neighbor 2.2.2.2 update-source Loopback0
neighbor 4.4.4.4 remote-as 65002
neighbor 4.4.4.4 update-source Loopback0
neighbor 4.4.4.4 ebgp-multihop 10
address-family ipv4
neighbor 2.2.2.2 activate
neighbor 2.2.2.2 send-label
network 1.1.1.1 mask 255.255.255.255
exit-address-family
address-family vpnv4
neighbor 4.4.4.4 activate
neighbor 4.4.4.4 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.1.1.2 remote-as 100
neighbor 10.1.1.2 activate
exit-address-family

router bgp 65001

```

```
address-family ipv4
no neighbor 4.4.4.4 activate
```

```
end
```

```
clear ip bgp 2.2.2.2
clear ip bgp 4.4.4.4
```

En PE1, la configuración mantiene una base similar a la del apartado 5.8.1, ya que el router sigue actuando como borde del proveedor frente a CE1. Se crea la VRF EMPRESA, se definen el rd y los route-target, y la interfaz e0/0 queda asociada a esa VRF mediante `ip vrf forwarding EMPRESA`. Esta interfaz mantiene la dirección 10.1.1.1/30 hacia el cliente, mientras que la interfaz e0/1 ya no enlaza directamente con ASBR1, sino con P1. Por ello, e0/1 se configura como enlace global de núcleo con `mpls ip`, siguiendo la misma separación entre acceso del cliente y transporte interno del proveedor que ya se ha visto en apartados anteriores. Cisco muestra esta lógica al asociar las interfaces de cliente a una VRF y habilitar `mpls ip` en las interfaces de núcleo hacia routers P (Cisco, 2022b).

La parte más importante de PE1 está en BGP. Dentro de router bgp 65001, se configura una relación IPv4 etiquetada con ASBR1 mediante la loopback 2.2.2.2, activando el vecino en `address-family ipv4` y aplicando `send-label`. Además, se anuncia la propia loopback 1.1.1.1/32, que será necesaria para que el extremo remoto pueda alcanzar a PE1. Paralelamente, se configura una sesión VPNv4 multihop directa con PE2, identificado por la loopback 4.4.4.4, usando `update-source Loopback0` y `ebgp-multihop 10`. De esta forma, PE1 separa claramente el transporte, basado en rutas IPv4 etiquetadas, del servicio VPN, que se intercambia directamente con el PE remoto. RFC 3107 permite entender esta parte porque describe cómo BGP puede transportar información de etiquetas asociada a una ruta (Rekhter & Rosen, 2001).

5.9.2. Router P1

```
conf term
hostname P1

ip cef
mpls label protocol ldp
```

```

int Loopback0
ip address 11.11.11.11 255.255.255.255

int e0/0
description P1-PE1
ip address 11.11.11.2 255.255.255.252
mpls ip
no shut
int e0/1
description P1-ASBR1
ip address 12.12.12.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 11.11.11.11
network 11.11.11.11 0.0.0.0 area 0
network 11.11.11.0 0.0.0.3 area 0
network 12.12.12.0 0.0.0.3 area 0

end

```

El router P1 cumple la función de nodo interno de tránsito dentro del AS 65001. Su configuración es prácticamente igual a la de otros routers P ya explicados, por lo que no incorpora VRF ni proceso BGP. En este caso se activan ip cef y mpls label protocol ldp, se define una Loopback0 con la dirección 11.11.11.11/32 y se configuran dos interfaces de núcleo con mpls ip, una hacia PE1 y otra hacia ASBR1. Después, mediante router ospf 1, se anuncian la loopback y los enlaces internos 11.11.11.0/30 y 12.12.12.0/30. Así, P1 permite que el tráfico avance entre el borde del proveedor y la frontera inter-AS sin participar directamente en las rutas VPN del cliente, actuando únicamente como parte del núcleo MPLS interno (Cisco, 2022b).

5.9.3. Router ASBR1

