



**UNIVERSIDAD TÉCNICA DEL NORTE
FACULTAD DE INGENIERÍA EN CIENCIAS APLICADAS
CARRERA DE TELECOMUNICACIONES**

**INFORME FINAL DEL TRABAJO DE INTEGRACIÓN
CURRICULAR, PROYECTO DE INVESTIGACIÓN**

TEMA:

**“TESTBED DE EVALUACIÓN DE DESEMPEÑO DE LA
TECNOLOGÍA VXLAN EN REDES DE CENTRO DE DATOS CON
ARQUITECTURA SPINE-LEAF”**

**Trabajo de titulación previo a la obtención del título de Ingeniero en
Telecomunicaciones**

Línea de investigación: Desarrollo, aplicación de software y cibersecurity (seguridad cibernética)

AUTOR:

Nupan Guzmán Ederson Germán

DIRECTOR:

Msc. Domínguez Limaico Hernán Mauricio

Ibarra, Ecuador 2025

UNIVERSIDAD TÉCNICA DEL NORTE
BIBLIOTECA UNIVERSITARIA

IDENTIFICACIÓN DE LA OBRA

La Universidad Técnica del Norte dentro del proyecto Repositorio Digital Institucional, determinó la necesidad de disponer de textos completos en formato digital con la finalidad de apoyar los procesos de investigación, docencia y extensión de la Universidad.

Por medio del presente documento dejo sentada mi voluntad de participar en este proyecto, para lo cual pongo a disposición la siguiente información:

DATOS DE CONTACTO			
CÉDULA DE IDENTIDAD:	0401657853		
APELLIDOS Y NOMBRES:	Nupan Guzmán Ederson Germán		
DIRECCIÓN:	IBARRA, PANAMERICA NORTE Y DR. PLUTARCO LARREA		
EMAIL:	egnupang@utn.edu.ec / edersonnupan12@gmail.com		
TELÉFONO FIJO:	062986622	TELF. MOVIL	0963977898

DATOS DE LA OBRA	
TÍTULO:	TESTBED DE EVALUACIÓN DE DESEMPEÑO DE LA TECNOLOGÍA VXLAN EN REDES DE CENTRO DE DATOS CON ARQUITECTURA SPINE-LEAF.
AUTOR (ES):	EDERSON GERMÁN NUPAN GUZMÁN
FECHA: AAAAMMDD	2025/07/31
SOLO PARA TRABAJOS DE INTEGRACIÓN CURRICULAR	
CARRERA/PROGRAMA:	GRADO ☒ POSGRADO ☐
TÍTULO POR EL QUE OPTA:	INGENIERO EN TELECOMUNICACIONES
DIRECTOR:	MSC. HERNÁN MAURICIO DOMÍNGUEZ LIMAICO
ASESOR:	MSC. EDGAR ALBERTO MAYA OLALLA

AUTORIZACIÓN DE USO A FAVOR DE LA UNIVERSIDAD

Yo, NUPAN GUZMÁN EDERSON GERMÁN, con cédula de identidad Nro. 0401657853, en calidad de autor y titular de los derechos patrimoniales de la obra o trabajo de integración curricular descrito anteriormente, hago entrega del ejemplar respectivo en formato digital y autorizo a la Universidad Técnica del Norte, la publicación de la obra en el Repositorio Digital Institucional y uso del archivo digital en la Biblioteca de la Universidad con fines académicos, para ampliar la disponibilidad del material y como apoyo a la educación, investigación y extensión; en concordancia con la Ley de Educación Superior Artículo 144.

Ibarra, a los 31 días del mes de Julio del 2025.

EL AUTOR:



Nupan Guzmán Ederson Germán

CONSTANCIAS

El autor manifiesta que la obra objeto de la presente autorización es original y se la desarrolló, sin violar derechos de autor de terceros, por lo tanto, la obra es original y que es el titular de los derechos patrimoniales, por lo que asume la responsabilidad sobre el contenido de la misma y saldrá en defensa de la Universidad en caso de reclamación por parte de terceros.

Ibarra, a los 31 días, del mes de Julio del 2025.

EL AUTOR:



Nupan Guzmán Ederson Germán

**CERTIFICACIÓN DEL DIRECTOR DEL TRABAJO DE INTEGRACIÓN
CURRICULAR**

Ibarra, 31 de Julio del 2025

MSC. HERNÁN MAURICIO DOMÍNGUEZ LIMAICO

DIRECTOR DEL TRABAJO DE INTEGRACIÓN CURRICULAR

CERTIFICA:

Haber revisado el presente informe final del trabajo de Integración Curricular, el mismo que se ajusta a las normas vigentes de la Universidad Técnica del Norte; en consecuencia, autorizo su presentación para los fines legales pertinentes.

(f) 

Msc. Hernán Mauricio Domínguez Limaico

C.C.: 1002379301

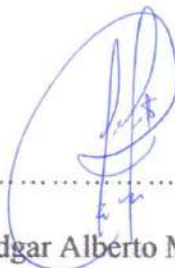
APROBACIÓN DEL COMITÉ CALIFICADOR

El Comité Calificado del trabajo de Integración Curricular “TESTBED DE EVALUACIÓN DE DESEMPEÑO DE LA TECNOLOGÍA VXLAN EN REDES DE CENTRO DE DATOS CON ARQUITECTURA SPINE-LEAF” elaborado por NUPAN GUZMÁN EDERSON GERMÁN, previo a la obtención del título de INGENIERO EN TELECOMUNICACIONES aprueba el presente informe de investigación en nombre de la Universidad Técnica del Norte:

(f) 

Msc. Hernán Mauricio Domínguez Limaico

C.C.: 1002379301

(f) 

Msc. Edgar Alberto Maya Olalla

C.C.: 1002702197

DEDICATORIA

El presente trabajo está dedicado a quienes creyeron en mí, y especialmente a quienes, con su amor incondicional y su ejemplo constante, me enseñaron que el conocimiento abre caminos y que la constancia convierte sueños en logros.

Con todo mi amor y gratitud, a mis padres, Luis Nupan y Gloria Guzmán, ya que son mi gran valioso tesoro e inspiración, gracias por brindarme apoyo incondicional y por estar siempre presentes, incluso a la distancia. Gracias por sostenerme con su amor desde el día en que salí de casa con el corazón lleno de sueños, la mochila cargada de ilusiones y los ojos humedecidos por la despedida. Gracias por estar siempre pendientes de mí, por su fe en mí y su ejemplo inquebrantable.

También dedico este logro a mis abuelos, que ahora descansan en el cielo. Aunque sus voces ya no resuenen en esta tierra, su amor y enseñanzas siguen acompañándome en cada paso, este triunfo también les pertenece, porque mucho de lo que soy nace de su legado. A mis tíos, tías, primos y primas, gracias por estar pendiente de mí, por esas preguntas llenas de cariño como: ¿Cómo estás?, ¿Cómo te va?, o ese ¡Dale, mete ñeque que ya te falta poco! que tanto me animaba. Sus palabras fueron una fuente constante de aliento y motivación, por las cuales estaré eternamente agradecido. A todos ustedes, gracias por ser parte de esta historia.

Finalmente, este logro también se lo dedico a una personita muy especial, Carlita Ojeda. Aunque no estuvo desde el primer día de este camino, llegaste en el momento justo, cuando más necesitaba una mano, un abrazo y palabras de aliento. Tu amor ha sido refugio en mis días difíciles y tú fe en mí una luz que me guió cuando quise rendimiento. Este logro no solo lleva mi esfuerzo, también lleva tu amor.

Nupan Guzmán Ederson Germán

AGRADECIMIENTO

“Encomienda tu camino al Señor, confía en Él, y Él actuará” - Salmo 37:5

En primer lugar, agradezco profundamente a Dios, por haber sido mi guía, mi fortaleza y mi refugio en cada momento de este camino. Porque incluso en días más difíciles, cuando el cansancio o la duda estaba presente, sentí su presencia brindándome paz, dirección y fuerza para seguir adelante. Quiero expresar mi más profundo agradecimiento a quienes han sido mi apoyo constante. A mis padres, por estar siempre a mi lado, brindándome confianza y amor. A mi tía Marlene, por ser una fuente constante de cariño y apoyo, también a unas personas que les considero como mis hermanos Gaby, Mayra y Cristian que han sido importantes en toda mi vida por su compañía, risas y ese afecto sincero que siempre me han brindado, como no también a Juan y Armando por su energía positiva, palabras de aliento y el cariño que me han demostrado ha sido valioso para conseguir este logro.

No puedo dejar de expresar mis agradecimientos a todos los profesores que me guiaron y me brindaron sus valiosas enseñanzas. Quiero hacer una mención especial al Msc. Mauricio Domínguez, director de mi proyecto de titulación, y al Msc. Edgar Maya, asesor del mismo. Su conocimiento, orientación y dedicación, junto a la paciencia que me han tenido en todo momento, han sido esenciales en cada etapa de este proceso, brindándome el apoyo necesario para llegar hasta aquí. También me siento agradecido con mis amigos no los puedo nombrar porque me faltara hojas para nombrarles a todos con quienes compartí no solo aulas y conocimientos, sino también risas, sufrimientos y momentos que dejan anécdotas importantes que marcaran una de las mejores etapas de esta vida, gracias por todo Pepas.

RESUMEN EJECUTIVO

El presente estudio se centra en establecer un entorno de pruebas que permita evaluar el comportamiento y la eficiencia de VXLAN (Virtual eXtensible Area Local Network) en redes de centro de datos diseñadas bajo la arquitectura Spine-Leaf. VXLAN es un protocolo de que facilita la extensión de redes de capa 2 sobre una infraestructura de capa 3 que utiliza cabeceras UDP para crear túneles entre puntos finales de red (VTEP), permitiendo crear redes virtuales extensibles sobre una infraestructura IP tradicional. La creciente demanda de escalabilidad, segmentación eficiente y optimización del tráfico de datos ha impulsado la adopción de tecnologías de superposición como VXLAN, cuyo análisis es fundamental para validar su aplicabilidad en entornos de producción reales. Para la implementación del entorno de pruebas se utiliza el simulador de redes GNS3, donde se diseña una topología Spine-Leaf utilizando dispositivos virtuales de Nvidia con sistema operativo Cumulus Linux y servidores de pruebas. Sobre esta infraestructura se desplegó VXLAN y se evaluó su comportamiento aplicando métricas basadas en la recomendación UIT-T Y.1540 que permiten valorar el desempeño real de VXLAN en una arquitectura Spine-Leaf y constituyan un aporte técnico relevante para la planificación, evaluación y toma de decisiones en la implementación de tecnologías de superposición en redes de centros de datos.

Palabras clave: VXLAN, Topología, Spine, Leaf, Rendimiento, Cumulus Linux, Testbed, VNI, BGP, Loopback, Wireshark, iPerf.