```

conf term
hostname ASBR1

ip cef
mpls label protocol ldp

int Loopback0
ip address 2.2.2.2 255.255.255.255

int e0/0

```

```

description ASBR1-P1
ip address 12.12.12.2 255.255.255.252
mpls ip
no shut

int e0/1
description ASBR1-ASBR2
ip address 100.0.0.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 2.2.2.2
network 2.2.2.2 0.0.0.0 area 0
network 12.12.12.0 0.0.0.3 area 0

router bgp 65001
bgp router-id 2.2.2.2
no bgp default route-target filter
no bgp default ipv4-unicast
neighbor 1.1.1.1 remote-as 65001
neighbor 1.1.1.1 update-source Loopback0
neighbor 100.0.0.2 remote-as 65002
address-family ipv4
neighbor 1.1.1.1 activate
neighbor 1.1.1.1 next-hop-self
neighbor 1.1.1.1 send-label
neighbor 100.0.0.2 activate
neighbor 100.0.0.2 send-label
exit-address-family

end

```

En ASBR1, la configuración confirma la lógica propia del escenario C. A diferencia de los escenarios A y B, este router no tiene configurada ninguna VRF EMPRESA ni una familia VPNv4 para intercambiar rutas del cliente. Su papel consiste en formar parte del transporte etiquetado entre PE1 y PE2 (Rosen & Rekhter, 2006). Por eso se activan ip cef, mpls label protocol ldp, una Loopback0 con la dirección 2.2.2.2/32 y dos interfaces globales con mpls ip: una hacia P1 y otra hacia ASBR2. También se configura OSPF para anunciar la loopback y el enlace interno hacia P1, integrando el router en el núcleo del AS 65001 (Cisco, 2022b). Esta separación coincide con RFC 4364, que para esta opción indica que los ASBR no mantienen ni distribuyen rutas VPN-IPv4 (Rosen & Rekhter, 2006).

En la parte BGP, ASBR1 establece una sesión interna con PE1 mediante la loopback 1.1.1.1 y una sesión externa con ASBR2 mediante 100.0.0.2. Ambas relaciones se activan dentro de address-family ipv4 y utilizan send-label, de forma que las rutas IPv4 puedan ir acompañadas de etiquetas MPLS. Además, se aplica next-hop-self hacia PE1 para controlar el siguiente salto dentro del AS local. Con esta configuración, ASBR1 no distribuye directamente rutas VPNv4, sino que ayuda a construir el camino etiquetado que permitirá a PE1 y PE2 intercambiar el servicio VPN mediante su sesión multihop. Esta idea encaja con RFC 3107, que define el transporte de etiquetas junto a rutas BGP (Rekhter & Rosen, 2001).

5.9.4. Router ASBR2

```
conf term
hostname ASBR2

ip cef
mpls label protocol ldp

int Loopback0
ip address 3.3.3.3 255.255.255.255

int e0/0
description ASBR2-ASBR1
ip address 100.0.0.2 255.255.255.252
mpls ip
no shut

int e0/1
description ASBR2-P2
ip address 23.23.23.1 255.255.255.252
mpls ip
no shut

router ospf 1
router-id 3.3.3.3
network 3.3.3.3 0.0.0.0 area 0
network 23.23.23.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 3.3.3.3
no bgp default route-target filter
no bgp default ipv4-unicast
neighbor 4.4.4.4 remote-as 65002
neighbor 4.4.4.4 update-source Loopback0
```

```
neighbor 100.0.0.1 remote-as 65001
address-family ipv4
neighbor 4.4.4.4 activate
neighbor 4.4.4.4 next-hop-self
neighbor 4.4.4.4 send-label
neighbor 100.0.0.1 activate
neighbor 100.0.0.1 send-label
exit-address-family

end
```

La configuración de ASBR2 sigue la misma lógica que la de ASBR1, adaptada al AS 65002. De nuevo no se crea ninguna VRF ni se configura una familia VPNv4, ya que este router se utiliza como elemento de transporte y no como punto de intercambio directo del servicio VPN. Se activan ip cef, mpls label protocol ldp, una Loopback0 con la dirección 3.3.3.3/32 y dos interfaces globales con mpls ip, una hacia ASBR1 y otra hacia P2. En BGP, se establece una sesión interna con PE2 y otra externa con ASBR1, ambas dentro de address-family ipv4 con send-label. Por tanto, ASBR2 debe entenderse como la réplica funcional de ASBR1 en el lado derecho de la topología, cambiando únicamente el direccionamiento, el router-id y el sistema autónomo al que pertenece.

5.9.5. Router P2