ABSTRACT

This paper focuses on establishing a testing environment that allows for the evaluation of the behavior and efficiency of VXLAN (Virtual eXtensible Area Local Network) in data center networks designed under the Spine-Leaf architecture. VXLAN is a protocol that allows for the extension of Layer 2 networks over a Layer 3 infrastructure that uses UDP headers to create tunnels between network endpoints (VTEP), enabling the creation of extensible virtual networks over a traditional IP infrastructure. The growing demand for scalability, efficient segmentation, and data traffic optimization has driven the adoption of overlay technologies such as VXLAN, whose analysis is essential to validate its applicability in real production environments. For the implementation of the test environment, the GNS3 network simulator is used, where a Spine-Leaf topology is designed using virtual devices from Nvidia with the Cumulus Linux operating system and test servers. VXLAN was deployed on this infrastructure and its performance was evaluated by applying metrics based on the ITU-T Y.1540 recommendation, which allow the actual performance of VXLAN in a Spine-Leaf architecture to be assessed and constitute a relevant technical contribution to the design, analysis, and decision-making in the implementation of overlay technologies in data center networks.

Keywords: VXLAN, Topology, Spine, Leaf, Performance, Cumulus Linux, Testbed, VNI, BGP, Loopback, Wireshark.

LISTA DE SIGLAS

VXLAN. Virtual eXtensible Local Area Network

VNI Virtual Network Identifier

VTEP. VXLAN Tunnel Endpoint

VLAN. Virtual Local Area Network

TCP. Transmission Control Protocol

UDP. User Datagram Protocol

FTP. File Transfer Protocol

GNS3. Graphical Network Simulator-3

IP. Internet Protocol

MAC. Media Access Control

STP. Spanning Tree Protocol

ZTP. Zero Touch Provisioning

IANA. Internet Assigned Numbers Authority

IEEE. Institute of Electrical and Electronics Engineers

IOS. Internetwork Operating System

BGP. Border Gateway Protocol

GRE. Generic Routing Encapsulation

RAM. Random Access Memory

CLI. Command Line Interface

IPLR: Tasa de pérdida de paquetes

IPDR: Tasa de duplicación de paquetes

IPDV: Variación del retardo de paquetes

ÍNDICE DE CONTENIDOS

1. CAPÍTULO I – ANTECEDENTES.....	1
1.1 PROBLEMA DE INVESTIGACIÓN	1
1.2 JUSTIFICACIÓN	2
1.3 OBJETIVOS	5
1.3.1 <i>Objetivo General</i>	5
1.3.2 <i>Objetivos Específicos</i>	5
2 CAPÍTULO II – MARCO TEÓRICO	6
2.1 VIRTUAL AREA NETWORK (VLAN)	6
2.2 VIRTUAL EXTENDIBLE LOCAL AREA NETWORK.....	7
2.2.1 <i>Componentes De VXLAN</i>	8
2.2.2 <i>Estructura De La Trama VXLAN</i>	9
2.2.3 <i>Mecanismos De Difusión VXLAN</i>	14
2.3 CENTRO DE DATOS	15
2.3.1 <i>Diseños Estratégicos Para Infraestructuras En Centro De Datos</i>	16
2.3.2 <i>Arquitectura De Red Spine-Leaf</i>	18
2.3.3 <i>Redes De Nueva Generación</i>	20
2.3.4 <i>Análisis Comparativo Entre Redes Overlay y Underlay</i>	22
3 CAPÍTULO III – DESPLIEGUE Y FUNCIONAMIENTO	24
3.1 EVALUACIÓN DE REQUISITOS DE EQUIPOS DE RED COMPATIBLES CON VXLAN	24
3.2 PLATAFORMA DE VIRTUALIZACIÓN: PROXMOX VE	27
3.3 HERRAMIENTAS ENFOCADAS EN EL ANÁLISIS DE PROTOCOLOS DE RED	28
3.3.1 <i>Wireshark</i>	29
3.3.2 <i>iPerf</i>	29
3.4 IMPLEMENTACIÓN DE TOPOLOGÍA PARA EL ENTORNO DE PRUEBAS	30
3.5 DETERMINACIÓN DEL SOFTWARE DE SIMULACIÓN	32

3.6 DETERMINACIÓN DEL SISTEMA OPERATIVO PARA EL USO DE GNS3	36
3.7 DESPLIEGUE DE DISPOSITIVOS Y RED EN GNS3.....	37
3.8 CONFIGURACIONES INICIALES DE DISPOSITIVOS	46
3.8.1 Asignación De Dirección Loopback En Equipos Cumulus VX.....	50
3.8.2 Configuración De Interfaces De Acceso o Puente.....	51
3.9 SELECCIÓN DE PROTOCOLO DE ENRUTAMIENTO	53
3.9.1 Border Gateway Protocol.....	55
3.10 CONFIGURACIÓN DE BGP EN CUMULUS VX	56
3.11 CONFIGURACIÓN DE VIRTUAL EXTENDIBLE LOCAL AREA NETWORK	59
3.12 CONFIGURACIÓN DE SERVICIOS PARA PRUEBAS DE OPERACIÓN DE RED	60
3.12.1 Servidor Web.....	61
3.12.2 Servidor FTP.....	63
3.12.3 Servidor De Base De Datos	65
4 CAPÍTULO IV – EVALUACIÓN DE RESULTADOS	68
4.1 VALIDACIÓN DE FUNCIONAMIENTO DEL PROTOCOLO VXLAN EN EL ENTORNO DE PRUEBAS	68
4.2 CAPTURA Y ANÁLISIS DE TRÁFICO VXLAN USANDO WIRESHARK.....	77
4.3 PRUEBAS DE RENDIMIENTO DEL ENTORNO DE PRUEBAS BASADAS EN LA UIT-T Y.1540	81
4.3.1 Planteamiento De Pruebas.....	83
4.4 PRUEBAS DE RENDIMIENTO EN EL ENTORNO VXLAN.....	88
4.4.1 Tasa De Pérdida De Paquetes (IPLR)	89
4.4.2 Disponibilidad De La Red	100
4.4.3 Tasa De Duplicación De Paquetes (IPDR)	101
4.4.4 Variación Del Retardo De Paquetes (IPDV)	111
4.4.5 Uso De Recursos.....	115
4.4.6 Medición Del Ancho De Banda	119
5 CAPITULO V – CONCLUSIONES Y RECOMENDACIONES	123
5.1 CONCLUSIONES.....	123
5.2 RECOMENDACIONES	124

6 BIBLIOGRAFÍA	125
ANEXOS.....	130
ANEXO 1: INSTALACIÓN DE DISCO SSD EN SERVIDOR DELL POWEREDGE R750XS	130
ANEXO 2: IMPORTACIÓN Y CONFIGURACIÓN DE VMs EN PROXMOX VE	135
ANEXO 3: CONFIGURACIONES DEL ENTORNO DE PRUEBAS.....	141
ANEXO 4: CONFIGURACIONES DE ACCESO Y CONEXIÓN DE EQUIPOS FÍSICOS A LA RED	147

ÍNDICE DE TABLAS

<i>Tabla 1 Comparación entre redes Overlay y Underlay.</i>	<i>22</i>
<i>Tabla 2 Equipos de red de varios fabricantes compatibles con VXLAN.</i>	<i>26</i>
<i>Tabla 3 Características técnicas del servidor DELL EMC R750xs.</i>	<i>28</i>
<i>Tabla 4 Comparación de software de simulación de redes.</i>	<i>33</i>
<i>Tabla 5 Requisitos mínimos y recomendados para la instalación de GNS3.</i>	<i>35</i>
<i>Tabla 6 Requisitos para instalación de Debian 12.</i>	<i>37</i>
<i>Tabla 7 Direccionamiento IP para interfaces y equipos del entorno de pruebas.</i>	<i>43</i>
<i>Tabla 8 Comparación de protocolos de enrutamiento.</i>	<i>54</i>
<i>Tabla 9 Asociación de VLAN con VNI.</i>	<i>73</i>
<i>Tabla 10 Terminología para establecer en el diagrama del entorno de pruebas.</i>	<i>84</i>
<i>Tabla 11 Herramientas de medición.</i>	<i>87</i>
<i>Tabla 12 Parámetros de evaluación.</i>	<i>88</i>
<i>Tabla 13 Dirección IP y puerto de servicios utilizados para las pruebas.</i>	<i>92</i>
<i>Tabla 14 Resultados de IPRL del servicio WEB.</i>	<i>95</i>
<i>Tabla 15 Resultados de IPRL del servicio FTP.</i>	<i>97</i>
<i>Tabla 16 Resultados de IPRL del servicio de Base de datos.</i>	<i>99</i>
<i>Tabla 17 Disponibilidad de la red por servicio en VXLAN y GRE.</i>	<i>101</i>
<i>Tabla 18 Resultados del servicio WEB para IPDR.</i>	<i>106</i>
<i>Tabla 19 Resultados del servicio FTP para IPDR.</i>	<i>108</i>
<i>Tabla 20 Resultados del servicio de Base de Datos para IPDR.</i>	<i>110</i>
<i>Tabla 21 Resultados obtenidos de Jitter de VXLAN y GRE.</i>	<i>113</i>
<i>Tabla 22 Base de datos de resultados de pruebas de VXLAN.</i>	<i>118</i>
<i>Tabla 23 Base de datos de resultados de pruebas de GRE.</i>	<i>118</i>

ÍNDICE DE FIGURAS

<i>Figura 1 Formato de trama Ethernet con etiquetado IEEE 802.1Q.</i>	6
<i>Figura 2 Estructura de la trama VXLAN.</i>	9
<i>Figura 3 Taxonomía de topologías para centro de datos.</i>	16
<i>Figura 4 Arquitectura Spine-Leaf.</i>	19
<i>Figura 5 Funcionamiento de las redes de nueva generación.</i>	21
<i>Figura 6 Topología base para la implementación del entorno de pruebas.</i>	32
<i>Figura 7 Proceso de instalación y configuración de equipos de red en GNS3.</i>	37
<i>Figura 8 Recursos por defecto de Cumulus VX.</i>	38
<i>Figura 9 Entorno de pruebas implementado en GNS3.</i>	40
<i>Figura 10 Configuración de nuevas credenciales de inicio de sesión.</i>	46
<i>Figura 11 Configuración de interfaz de red en Network Automation.</i>	47
<i>Figura 12 Instalación del complemento dnsmasq dentro de Network Automation.</i>	48
<i>Figura 13 Configuración de interfaces de gestión.</i>	48
<i>Figura 14 Asignación de nombres a equipos Cumulus VX.</i>	49
<i>Figura 15 Supresión del modo ZTP.</i>	50
<i>Figura 16 Establecimiento de direcciones de Loopback en Cumulus Vx.</i>	51
<i>Figura 17 Configuración de puentes en equipos LEAF.</i>	53
<i>Figura 18 Habilitación de protocolos de enrutamiento en Cumulus Vx.</i>	57
<i>Figura 19 Configuración de BGP sin enumerar en equipos Cumulus Vx.</i>	58
<i>Figura 20 Verificación de rutas BGP.</i>	58
<i>Figura 21 Configuración de VXLAN estático en switches LEAF.</i>	60
<i>Figura 22 Instalación de Apache2.</i>	61
<i>Figura 23 Reglas de firewall que permiten el tráfico HTTP.</i>	62