```
conf term
hostname P2

ip cef
mpls label protocol ldp

int Loopback0
ip address 22.22.22.22 255.255.255.255

int e0/0
description P2-ASBR2
ip address 23.23.23.2 255.255.255.252
mpls ip
no shut
int e0/1
description P2-PE2
ip address 24.24.24.1 255.255.255.252
mpls ip
no shut
```

```

router ospf 1
router-id 22.22.22.22
network 22.22.22.22 0.0.0.0 area 0
network 23.23.23.0 0.0.0.3 area 0
network 24.24.24.0 0.0.0.3 area 0

```

```
end
```

El router P2 desempeña en el AS 65002 la misma función que P1 en el AS 65001. Su configuración se limita al plano de transporte, por lo que no aparecen ni VRF ni BGP. Se activan ip cef y mpls label protocol ldp, se define una Loopback0 con la dirección 22.22.22.22/32 y se configuran dos interfaces globales con mpls ip, una hacia ASBR2 y otra hacia PE2. Después, mediante router ospf 1, se anuncian la loopback y los enlaces internos del core. De este modo, P2 mantiene la simetría del escenario y permite que el proveedor derecho tenga también un núcleo MPLS entre su PE y su ASBR.

5.9.6. Router PE2

```

conf term
hostname PE2

ip cef
mpls label protocol ldp

ip vrf EMPRESA
rd 65002:1
route-target export 65002:1
route-target import 65002:1
route-target import 65001:1

int Loopback0
ip address 4.4.4.4 255.255.255.255

int e0/0
description PE2-P2
ip address 24.24.24.2 255.255.255.252
mpls ip
no shut

int e0/1
description PE2-CE2
ip vrf forwarding EMPRESA
ip address 10.2.2.1 255.255.255.252
no shut

```

```

router ospf 1
router-id 4.4.4.4
network 4.4.4.4 0.0.0.0 area 0
network 24.24.24.0 0.0.0.3 area 0

router bgp 65002
bgp router-id 4.4.4.4
no bgp default ipv4-unicast
neighbor 3.3.3.3 remote-as 65002
neighbor 3.3.3.3 update-source Loopback0
neighbor 1.1.1.1 remote-as 65001
neighbor 1.1.1.1 update-source Loopback0
neighbor 1.1.1.1 ebgp-multihop 10
address-family ipv4
neighbor 3.3.3.3 activate
neighbor 3.3.3.3 send-label
network 4.4.4.4 mask 255.255.255.255
exit-address-family
address-family vpnv4
neighbor 1.1.1.1 activate
neighbor 1.1.1.1 send-community both
exit-address-family
address-family ipv4 vrf EMPRESA
neighbor 10.2.2.2 remote-as 200
neighbor 10.2.2.2 activate
exit-address-family