<i>Figura 24 Determinación del nombre del host para el servicio web.</i>	62
<i>Figura 25 Comando usado para la instalación del servicio FTP.</i>	63
<i>Figura 26 Configuración para desactivar el acceso anónimo al servidor FTP.</i>	64
<i>Figura 27 Configuración de usuarios para acceso al servicio FTP.</i>	64
<i>Figura 28 Servicio de MariaDB funcionando.</i>	65
<i>Figura 29 Configuración inicial de MariaDB.</i>	66
<i>Figura 30 Creación de base de datos.</i>	66
<i>Figura 31 Agregación de información en la Tabla de la Base de datos.</i>	67
<i>Figura 32 Creación de usuarios para ingreso a la base de datos.</i>	67
<i>Figura 33 Comprobación de vecinos BGP en el entorno de pruebas.</i>	70
<i>Figura 34 Información de túnel VXLAN en LEAF 4.</i>	70
<i>Figura 35 Información detallada de puente MAC en LEAF 1.</i>	71
<i>Figura 36 Captura de paquete VXLAN con tcpdump.</i>	72
<i>Figura 37 Trama Ethernet originada por el servidor de base de datos.</i>	73
<i>Figura 38 Encapsulación VXLAN en la comunicación de servidor WEB hasta el Cliente 1.</i>	75
<i>Figura 39 Proceso de comunicación de VXLAN.</i>	76
<i>Figura 40 Paquete Capturado en Wireshark.</i>	77
<i>Figura 41 Identificación de encapsulación de información.</i>	78
<i>Figura 42 Capa de enlace de datos externa.</i>	78
<i>Figura 43 Información de la capa de red externa.</i>	79
<i>Figura 44 Capa transporte.</i>	79
<i>Figura 45 Información de encapsulamiento de VXLAN.</i>	80
<i>Figura 46 Trama Ethernet Original.</i>	81
<i>Figura 47 Esquema de pruebas con nomenclatura según la UIT-T Y.1540.</i>	86
<i>Figura 48 Proceso para la obtención del parámetro IPRL.</i>	91

<i>Figura 49 Paquetes Totales enviados.</i>	94
<i>Figura 50 Paquetes perdidos.</i>	94
<i>Figura 51 Paquetes perdidos haciendo uso del servicio WEB.</i>	96
<i>Figura 52 Paquetes perdidos haciendo uso del servicio FTP.</i>	98
<i>Figura 53 Paquetes perdidos haciendo uso del servicio de Base de datos.</i>	100
<i>Figura 54 Proceso para la medición de IPDR en tráfico TCP.</i>	103
<i>Figura 55 Comando de iPerf3 para generar tráfico.</i>	104
<i>Figura 56 Cantidad de paquetes retransmitidos o duplicados.</i>	104
<i>Figura 57 Total de paquetes TCP transmitidos.</i>	105
<i>Figura 58 Paquetes duplicados entre VXLAN y GRE al usar el servicio WEB.</i>	107
<i>Figura 59 Paquetes duplicados en tráfico FTP.</i>	109
<i>Figura 60 Paquetes duplicados en el servicio de Base de Datos.</i>	111
<i>Figura 61 Proceso para obtener el parámetro de IPDV.</i>	112
<i>Figura 62 Obtención del valor del Jitter con iPerf3.</i>	113
<i>Figura 63 Comparación del Jitter entre VXLAN y GRE.</i>	115
<i>Figura 64 Consumo de vCPU en VXLAN.</i>	116
<i>Figura 65 Consumo de vCPU en GRE.</i>	116
<i>Figura 66 Prueba de ancho de banda con iPerf.</i>	119
<i>Figura 67 Ancho de banda servicio WEB.</i>	120
<i>Figura 68 Ancho de banda Servicio FTP.</i>	121
<i>Figura 69 Ancho de banda Servicio de Base de Datos.</i>	122
<i>Figura 70 Disco SSD listo para acoplar al servidor Dell PowerEdge R750xs.</i>	131
<i>Figura 71 Cambio de formato de disco de MBR a GPT.</i>	132
<i>Figura 72 Montaje de disco SSD en servidor DELL PowerEdge R750xs.</i>	132
<i>Figura 73 Creación del nuevo disco virtual.</i>	133

<i>Figura 74 Creación de una nueva partición en Proxmox VE.....</i>	<i>134</i>
<i>Figura 75 Verificación de virtualización aninada de Proxmox.....</i>	<i>135</i>
<i>Figura 76 Conexión de FTP entre cliente y servidor Proxmox.</i>	<i>136</i>
<i>Figura 77 Transferencia de la máquina virtual a Proxmox VE.....</i>	<i>136</i>
<i>Figura 78 Descomprimir archivo en Proxmox VE.</i>	<i>137</i>
<i>Figura 79 Importación de máquina virtual en Proxmox VE.</i>	<i>137</i>
<i>Figura 80 Asignación de nombre a la máquina virtual en Proxmox VE.....</i>	<i>138</i>
<i>Figura 81 Asignación de memoria RAM.</i>	<i>138</i>
<i>Figura 82 Asignación de números de procesadores.....</i>	<i>139</i>
<i>Figura 83 Configuración de interfaz de red de la máquina virtual.....</i>	<i>139</i>
<i>Figura 84 Configuración de dirección IP de GNS3 VM.....</i>	<i>140</i>
<i>Figura 85 Configuración inicial de GNS3.....</i>	<i>140</i>
<i>Figura 86 Conexión con el servidor remoto de GNS3 VM.....</i>	<i>141</i>
<i>Figura 87 Topología de red modificada para el acceso desde la red física.....</i>	<i>148</i>
<i>Figura 88 Configuración de direcciones IP en interfaces de Router.....</i>	<i>148</i>
<i>Figura 89 Configuración de la interfaz WAN.....</i>	<i>149</i>
<i>Figura 90 Configuración de ACL.....</i>	<i>149</i>
<i>Figura 91 Configuración de traducción de interfaces de redes internas.....</i>	<i>150</i>
<i>Figura 92 Ingreso al servidor WEB desde un terminal móvil.....</i>	<i>150</i>
<i>Figura 93 Captura del tráfico generado desde el terminal móvil al servidor Web.....</i>	<i>151</i>
<i>Figura 94 Dirección IP de laptop asignada por DHCP de Eduroam.....</i>	<i>152</i>
<i>Figura 95 Captura de Wireshark desde computadora física hacia simulación por Wi-Fi....</i>	<i>152</i>
<i>Figura 96 Prueba de ingreso a servidor FTP mediante conexión alámbrica de red física. .</i>	<i>153</i>

1. CAPÍTULO I – ANTECEDENTES

1.1 Problema De Investigación

Los requerimientos de las aplicaciones han sido el motor principal que ha impulsado la evolución de las tecnologías en los centros de datos. De acuerdo con (Davoli & Goldweber, 2015) se ha señalado que los diseños de redes tradicionales de capa 2, ya no poseen la capacidad necesaria para atender todas las demandas de los usuarios en la actualidad.

En los centros de datos que albergan información de varios usuarios, uno de los desafíos radica en que los administradores de red buscan de mantener la mayor segregación posible, gestionando el tráfico de manera individual. Actualmente, esto se logra a través del uso del estándar IEEE 802.1Q. En entornos con equipos virtualizados y clientes que demandan un extenso uso de VLAN, surgen limitaciones significativas. Esto se debe a que el número de VLAN que pueden crearse es restringido, con un límite de 4094. Este límite puede resultar insuficiente para las necesidades crecientes de estos entornos (Razo Achig, 2022).

Las infraestructuras de red frecuentemente encaran el desafío de habilitar la interconexión entre dos o más ubicaciones geográficas mediante el empleo de una VLAN común con el objetivo de mejorar el rendimiento de las aplicaciones de tráfico y segmentación de la red. Según (Lagla Gallardo, 2023) VXLAN sobresale como uno de los protocolos más avanzados que ayudan a aumentar la eficiencia de las aplicaciones de tráfico y a organizar mejor la red en los centros de datos.

Desde la perspectiva de (Salazar-Chacón & Marrone, 2022) mencionan que las VLAN presentan limitaciones en términos de alcance y despliegue. Como los routers no pueden difundir información de cierto tipo, generalmente se conectan las VLAN que se necesitan a través de varios

switches. Esta situación puede generar complicaciones al expandir considerablemente la red, dificultando su gestión. Además, este escenario implica altos costos para los centros de datos al adquirir nuevos equipos para resolver esta problemática. La proliferación de dispositivos aumenta el riesgo de posibles fallos en la red debido a errores de configuraciones, inconsistencias o interrupciones potenciales.

En la Universidad Técnica del Norte, se enfrenta una carencia significativa en cuanto a investigación y desarrollo de tecnología VXLAN. No se dispone de ningún estudio o ambiente específico destinado para probar, evaluar o realizar experimentos relacionados con la tecnología VXLAN. Esta ausencia de un entorno de diseño dedicado limita la capacidad de la institución para explorar, comprender y aplicar eficazmente esta tecnología en el contexto de sus actividades académicas e investigativas. La falta de este entorno especializado restringe la oportunidad de realizar pruebas controladas o análisis detalladas que podrían ser beneficiosos para personas interesadas en el campo de las redes y tecnologías de la información.

1.2 Justificación

El proyecto de investigación no solo se enfoca en resolver los desafíos cruciales de escalabilidad, rendimiento y flexibilidad que enfrentan los centros de datos en el dinámico ámbito de las telecomunicaciones. También representa una valiosa oportunidad para los estudiantes de Ingeniería en Telecomunicaciones de la Universidad Técnica del Norte, ya que les proporciona la posibilidad de sumergirse en tecnologías emergentes que tienen un impacto directo en el campo de las telecomunicaciones y la ingeniería en redes. Estas tecnologías emergentes son vitales para abordar problemas actuales en el ámbito de las redes, lo que les brinda a los estudiantes una visión práctica y aplicada que complementa su formación académica y les permite desarrollar habilidades

relevantes para la resolución de desafíos tecnológicos en este campo en constante evolución. En la opinión de (Tobón, 2013) al proporcionar un entorno de aprendizaje práctico y aplicado, se fortalece la formación académica de los estudiantes al permitirles desarrollar habilidades que complementan la teoría. Esta metodología los prepara para enfrentar desafíos reales y les otorga una perspectiva más amplia y práctica de las complejidades presentes en diversos ámbitos.

La implementación del protocolo de superposición de redes VXLAN representa una solución innovadora que elimina la necesidad de inversiones significativas en equipos de capa 2, una inversión tradicional requerida para abordar estos problemas. Esta transición representa un avance significativo al reducir considerablemente la carga financiera asociada a la expansión o actualización de la infraestructura de red (Tripathi et al., 2015). En este sentido, la adopción de VXLAN no solo aborda las limitaciones actuales relacionadas con la virtualización en la infraestructura de los centros de datos, sino que también se ofrece soluciones más eficientes y escalables para satisfacer las demandas actuales de conectividad y rendimiento.

El empleo de una arquitectura convencional de centro de datos restringe las capacidades y el rendimiento, lo que conlleva dificultades al brindar servicios a los usuarios, es por esa razón que se busca mitigar los efectos de latencia y congestión de red usando tecnología VXLAN (You et al., 2020). En este contexto, se está desarrollando un entorno de pruebas diseñado para demostrar de manera concreta y precisa las ventajas que esta innovadora tecnología puede aportar a la optimización de los centros de datos, ofreciendo así soluciones más eficiente y escalables para atender las demandas actuales de conectividad y rendimiento.

VXLAN se centra en abordar las limitaciones actuales relacionadas con la virtualización de la infraestructura de los centros de datos. Su objetivo principal es superar los desafíos asociados con la asignación y segmentación de VLAN. La adopción del protocolo de superposición de redes

VXLAN, amplia exponencialmente la capacidad hasta 16 millones de VLAN, atenuando eficazmente el problema de escalabilidad en los centros de datos y perfeccionando considerablemente el direccionamiento lógico (Naranjo & Salazar Ch, 2018).

Uno de los Objetivos de Desarrollo Sostenible propuestos por las Naciones Unidas busca “Establecer infraestructuras resilientes, fomentar la industrialización y promover la innovación”. En este contexto, VXLAN emerge como una innovación tecnología destacada al abordar los desafíos presentes en las infraestructuras tradicionales. Al ofrecer una solución más adaptable y flexible para las demandas de los centros de datos, VXLAN no solo optimiza la infraestructura, sino que también fomenta la evolución y el desarrollo de redes avanzadas. (Chavarro et al., 2017)

La realización de este proyecto no solo responde a la falta identificada en investigación y desarrollo de VXLAN en la Universidad Técnica del Norte, sino que también proporcionara a los interesados en el campo de las redes y tecnologías de la información una oportunidad de adquirir conocimiento más profundo y aplicable. Esta iniciativa no solo aborda la carencia de recursos para la investigación de VXLAN, sino que también posicionara a la universidad en la vanguardia de este campo, beneficiando a estudiantes en la adquisición de habilidades prácticas y experiencias relevantes para el mercado laboral.

1.3 Objetivos

1.3.1 Objetivo General

Establecer un entorno de pruebas que permita evaluar el comportamiento y la eficiencia de la tecnología VXLAN en entornos de redes de centros de datos con arquitectura Spine-Leaf.

1.3.2 Objetivos Específicos

1. Estudiar los principios fundamentales de la tecnología VXLAN, para comprender sus conceptos esenciales, funcionalidades y estructura que permita su aplicabilidad efectiva en entornos de redes de centro de datos.
2. Implementar una simulación de una red de centro de datos basada en la arquitectura Spine-Leaf, integrando de manera efectiva la tecnología VXLAN para evaluar su desempeño y comportamiento en un entorno controlado.
3. Evaluar los beneficios derivados de la implementación de VXLAN en entornos de redes de centros de datos a través de pruebas de rendimiento.

2 CAPÍTULO II – MARCO TEÓRICO

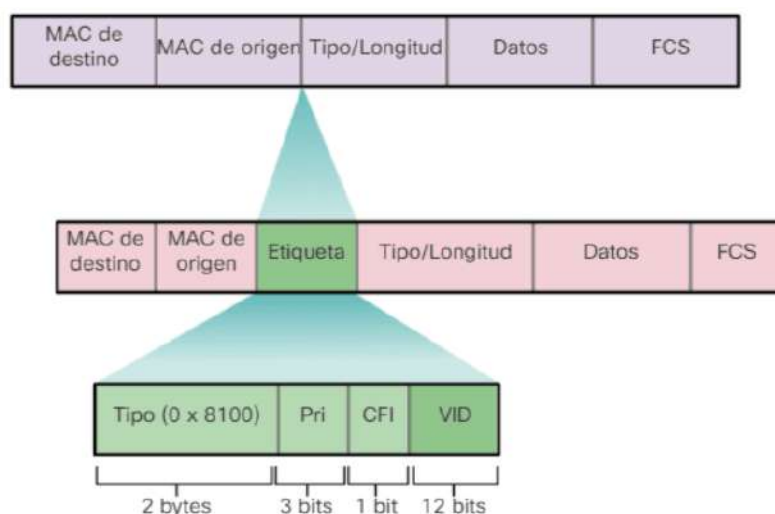
2.1 Virtual Area Network (VLAN)

El protocolo IEEE 802.1Q, llamado también como dot1Q en el área de las redes, fue creado por el grupo 802 de la IEEE. Este protocolo fue el inicio del concepto de Virtual Local Area Network (VLAN), lo que deja segmentar redes lógicas. La segmentación dada por la VLAN ayuda a gestionar el tráfico de manera eficiente, puesto que mejora el intercambio entre varios segmentos de red sin interferencias. Esto lo hace esencial en lugares donde la virtualización de redes es importante (Capella Hernández & Vicente, 2012).

De acuerdo como se muestra en la Figura 1, la trama Ethernet con la modificación del encabezado por el protocolo 802.1Q, en donde se agrega una longitud de 4 bytes. Esto crea una manera eficaz de identificar y gestionar el flujo de información en equipos de capa 2 dentro de la red.

Figura 1

Formato de trama Ethernet con etiquetado IEEE 802.1Q.



Nota. Imagen tomada de (Interpolados, 2017).

En la Figura mostrada previamente, se observa el campo de etiqueta VLAN, el cual está compuesto por cuatro elementos fundamentales. El primer elemento, se encuentra el campo “TPID (Tipo de Protocolo de Etiqueta)”, que ocupa un valor de 2 bytes designado como 0x8100 en formato hexadecimal. Cabe destacar que cada VLAN tiene un número único, con un rango de 1 a 4094, ya que el 0 (0x000) y el 4095 (0xFFFF) se encuentran reservados.

Luego, el campo es el “Pri (Prioridad de Usuario)”, el cual permite establecer prioridades en el tráfico de red durante la transmisión de la trama, ayudando a la gestión eficiente de la red. Seguidamente, se tiene el “CFI (Identificador de Formato Canónico)”, este consiste por un bit, el cual es utilizado para optimizar la transferencia de tramas Token Ring a través de enlaces ethernet.

Para finalizar, se tiene el “VID (Identificador de VLAN)”, es uno de los componentes más importantes que existe dentro del campo de etiqueta VLAN, este ocupa un valor de 12 bits. La función principal es asignar el identificador único para cada VLAN, lo cual resulta ser un papel fundamental para segmentar y aislar el tráfico en redes virtuales. Ayudando a gestionar de manera correcta cada uno de los segmentos de red y también permitiendo enrutar el tráfico entra las distintas VLAN que exista dentro de una red. (Capella Hernández & Vicente, 2012)

2.2 Virtual eXtensible Local Area Network

El RFC 7348, es el que define el protocolo VXLAN (Virtual eXtensible Local Area Network), en la actualidad es una solución de redes superpuestas que ha ganado importancia en infraestructura de redes de capa 2 y capa 3, especialmente en entornos multiusuario, como los centros de datos que manejan principalmente máquinas virtuales. VXLAN facilita la extensión de redes de capa 2 a través de una infraestructura de capa 3, lo que permite crear fácilmente segmentos VXLAN. Cada uno de los segmentos funciona como una red independiente, lo que permite tener

una comunicación segura y eficiente entre dispositivos que comparten el mismo segmento. Un elemento fundamental dentro de VXLAN es el identificador de segmento de 24 bits, conocido como "VXLAN ID" o "VNI (Virtual Network Identifier)", el cual permite reconocer de forma única hasta dentro de un mismo dominio administrativo. VXLAN permite tener 16 millones de segmentos dentro de un mismo entorno administrativo. En la actualidad, VXLAN se ha convertido en un protocolo fundamental dentro de la creación y gestión de redes superpuestas dentro de entornos modernos y dinámicos, al adaptarse fácilmente las necesidades de desarrollo en el ámbito de centro de datos (Dutt et al., 2014).

2.2.1 Componentes De VXLAN

- ✓ **VTEP (Virtual Tunnel End Point):** es un componente clave dentro de VXLAN, ya que puede implementarse como un dispositivo hardware especializado o como una entidad de software dentro de un hipervisor. Su función primaria es crear y gestionar los túneles VXLAN, los cuales son importantes para poder realizar la encapsulación como también la desencapsulación de los datos que se transporta en una red, lo que permite que los datos sean transmitidos de manera segura a través de la red (Naranjo & Salazar Ch, 2018).
- ✓ **VNI (Virtual Network Identifier):** es el identificador único para cada segmento VXLAN, la principal función es asignar una identificación independiente a cada segmento, permitiendo así su distinción dentro de la red. El VNI también ayuda a mitigar las colisiones de direcciones MAC, este proceso lo realiza mediante el encapsulamiento de las tramas Ethernet originales en un nuevo formato de trama VXLAN en donde se incorpora el VNI. Esto hace que las direcciones MAC

repetidas en diferentes segmentos pueda guardar su unicidad gracias a la combinación que existe entre el VNI y la MAC, permitiendo así tener un aislamiento efecto del flujo de datos entre distintas redes lógicas (Dutt et al., 2014).

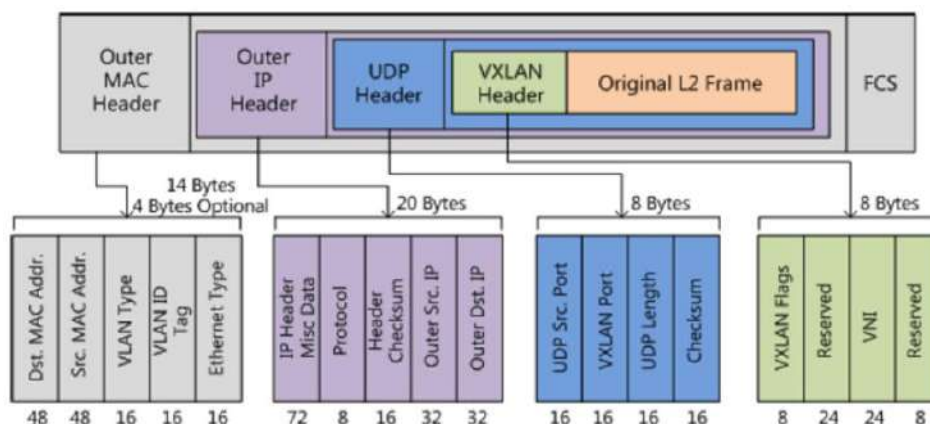
- ✓ **Túnel VXLAN:** esta es la ruta lógica que VXLAN utiliza para realizar el envío de datos, percibida dentro de la red como una conexión directa entre dos VTEP, facilitando la conectividad entre segmentos VXLAN en un entorno de red (Naranjo & Salazar Ch, 2018).