router bgp 65002
address-family ipv4
no neighbor 1.1.1.1 activate

end

clear ip bgp 1.1.1.1
clear ip bgp 3.3.3.3

```

En PE2 se repite la misma estructura que en PE1, adaptada al extremo derecho del escenario. Se define la VRF EMPRESA, se configura la interfaz hacia CE2 dentro de esa VRF y se mantiene la relación BGP con el cliente mediante address-family ipv4 vrf EMPRESA. La interfaz hacia P2 queda como enlace global con mpls ip, y OSPF anuncia tanto la loopback 4.4.4.4/32 como el enlace interno hacia el router P. En BGP, PE2 establece una sesión IPv4 etiquetada con ASBR2 mediante send-label y, al mismo tiempo, una sesión VPNv4 multihop directa con PE1, usando la loopback remota 1.1.1.1. Por ello, PE2 debe entenderse como la réplica funcional de PE1, cambiando únicamente el direccionamiento, el AS y el lado de la topología en el que se ubica.

6. ACRÓNIMOS

ARP	Address Resolution Protocol
AS	Autonomous System
ASBR	Autonomous System Boundary Router
BDR	Backup Designated Router
BGP	Border Gateway Protocol
CE	Customer Edge Router
DR	Designated Router
eBGP	External Border Gateway Protocol
GNS3	Graphical Network Simulator-3
ICMP	Internet Control Message Protocol
IGP	Interior Gateway Protocol
IOS	Internetwork Operating System
IP	Internet Protocol
IPv4	Internet Protocol version 4
LDP	Label Distribution Protocol
MPLS	Multiprotocol Label Switching
OSPF	Open Shortest Path First
P	Provider Router
PE	Provider Edge Router
QEMU	Quick Emulator
RD	Route Distinguisher
RT	Route Target
TCP	Transmission Control Protocol
TTL	Time To Live
VPN	Virtual Private Network
VPNv4	VPN IPv4 Address Family
VRF	Virtual Routing and Forwarding

7. BIBLIOGRAFÍA

Andersson, L., Minei, I., & Thomas, B. (2007). LDP specification (RFC 5036). RFC Editor. <https://www.rfc-editor.org/info/rfc5036/>

Bates, T., Chandra, R., Katz, D., & Rekhter, Y. (2007). Multiprotocol extensions for BGP-4 (RFC 4760). RFC Editor. <https://www.rfc-editor.org/rfc/rfc4760.html>

Cisco. (2001). Cisco 3600 Series – Cisco IOS Release 12.2 XG release notes. Cisco. https://www.cisco.com/c/en/us/td/docs/ios/12_2/12_2x/12_2xg/release/notes/rn3600xg.html

Cisco. (2016). Configuration and verification of Layer 3 INTER-AS MPLS VPN Option B using IOS and IOS-XR. Cisco. <https://www.cisco.com/c/en/us/support/docs/multiprotocol-label-switching-mpls/mpls/200557-Configuration-and-Verification-of-Layer.html>

Cisco. (2018a). Configuración del servicio MPLS L3VPN en el router PE mediante REST-API (IOS-XE). Cisco. https://www.cisco.com/c/es_mx/support/docs/multiprotocol-label-switching-mpls/mpls/212655-configure-mpls-l3vpn-service-on-pe-route.html

Cisco. (2018b). Configure Inter-AS Option C MPLS VPN with Cisco IOS and Cisco IOS-XR. Cisco. <https://www.cisco.com/c/en/us/support/docs/multiprotocol-label-switching-mpls/mpls/200523-Configuration-and-Verification-of-Layer.html>

Cisco. (2020a). Implementing MPLS Layer 3 VPNs. En L3VPN Configuration Guide for Cisco NCS 5500 Series Routers, IOS XR Release 7.2.x. Cisco. <https://www.cisco.com/c/en/us/td/docs/iosxr/ncs5500/vpn/72x/b-l3vpn-cg-ncs5500-72x/implementing-mpls-layer-3-VPNs.html>

Cisco. (2020b). Configuring MPLS VPN InterAS options. En Multiprotocol Label Switching (MPLS) Configuration Guide, Cisco IOS XE Amsterdam 17.3.x. Cisco. https://www.cisco.com/c/en/us/td/docs/switches/lan/catalyst9300/software/release/17-3/configuration_guide/mpls/b_173_mpls_9300_cg/configuring_mpls_vpn_interas_options.html

Cisco. (2022a). Implementing OSPF. En Routing Configuration Guide for Cisco NCS 5500 Series Routers, IOS XR Release 7.7.x. Cisco. <https://www.cisco.com/c/en/us/td/docs/iosxr/ncs5500/routing/77x/b-routing-cg-ncs5500-77x/implementing-ospf.html>

Cisco. (2022b). Configure a basic MPLS VPN network. Cisco. <https://www.cisco.com/c/en/us/support/docs/multiprotocol-label-switching-mpls/mpls/13733-mpls-vpn-basic.html>