2.2.2 Estructura De La Trama VXLAN

El formato básico de la trama de VXLAN es esencial para comprender como funciona este protocolo. Cuando se utiliza un proceso de encapsulación específico, es necesario garantizar que el tráfico de la red se transmita de manera eficiente a través de una infraestructura de red. Esta estructura está compuesta de varios campos tal, como se muestra en la Figura 2, que contienen información crucial que permite la transferencia eficaz de dicha trama.

Figura 2

Estructura de la trama VXLAN.



Nota: Imagen tomada de (Huawei, 2022).

El VTEP, como punto crítico en las infraestructuras, incorpora la trama VXLAN al agregar el Virtual Network Identifier (VNI) al paquete de datos de capa 2. Posteriormente, este paquete se introduce en el túnel VXLAN, donde se lleva a cabo la encapsulación de la cabecera UDP, la asignación de una nueva dirección IP y la modificación de la cabecera MAC. Este proceso garantiza la integración eficiente del VNI y la encapsulación adecuada de la información, preparando el paquete para su transmisión segura a través de la red (Razo Achig, 2022).

2.3.2.1 Cabecera VXLAN.

La cabecera de la trama VXLAN, se compone por un tamaño de 8 Bytes, organizados en una disposición específica, como se muestra a continuación:

- ✓ Ocho bits conforman el campo “VXLAN Flags (Banderas VXLAN)”, de los cuales un bit se utiliza específicamente para la validación del VNI. Los siete bits restantes son campos reservados que deben establecerse en 0 durante la transmisión del paquete y de los cuales son ignorados en la recepción.
- ✓ El campo “VNI (Identificador de red virtual)” se caracteriza por contar con una longitud de 24 bits, utilizado para designar la red superpuesta VXLAN individual en la que se encuentran las máquinas virtuales en comunicación.
- ✓ El campo de “Reserved (Reservados)” está formado por 24/8 bits, el funcionamiento de esto es que al momento de realizar la transmisión del paquete este valor se debe de poner en cero y ser ignorado en el destino (Dutt et al., 2014).

2.3.2.1 Cabecera UDP.

La cabecera UDP, esencial en la transmisión de datos y se compone de un total de 8 Bytes con una disposición específica, como se detalla a continuación:

- ✓ Los primeros 16 bits están dedicados a identificar el “UDP Src. Port (Puerto de origen UDP)”, cuyo valor es calculado a partir de un hash que involucra a campos internos del paquete. Este proceso tiene como objetivo facilitar un equilibrio de carga del tráfico a lo largo de una red superpuesta VXLAN. En el RFC 7349, se asignan puertos dinámicos y también privados, los valores se encuentran en los rangos de 49152 y 65535.
- ✓ El “VXLAN Port (Puerto VXLAN)”, es un campo de 16 bits que permite identificar el puerto, el cual será el destino de la comunicación. El valor 4789 es un puerto UDP predeterminado para VXLAN, definido por la IANA (Internet Assigned Numbers Authority). Sin embargo, el RFC 7348 aclara que el valor del puerto puede cambiar si el administrador de una red lo requiere.
- ✓ También, existe 16 bits más que muestran el “UDP Length (Longitud del paquete)”. Este campo es primordial porque permite que el receptor pueda reconocer la cantidad de datos en bytes que debe de procesar al recibir el paquete. Esto significa que se tiene una recepción y decodificación precisa de la información enviada.
- ✓ El campo de “UDP Checksum (Suma de comprobación UDP)” es un valor de 16 bits, la función principal es para asegurarse de la integridad del paquete, si los datos se modifiquen no y habría un error en los datos. Siguiendo el RFC 7348, sugiere que la transmisión de este campo casi siempre estaba un valor de cero. Si el receptor recibe un paquete esté el campo checksum con valor “0”, el paquete debe estar aceptado para iniciar el proceso de desencapsulación. Si el checksum no tiene el valor “0”, es importante asegurar de que se haya calculado correctamente todo el paquete para así evitar fallas en la transmisión. En los casos en que se opte por

verificar este valor y resulte ser incorrecta, el paquete será descartado para así poder evitar problemas de integridad (Dutt et al., 2014).

2.3.2.1 Cabecera IP.

La cabecera IP consta de 20 bytes está directamente relacionado con la dirección IP del VTEP. Este funciona como identificador de una máquina virtual que inicia la comunicación, lo que permite transportar paquetes de datos a través de la red VXLAN. Esta dirección IP puede adoptar la forma de una dirección IP unicast o multicast. La estructura y distribución de la cabecera se muestran a continuación:

- ✓ El campo “Ip Header Misc Data (Datos varios del encabezado Ip)” está compuesto por la cantidad de 72 bits y proporciona información importante que detalla el enrutamiento y la transmisión de datos a través de la red.
- ✓ El campo “Protocol (Protocolo)” consta de 8 bits que son utilizados para el protocolo que se transporta, especifica el tipo de datos que se encuentra encapsulado, facilitando la correcta interpretación y manipulación en los puntos de origen y destino.
- ✓ Se emplean 16 bits en el campo “Header Checksum (Suma de comprobación del encabezado)”, el cual posibilita la verificación de la integridad de la cabecera IP.
- ✓ El campo “Outer Src. IP (Ip de origen externo)”, se compone por un valor de 32 bits, los cuales son utilizados para la dirección IP del VTEP origen, proporcionando la información necesaria para la identificación del punto de inicio de la transmisión.
- ✓ El campo “Outer Dst. IP (Ip de destino exterior)”, está compuesto por 32 bits para la dirección IP del VTEP destino, esta información es esencial para que el paquete enviado sea encaminado hacia el punto final previsto en la red (Razo Achig, 2022).

2.3.2.1 Cabecera MAC.

Dicha cabecera consta de 14 Bytes y se encuentran distribuidos de la siguiente manera:

- ✓ El primer campo consta de 48 bits, el cual nos indica el “Dst. MAC Addr. (Dirección MAC de destino)” que corresponde al VTEP donde reside la máquina virtual de destino.
- ✓ El campo “Src. MAC Addr. (Dirección MAC de origen)” está compuesto por 48 bits, representando la dirección MAC del VTEP donde se encuentra alojada la máquina virtual de origen.
- ✓ “VLAN Type (Tipo de VLAN)” este campo tiene 16 bits y es el que especifica el tipo de VLAN que se asocia al paquete VXLAN. Aunque este puede ser opcional, cuando se emplea, su valor es 0x8100.
- ✓ El “VLAN ID Tag (Etiqueta de identificación de VLAN)”, este también se compone por el valor de 16 bits, brinda una identificación adicional sobre la VLAN asignada a un paquete, permite tener una segmentación detallada de la red y un mapeo eficiente de datos.
- ✓ El “Ethernet Type (Tipo de Ethernet)”, es un campo conformado por la cantidad de 16 bits, el cual nos muestra el tipo de información que contiene el paquete, permitiendo tener una correcta interpretación y manejo de contenido (Razo Achig, 2022).

2.2.3 Mecanismos De Difusión VXLAN

3.3.2.1 Difusión Unicast.

La difusión o transmisión unicast en redes de comunicación se refiere a la realización de una entrega de paquetes desde un dispositivo origen hacia un dispositivo destino y para eso se hace la identificación por su dirección MAC o IP específica. Esta modalidad permite tener una comunicación directa entre dos dispositivos, evitando la distribución masiva de información los dispositivos dentro de la red, característica de difusión común.

Tomando en cuenta a (Dutt et al., 2014), las máquinas virtuales intercambian paquetes encapsulados a través de túneles VXLAN sin necesidad de conocer la red subyacente. Este mecanismo hace que la comunicación sea eficiente entre los equipos virtuales, ya que asegura la entrega precisa de la información. Además, la encapsulación en VXLAN evita que los paquetes lleguen a dispositivos que no actúan en la transmisión, lo que hace que mejore la privacidad y también optimiza la cantidad de recursos utilizados por los equipos de la red.

3.3.2.1 Difusión Multicast.

Dentro de una red VXLAN, la comunicación entre los hosts virtuales debe ser eficiente porque se debe de asegurar la correcta transmisión de los datos. Una de las técnicas que se utiliza dentro del ámbito de las redes es la difusión multicast, en donde se emplea direcciones de grupo multicast IP para realizar la transmisión de datos simultáneamente a varios destinos. Este enfoque permite optimizar el uso de ancho de banda y reduce la carga en la red, lo que garantiza tener una transmisión eficiente y de gran rendimiento.

El RFC 7348 indica el proceso de comunicación de dispositivos que se encuentran en la misma subred en una red VXLAN. Cuando una máquina virtual intenta conectarse con otra de la

misma subred, envía un mensaje en donde contiene el encabezado VXLAN, el cual agrega la información necesaria según lo descrito en la sección 2.2.2 de este documento.

Para que el tráfico multicast se dirija correctamente, es esencial asociar el VXLAN VNI con el grupo multicast IP correspondiente. Esta asociación se configura en una capa de administración y se comunica a los VXLAN Tunnel End Points (VTEP) mediante un canal de gestión. Los VTEP, basándose en este mapeo, seleccionan de manera correcta los flujos de tráfico multicast en los cuales están participando, lo que permite tener una entrega específica y eficiente de los datos entre los dispositivos de la red (Dutt et al., 2014).

2.3 Centro De Datos

Según (Arregoces & Portolani, 2004) los centros de datos son infraestructuras sofisticadas que incorporan una variedad de tecnologías en constante evolución. Su funcionamiento requiere una coordinación precisa de diversos componentes que permiten el almacenamiento, el procesamiento y la gestión de grandes cantidades de datos de forma segura y eficiente.

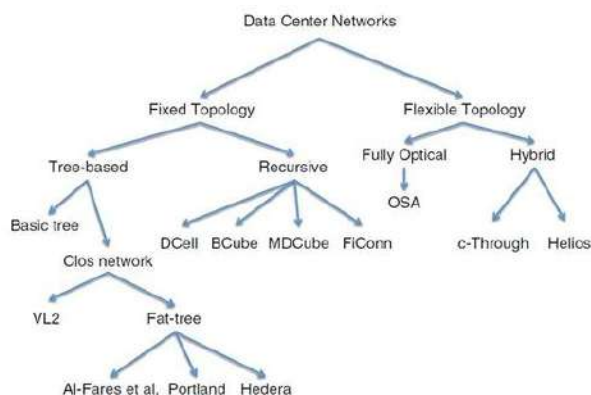
Para diseñar y operar adecuadamente una red de centro de datos, se requiere combinar conocimientos técnicos y experiencia práctica. Esto abarca desde la configuración de redes mediante protocolos de enrutamiento y conmutación, hasta la adopción de métodos de protección, balanceo de carga y administración de servidores. La comprensión integral de estos elementos es importante para garantizar el comportamiento adecuado en un entorno de TI robusto y escalable.

2.3.1 Diseños Estratégicos Para Infraestructuras En Centro De Datos

En la actualidad, existen diferentes diseños estratégicos para las redes en los centros de datos, se categorizan en varios tipos, según lo ilustrado en la Figura 3. Según (Liu et al., 2013) un diseño o topología se define como la estructura, tanto física como lógica, que determina la disposición de los dispositivos que conforman una red.

Figura 3

Taxonomía de topologías para centro de datos.



Nota: Imagen tomada de (Liu et al., 2013)

Las topologías en los centros de datos pueden fluctuar de manera significativa dependiendo de varios factores, como las necesidades particulares de una organización, la capacidad de adaptación y crecimiento, la redundancia, seguridad y eficiencia dentro de las operaciones. La selección de una topología en un centro de datos está determinada por diferentes aspectos como la capacidad de expansión, tolerancia a fallos, rendimiento, costo y la facilidad de administración. Según (Lebiednik et al., 2016) se pueden identificar diferentes topologías usuales y útiles en un centro de datos de la actualidad.

Una de las principales configuraciones que se utiliza es la fija basada en árbol, conocida como Fixed Topology Tree-based, se puede definir como una estructura jerárquica donde los

dispositivos de red se organizan en niveles parecido a las ramas de un árbol, de ahí nace su nombre, comenzando desde un nodo raíz y así se despliegan conexiones hacia nodos secundarios en niveles inferiores de nodos secundarios y también de dispositivos finales. Este modelo es muy popular dentro de las redes de centro de datos, debido a que permite gestionar una gran cantidad de volumen de tráfico y también facilita la administración de la red.

Existen diferentes variantes y dentro de esta categoría se destaca estructuras importantes como el Clos Network y el Fat-Tree. El Clos Network se destaca por ser una estructura que utiliza varias conexiones y caminos de forma paralela, los cuales hacen que los datos se transporten de forma correcta a través de diferentes rutas las cuales permiten la comunicación entre dispositivos. El Fat-Tree utiliza una estructura más densa a la comparación de Clos Network, esta jerarquía permite tener escalabilidad y redundancia dentro de la red, algunos de los aspectos claves que tiene este tipo de topología es que permite tener una gestión clara de centro de datos en donde se maneja una gran cantidad de tráfico de datos (Liu et al., 2013).

Según (Chkirbene et al., 2020) define una topología fija recursiva (Fixed Topology Recursive) es una de las arquitecturas más novedosas en la actualidad, esta se la define como una estructura la cual tiene un patrón que se repite a diferentes niveles de la red, haciendo que la red sea mucho más escalable y redundante. Su principal ventaja es que es una topología que brinda confiabilidad y facilidad de gestión, siendo así una de las estructuras más adecuadas para utilizar en un centro de datos moderno.

Como expresa (Chen et al., 2014) explica una topología flexible totalmente óptica (Flexible Topology Fully Optical), dentro de esta topología se elimina totalmente los componentes eléctricos y solo se opta por la utilización de dispositivos de transmisión óptica de datos mediante señales de luz. Este tipo de enfoque permite proporcionar velocidades de transferencia de datos muy altas y

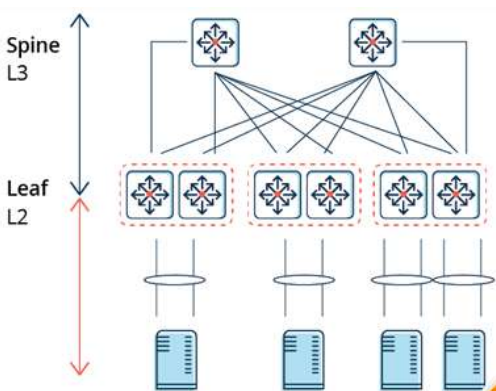
con un ancho de banda amplio, la desventaja de implementar este tipo de topología es que conlleva altos costos y tiene una complejidad de implementación.

Por último, es conveniente tener en cuenta las palabras de (Luo et al., 2015) el cuál define una topología flexible híbrida (Flexible Topology Hybrid), esta estructura jerárquica integra varios componentes eléctricos y ópticos para poder brindar una mayor flexibilidad dentro de la red para así adaptarse a las necesidades que ofrece un centro de datos. Tiene la ventaja de combinar elementos de configuración fija y flexible, permitiendo a que los entornos complejos y cambiantes de los centros de datos tengan un rendimiento optimo.

2.3.2 Arquitectura De Red Spine-Leaf

La Figura 4 indica un ejemplo referente a la arquitectura de red Spine-Leaf, este tipo de diseño es muy relevante en los centros de datos modernos. Esta arquitectura principalmente se caracteriza por tener dos niveles de conmutación: SPINE (Columna Vertebral) y LEAF (Ramificaciones), al tener los dos niveles el centro de datos se vuelve escalable y eficiente.

En la capa L2 se encuentran los equipos de cada 2 (acceso), estos son los responsables en gestionar el tráfico de los servidores y lo dirigen directamente hacia el núcleo de la red que está ubicado en la capa L3 del modelo OSI. La configuración de la arquitectura de red Spine-Leaf es una opción altamente valorada en el ámbito de centro de datos porque cada switch LEAF establece conexión con cada uno de los equipos SPINE que conforman el núcleo central. Esto hace que los cuellos de botella sean reducidos, optimizando el flujo de datos y proporcionando un alto grado de escalabilidad en la red, asegurando tener un centro de datos de alta demanda (Aruba Networking, n.d.).

Figura 4*Arquitectura Spine-Leaf.*

Nota: Imagen tomada de (Aruba Networking, n.d.).

En una arquitectura de red Spine-Leaf, tiene la posibilidad de realizar una distribución uniforme de tráfico entre los enlaces disponibles, lo que garantiza que una red permitiendo que la red se encuentre disponible en caso de existir un fallo de switches de la capa superior, también permite realizar expansión en caso de ser necesario, ya que permite agregar switches en ambos niveles de la arquitectura haciendo que incremente el ancho de banda y reduciendo así la congestión que exista dentro de los enlaces. Según (Cisco, 2023), la arquitectura Spine-Leaf es la mejor opción para un centro de datos moderno, ya que tiene la capacidad de adaptarse y expandirse, permitiendo así tener un rendimiento óptimo de la red incluso si existe un crecimiento considerable en la demanda de necesidades de los usuarios.

Al ser una arquitectura que tiene conexión entre todas las conexiones de sus dos capas permite el uso de protocolos tanto de capa 2 como de capa 3. Cuando se utilizan protocolos de capa 2, se reemplaza el Spanning Tree Protocol (STP) por opciones como Trill (Transparent Interconnection of Lots of Links) y SPB (Shortest Path Bridging), ya que ofrecen tener un enrutamiento sin bucles al aprender la ubicación de cada equipo en la red.

Por otro lado, al considerar protocolos de capa 3, es importante destacar que en esta configuración todos los enlaces están enrutados, lo que hace que esta arquitectura sea especialmente adecuada para escenarios particulares, como la superposición de redes como VXLAN. (Darquea Endara, 2017)

2.3.3 Redes De Nueva Generación

Las redes de nueva generación (NGN, por su acrónimo en inglés) representan las infraestructuras de comunicación avanzadas, las cuales aprovechan la última generación de tecnología existente en el área de las telecomunicaciones. Este tipo de redes han tenido un cambio constante, mejorando tanto en eficiencia y calidad de servicio al proporcionar mayor ancho de banda, capacidad de administrar diferentes tipos de tráfico de datos y también seguridad mejorada. Dentro de este tipo de redes, se tiene dos conceptos arquitectónicos clave que se aplica en este tipo de redes de nueva generación.

En primer lugar, se tiene el termino Underlay, de acuerdo con (Huawei, 2021), este tipo de arquitectura describe lo que es la arquitectura física de la red, la cual permite realizar operaciones dentro de las redes superpuestas Overlay. Esta infraestructura se compone de equipos especializados como son los switches, routers, firewalls y balanceadores de carga, los cuales son utilizados para transmitir información entre varias redes. La Figura 5 ilustra la estructura física básica Underlay, la conectividad entre dispositivos de red se asegura por el uso de protocolos de enrutamiento los cuales garantizan una conexión IP estable entre ellos. Cabe recalcar que la capa subyacente puede implementarse tanto en la capa 2 como en la capa 3 de la arquitectura de red.

Por otro lado, tenemos las redes Overlay o redes superpuestas establecen las conexiones lógicas entre dispositivos, lo cual es relevante en redes WAN donde los dispositivos se encuentran

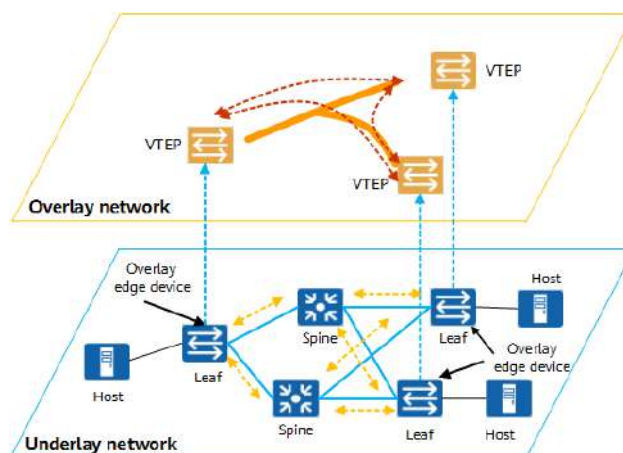
a grandes distancias. Las conexiones lógicas en redes Overlay trabajan de manera independiente de la infraestructura física que une los dos puntos, lo que significa que existe la facilidad de transportación de información a través de la conexión lógica. La Figura 5, muestra como la conexión física subyacente es respaldada por las conexiones lógicas entre los equipos de red, lo que garantiza tener una comunicación eficiente (Huawei, 2021).

Las redes Overlay componen sistemas de conexiones virtuales que se ejecutan sobre la infraestructura física mediante el uso de túneles que permiten transportar información hasta el destino final, lo cual hace que el cliente tenga el control sobre qué y de qué manera se transmiten los datos a través de la red.

En la actualidad se tiene una gran cantidad de ejemplos de túneles que se los implementa en redes Overlay, como Ipsec, GRE, y uno de los más destacados, VXLAN. Estos protocolos brindan varias formas de transmitir y encapsular información en una red, lo cual hace apropiado a las necesidades específicas de diferentes conexiones y permitiendo tener una mayor flexibilidad y gestión de la información que se envía (Franco, 2020).

Figura 5

Funcionamiento de las redes de nueva generación.



Nota: Imagen tomada de (Huawei, 2021).