Cisco. (2022c). Configuring MPLS VPN InterAS options. En Multiprotocol Label Switching (MPLS) Configuration Guide, Cisco IOS XE Cupertino 17.8.x. Cisco. https://www.cisco.com/c/en/us/td/docs/switches/lan/catalyst9600/software/release/17-8/configuration_guide/mpls/b_178_mpls_9600_cg/configuring_mpls_vpn_interas_options.html

Cisco. (2024a). Configuring MPLS VPN InterAS options. En MPLS Configuration Guide for Cisco Catalyst 9600 Series Switches. Cisco. https://www.cisco.com/c/en/us/td/docs/switches/lan/catalyst9600/software/release/17-15/configuration_guide/mpls/b_1715_mpls_9600_cg/configuring_mpls_vpn_interas_options.html

Cisco. (2024b). Configuring MPLS VPN InterAS options. En Multiprotocol Label Switching Configuration Guide, Cisco IOS XE 17.17.x (Catalyst 9300 Switches). Cisco. https://www.cisco.com/c/en/us/td/docs/switches/lan/catalyst9300/software/release/17-17/configuration_guide/mpls/b_1717_mpls_9300_cg/configuring_mpls_vpn_interas_options.html

Cisco. (s. f.). What is MPLS? Multiprotocol Label Switching. Cisco.
https://www.cisco.com/en/US/products/ps6557/products_ios_technology_home.html

GNS3. (s. f.-a). Getting started with GNS3. GNS3 Documentation.
<https://docs.gns3.com/docs/>

GNS3. (s. f.-b). Cisco IOS images for Dynamips. GNS3 Documentation.
<https://docs.gns3.com/docs/emulators/cisco-ios-images-for-dynamips/>

GNS3. (s. f.-c). VPCS. GNS3 Documentation.
<https://docs.gns3.com/docs/emulators/vpcs/>

González Ramírez, G. (2022). GNS3 como herramienta para el diseño y desarrollo de prácticas de laboratorio de redes [Trabajo Fin de Grado, Universidad de Valladolid]. UVaDOC. <https://uvadoc.uva.es/handle/10324/57310>

Juniper Networks. (2025). Interprovider VPNs. Juniper Networks.
<https://www.juniper.net/documentation/us/en/software/junos/vpn-l3/topics/topic-map/l3-vpns-interprovider.html>

Linux man-pages project. (2025). ifconfig(8) - Linux manual page. man7.org.
<https://man7.org/linux/man-pages/man8/ifconfig.8.html>

Linux man-pages project. (2026). ping(8) Linux manual page. man7.org.
<https://man7.org/linux/man-pages/man8/ping.8.html>

Microsoft. (2021). Windows Sockets error codes. Microsoft Learn.
<https://learn.microsoft.com/en-us/windows/win32/winsock/windows-sockets-error-codes-2>

Moy, J. (1998). OSPF version 2 (RFC 2328). RFC Editor. <https://www.rfc-editor.org/rfc/rfc2328.html>

Plummer, D. (1982). An Ethernet Address Resolution Protocol: Or Converting Network Protocol Addresses to 48.bit Ethernet Address for Transmission on Ethernet Hardware (RFC 826). RFC Editor. <https://www.rfc-editor.org/rfc/rfc826>

Rekhter, Y., & Rosen, E. (2001). Carrying label information in BGP-4 (RFC 3107). RFC Editor. <https://www.rfc-editor.org/rfc/rfc3107>

Rekhter, Y., Li, T., & Hares, S. (2006). A Border Gateway Protocol 4 (BGP-4) (RFC 4271). RFC Editor. <https://www.rfc-editor.org/rfc/rfc4271.html>

Rodríguez, J. (2019). Desarrollo de un entorno MPLS basado en GNS3 [Trabajo Fin de Grado, Universitat Politècnica de València]. RIUNET. <https://riunet.upv.es/bitstream/handle/10251/127859/Rodr%C3%ADguez%20-%20Desarrollo%20de%20un%20entorno%20MPLS%20basado%20en%20GNS3.pdf>

Rosen, E., & Rekhter, Y. (2006). BGP/MPLS IP virtual private networks (VPNs) (RFC 4364). RFC Editor. <https://www.rfc-editor.org/rfc/rfc4364.html>

Rosen, E., Viswanathan, A., & Callon, R. (2001). Multiprotocol Label Switching architecture (RFC 3031). RFC Editor. <https://www.rfc-editor.org/rfc/rfc3031.html>

Tiny Core Linux Wiki. (2011a). bootlocal.sh and shutdown.sh. Tiny Core Linux Wiki. https://wiki.tinycorelinux.net/doku.php?id=wiki:bootlocal.sh_and_shutdown.sh

Tiny Core Linux Wiki. (2011b). Backup. Tiny Core Linux Wiki. <https://wiki.tinycorelinux.net/doku.php?id=wiki:backup>