2.3.4 Análisis Comparativo Entre Redes Overlay y Underlay

Es fundamental identificar las diferencias entre las redes Overlay y Underlay. En la Tabla 1, se especifican algunas características fundamentales que existen de diferente entre las redes Overlay y Underlay, hay dos enfoques distintos en el ámbito de las comunicaciones, como es la estructura lógica, accesibilidad y el control que existe dentro de la red.

Tabla 1

Comparación entre redes Overlay y Underlay.

	Underlay Network	Overlay Network
<i>Transmisión de datos</i>	Emplea routers y switches para la transmisión.	Utiliza enlaces virtuales entre nodos.
<i>Encapsulación y sobrecarga de paquetes</i>	Encapsula paquetes en Capas 2 y 3.	Requiere encapsulación específica generando sobrecarga.
<i>Control de paquetes</i>	Orientado al hardware.	Orientado al software.
<i>Tiempo de despliegue</i>	Requiere configuraciones muy extensas, lo que lo hace más lento.	Permite cambios rápidos en la estructura de la red virtual.
<i>Reenvío de rutas múltiples</i>	Aumenta sobrecarga y complejidad.	Se admite el reenvío de rutas múltiples en redes virtuales.

<i>Escalabilidad</i>	Escalabilidad limitada una vez construida.	Elevada escalabilidad y flexibilidad; VXLAN permite hasta 16 millones de redes virtuales.
<i>Protocolos</i>	Conmutación Ethernet, VLAN y protocolos de enrutamiento (OSPF, IS-IS y BGP)	VXLAN, NVGRE, SST, GRE, NVO3 y EVPN
<i>Gestión de múltiples usuarios</i>	Aislamiento complejo en redes a gran escala (NAT o VRF).	Facilita gestión de direcciones IP solapadas.

Fuente: Requerimientos tomados de (Huawei, 2021).

3 CAPÍTULO III – DESPLIEGUE Y FUNCIONAMIENTO

En este nuevo capítulo, se aborda la implementación del entorno de pruebas de la tecnología VXLAN dentro de una arquitectura Spine-Leaf. El estudio incluye una evaluación de equipos de red de diferentes fabricantes con el fin de identificar la mejor opción para realizar el diseño y montaje del entorno de pruebas. Dentro de este proceso se considera la selección del firmware adecuado y también una planificación detallada de los elementos necesarios que permitan tener el mejor rendimiento de la red. Además, se abordarán aspectos clave relacionados con la configuración y el despliegue de cada elemento de la arquitectura Spine-Leaf, garantizando una integración fluida con otros servicios de red.

3.1 Evaluación De Requisitos De Equipos De red Compatibles Con VXLAN

Antes de realizar la puesta en marcha del proceso de implementación del entorno de pruebas relacionado con VXLAN, es fundamental llevar a cabo un análisis detallado de los requisitos de equipos de red de diferentes fabricantes que sean compatibles con VXLAN. Dentro de esta etapa, es muy esencial identificar la mejor opción de equipos para poder utilizar en el entorno de pruebas, dado que los recursos disponibles en términos de CPU y memoria RAM son limitados y es por esa razón que se debe de seleccionar cuidadosamente el hardware que mejor se adapte a estas restricciones.

Teniendo en cuenta a (Gartner Peer Insights, 2024) en su último informe del cuadrante de Gartner sobre DATACENTER NETWORKING en Latinoamérica, se destacan tres fabricantes muy reconocidos en el área de las redes: CISCO, Juniper Networks y Aruba Networks. Esta información resulta ser muy importante dentro del proceso de evaluación de los equipos de red

que se pueda utilizar para el despliegue del entorno de pruebas. Estos fabricantes siempre han demostrado ofrecer un rendimiento óptimo dentro de un centro de datos, ya que están en la habilidad para satisfacer con éxito las demandas del mercado con sus innovadores productos.

Sin embargo, el mercado de Datacenter Networking no se limita solo a los líderes establecidos. También existen otros fabricantes que ofrecen alternativas innovadoras. Por ejemplo, VMware y Huawei se destacan como desafiantes (Challengers), demostrando su capacidad para competir con los líderes porque de igual forma ofrecen herramienta y soluciones innovadoras. Además, dentro del panorama de jugadores especializados (Niche Players), se encuentra Extreme Networks y Arista Networks, quienes, si bien pueden tener la misma presencia que los líderes establecidos, destacan por su enfoque especializado y por atender las necesidades específicas de forma efectiva y eficiente.

Finalmente, en la categoría de visionarios (Visionaries), se encuentra Nvidia, una empresa conocida principalmente por innovaciones en el ámbito de la computación, pero que también se encuentra incursionando en el mundo del Networking con soluciones que prometen transformar la forma en que se construyen y gestionan los centros de datos.

En la Tabla 2 se detallan los modelos seleccionados de equipos de red que son compatibles con VXLAN, además de eso se tomó en cuenta los requisitos que garanticen el mejor desempeño para la implementación del entorno de simulación del trabajo de grado. El proceso para la selección de estos modelos se ha realizado de forma minuciosa para garantizar que los equipos cumplan con todos los estándares necesarios para poder llevar a cabo la correcta implementación del entorno de simulación, dentro de estos requisitos se analizó si cada uno de estos modelos tiene firmware adecuado para poder utilizarlo dentro de un software de simulación de redes. Estos criterios son fundamentales para asegurar un rendimiento óptimo en escenarios de simulación de red.

Tabla 2

Equipos de red de varios fabricantes compatibles con VXLAN.

MODELO	DESCRIPCIÓN
CISCO NX 9318	• Memoria RAM 8 o 16 GB • CPU 4 o 6 núcleos • 1 puerto de administración • 48 puertos SFP de 1/10/25GbE • 6 puertos QSFP28 de 40/100GbE • Compatible con VXLAN • Firmware NX-OSV(Cisco, 2020)
JUNIPER EX4400	• Memoria RAM 4 o 16 GB • CPU 1 o 4 núcleos • 1 puerto de acceso de 1/10GbE • 24 puertos de 10/100/1000BASE-T • 2 puertos QSFP28 de 100GbE • Compatible con VXLAN • Firmware JunOS V21.1R1.11 (Juniper, 2024)
ARISTA 7200	• Memoria RAM 2 o 4 GB • CPU 1 o 4 núcleos • 1 puerto de administración • 24 puertos SFP+ de 10GBASE-T • 2 puertos SFP+ de 100 GbE • Compatible con VXLAN • Firmware vEOS 4.20.1F (Arista, 2024)
Nvidia Spectrum SN2100	• Memoria RAM 1 o 4 GB • CPU 2 o 6 núcleos • 64 puertos QSFP56 200GbE • Compatible con VXLAN • Firmware Cumulus Linux 4.1 o superior (Nvidia, 2024)

3.2 Plataforma De Virtualización: PROXMOX VE

La plataforma Proxmox VE, es una herramienta tecnológica reconocida por la capacidad que tiene al facilitar la creación y administración de servidores virtuales. La ventaja de que este equipo sea considerado una de las mejores plataformas de virtualización, es la capacidad de compatibilidad con múltiples sistemas operativos, lo que le permite ser la solución para una gran variedad de aplicaciones dentro del ámbito tecnológico.

Esta herramienta permite crear servidores virtuales y servidores en dos niveles de virtualización: el primer nivel utiliza principalmente recursos virtuales, mientras que el segundo nivel utiliza los recursos del servidor físico. Esta doble capacidad de ofrecer flexibilidad, ya que permite a los usuarios modificar las configuraciones según las necesidades de cada uno de estos entornos.

La plataforma de virtualización Proxmox VE, brinda una variedad de funcionalidades que mejora la relación de costo-beneficio, siendo así una opción rentable y practica para usuarios y organizaciones que necesitan una solución integral y fácil de controlar. Según (Castro, 2023) está es la solución que incluye la capacidad de gestionar almacenamiento y seguridad de manera integrada, simplificando la complejidad operativa de entornos de red.

Dentro de los laboratorios de Telecomunicaciones de la Facultad de Ingeniera en Ciencias Aplicadas de la Universidad Técnica del Norte cuentan con un servidor DELL EMC R750xs, el cual contiene un entorno de virtualización de código abierto, basado en Debian y es PROXMOX VE. Utilizado por los estudiantes para poder aprender y desarrollar proyectos de investigación. En este caso, el entorno de pruebas de VXLAN se alojará dentro de este servidor, cuyas especificaciones técnicas se encuentran en la Tabla 3. Sin embargo, al ser un recurso compartido

por todos los estudiantes, los recursos que se dispone para la realización de este trabajo de grado son limitados.

Tabla 3

Características técnicas del servidor DELL EMC R750xs.

Componente	Especificación
Versión	7.4.3
Memoria RAM	94 GB
Núcleos de CPU	48 CPU(s)
Almacenamiento	2 TB
Versatilidad	Multiplataforma
Virtualización	Niveles Dual

3.3 Herramientas Enfocadas En El Análisis De Protocolos De Red

En esta sección se presenta las herramientas importantes en el ámbito de análisis de protocolos de red, especialmente con un enfoque en VXLAN. Algunas de las más utilizadas son Wireshark e iPerf, junto a otras que también se encuentran diseñadas para poder examinar el tráfico en topologías de red. Estas herramientas permiten visualizar lo que está pasando en una red en tiempo real, permitiendo que el análisis en escenarios de pruebas como el de este trabajo sea más fácil. Gracias al uso de estas herramientas, es posible detectar fallos en la transmisión de información, problemas de rendimiento, seguridad y en otras métricas importantes que influyan dentro del funcionamiento de una red.

3.3.1 Wireshark

Con base en (Mitkov Angelov, 2021), Wireshark es una de las herramientas importantes más populares e importantes dentro del ámbito de redes. Primeramente, se destaca por ser un software que cuenta con licencia libre y su gran capacidad de analizar protocolos como es VXLAN. Wireshark captura y examina cada paquete que existe en la red con gran detalle, lo que hace que los administradores de una red.

En el caso de VXLAN, Wireshark permite interpretar la información, ya que se puede visualizar el proceso de la comunicación entre los equipos de la red. Entre estos datos importantes se incluye la estructura del paquete Ethernet en donde se debe de tener VNI, VTEP, direcciones IP y otros campos que se relacionan con VXLAN, tal y como se detalla en el punto 2.2.2 de este documento.

La interfaz gráfica de este software es amigable con el usuario y facilita observar como se encapsula y se transporta el tráfico VXLAN. Además, ofrece herramientas visuales como gráficos de flujos y estadísticas de rendimiento que permite analizar sobre como se comporta VXLAN dentro del entorno de pruebas.

3.3.2 iPerf

iPerf es una herramienta muy utilizada en el ámbito de las redes, ya que permite evaluar el rendimiento de la red. También permite medir parámetros como la capacidad de dos equipos para transferir datos, lo que proporciona información crucial sobre la estabilidad de la conexión. Esta herramienta es de código abierto, altamente compatible con una gran variedad de plataformas, incluyendo sistemas Unix, Linux, macOS y Windows, además hay que recalcar que el modo de uso de esta herramienta es por medio de línea de comandos (Arbeláez Cardona, 2013).

En caso de redes las cuales contienen VXLAN, este software resulta ser importante para poder realizar las pruebas de rendimiento correspondientes, ya que se pueden configurar dos nodos de pruebas en la red que se encuentren conectados a través de VXLAN, es posible realizar las respectivas pruebas que permitan evaluar la tasa de transferencia de datos, latencia y otros parámetros de rendimiento clave. Este análisis se lo debe de realizar repetidas veces bajo diferentes condiciones de red como también de tráfico variable que permita identificar posibles puntos de congestión y así lograr la optimización de configuraciones de red para que así el administrador de red tenga una visión detallada del comportamiento de VXLAN dentro del escenario de red y además para que se pueda tomar decisiones informadas que garanticen la utilización de VXLAN en un centro de datos (Arbeláez Cardona, 2013).

3.4 Implementación De Topología Para El Entorno De Pruebas

Para implementar la topología del entorno de pruebas, es fundamental adoptar el modelo de arquitectura Spine-Leaf, reconocido por su eficiencia y capacidad de escalabilidad. Aunque no exista una directriz específica de organizaciones internacionales al respecto, se pueden encontrar publicaciones con información general que explican como opera este tipo de topología, pero no las directrices que se debe de tomar en cuenta para crear una topología de este tipo, además en algunas comunidades de personas toman en cuenta diferentes aspectos como el tráfico de la red que entra y sale, también que la cantidad de dispositivos que se implementa en la red viene dado por la densidad de puertos de los dispositivos y en gran medida se debe de considerar las necesidades y requerimientos específicos del centro de datos.

Previo al inicio de la implementación del entorno de pruebas, se ha realizado una consulta de requisitos técnicos, principalmente en memoria RAM y cantidad de CPU en diferentes equipos

de distintos proveedores de infraestructura de redes, tal y como se indica en la sección 3.1. Esto se debe a que el entorno de pruebas se configurará en un espacio del servidor PROXMOX de la carrera de Telecomunicaciones, donde los recursos asignados son limitados. Esta evaluación ha permitido definir la cantidad óptima de switches Spine-Leaf que se pueden utilizar en el entorno de pruebas, logrando un despliegue efectivo dentro de las restricciones del entorno.

Tras evaluar la Tabla 2, se puede decir que la opción más adecuada para esta implementación es Nvidia Spectrum SN2100, el cual opera con el sistema operativo de red Cumulus Linux. Según (Salazar-Chacón & Marrone, 2022) este sistema operativo destaca por su bajo consumo de recursos de memoria RAM y CPU, es por esa razón que sobresale como una solución ideal para la implementación de VXLAN, proporcionando eficiencia y rendimiento óptimos en entornos de pruebas.

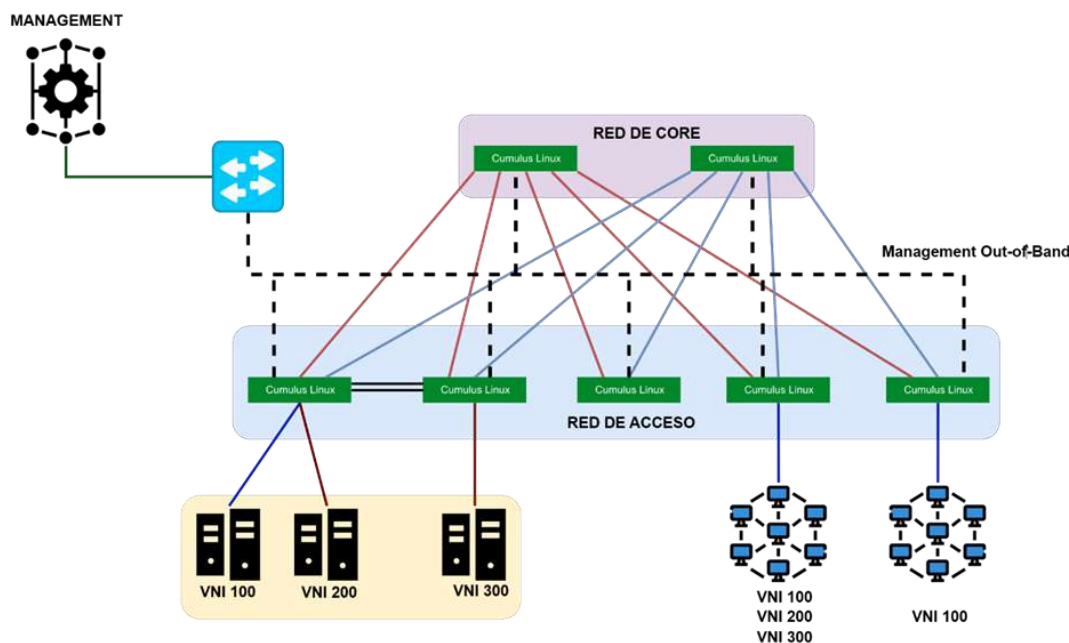
En la Figura 6 se muestra la arquitectura Spine-Leaf que será empleada para medir el rendimiento de VXLAN para este trabajo de titulación. La topología de red tiene 2 equipos SPINE y 5 equipos de acceso denominados LEAF, los cuales se conectan de manera estratégica para asegurar que la red sea confiable, distribuya el tráfico de manera equitativa y tenga alta disponibilidad al momento de realizar uso del entorno de pruebas. Con el propósito de asegurar que el servicio se encuentre siempre activo, se selecciona dos switches SPINE. Cada uno de estos equipos cuenta con conexiones de alta capacidad, lo que permite asegurar tener una gestión eficiente del tráfico generado por los dispositivos de la red de acceso.

Cabe recalcar que en la topología de red se emplea un canal de management out of band (Gestión fuera de banda), el cual se la realiza mediante un switch genérico dedicado. Este equipo se conecta a cada una de las interfaces de gestión de los equipos de red Cumulux Vx, asegurando la separación entre el tráfico de datos y el tráfico administrativo. Este enfoque evita generar un

tráfico innecesario en la red de Core, además permite que la administración de equipos sea centralizada y eficiente a través de asignación dinámica de direcciones IP a través del protocolo DHCP. Los switches LEAF forman la capa de acceso de la topología y son los responsables de conectar a los servidores y clientes finales.

Figura 6

Topología base para la implementación del entorno de pruebas.



3.5 Determinación Del Software De Simulación

La elección del software de simulación es una decisión técnica, ya que impacta de forma directa en la validez y fiabilidad del entorno de pruebas a desarrollar. Esta selección es especialmente crítica al abordar el tema de VXLAN, entonces es necesario un software que permita evaluar ciertos parámetros como rendimiento, interoperabilidad y eficacia de dicha tecnología en un entorno productivo.

Teniendo en cuenta lo anterior, existen diferentes opciones de software de simulación y emulación de redes, entre las cuales podemos destacar en este ámbito a Cisco Modeling Labs, EVE-NG y GNS3. Cada uno de estos entornos ofrecen diferentes funcionalidades y características que los hace compatibles para desarrollar las pruebas de funcionamiento. Aun así, la elección de la herramienta correcta para el entorno de pruebas no solo depende de la capacidad de soportar VXLAN, sino también de otros factores que permitan la facilidad de uso, accesibilidad de recursos y que sea rentable.

En la Tabla 4, se detalla información importante que se considera para la elección del software de simulación en donde se va a desarrollar el entorno de pruebas, esta tabla permite visualizar de manera objetiva ventajas y limitaciones de cada uno de los softwares planteados anteriormente. Los criterios que se van a evaluar incluyen compatibilidad con dispositivos que soporten VXLAN, facilidad de uso, costo, flexibilidad, escalabilidad y entre otros más.

Tabla 4

Comparación de software de simulación de redes.

Características	EVE-NG	Cisco Modeling Labs	GNS3
<i>Soporte de dispositivos</i>	Soporta diferentes dispositivos que son compatibles con VXLAN.	Es compatible solo con dispositivos del fabricante Cisco y tiene ciertas limitaciones	Es compatible con una amplia gama de dispositivos de diferentes fabricantes que soportan VXLAN.

<i>Facilidad de uso</i>	Interfaz profesional, puede ser complejo para principiantes.	Interfaz amigable, pero requiere licencias específicas.	Interfaz intuitiva, ideal para usuarios de todos los niveles.
<i>Costo</i>	Dispone de licencia de versión gratuita como también la versión profesional que es de pago.	Requiere licencias de alto costo, especialmente para nivel empresarial.	Mayormente gratuito, opciones de pago para funcionalidades avanzadas.
<i>Requisitos de sistema</i>	Requiere hardware flexible para entornos complejos.	Requiere hardware potente y licencias específicas.	Optimizado para operar en diferentes configuraciones de hardware.
<i>Integración de herramientas</i>	Buena integración de herramientas externas, pero con algunas limitaciones.	Integración limitada, solo acepta herramientas de Cisco.	Excelente integración con diversas herramientas y plataformas.
<i>Comunidad y Soporte</i>	Comunidad activa, soporte profesional disponible.	Soporte oficial de Cisco, comunidad más restringida	Amplia comunidad de usuarios, abundante documentación y recursos en línea

A través de la comparación de características en la Tabla 4, se demuestra que GNS3 se presenta como la mejor opción para desarrollar el entorno de pruebas de VXLAN para el presente trabajo de titulación, ya que principalmente es un software gratuito y que el consumo de recursos es bajo comparado a otras herramientas, además de que GNS3 es compatible con una variedad extensa de dispositivos de red y también permite agregar máquinas virtuales con diferentes sistemas operativos haciendo que las simulaciones sean lo más cercanas posible a un entorno de producción, facilitando pruebas y validaciones precisas de las configuraciones y el desempeño de un entorno de pruebas. También es beneficioso, ya que nos permite utilizar herramientas de monitoreo y análisis como es Wireshark, esta funcionalidad es esencial para evaluar el desempeño del VXLAN, ya que permite identificar y solucionar problemas de latencia, pérdida de paquetes y otros factores críticos en el desempeño de la red. Por otra parte, en la Tabla 5, se puede detallar los recursos necesarios para la instalación y utilización del software de GNS3.

Tabla 5

Requisitos mínimos y recomendados para la instalación de GNS3.

Descripción	Requerimientos	Requerimientos
	mínimos	recomendados
Procesador	2 o más núcleos lógicos	4 o más núcleos lógicos
Memoria RAM	4 GB de RAM	16 GB de RAM
Almacenamiento	1 GB	SSD > 1 GB
Virtualización	Si	Si

Fuentes: Requerimientos tomados de (GNS3, 2023).

3.6 Determinación Del Sistema Operativo Para El Uso De GNS3

En la actualidad se tiene una gran cantidad de sistemas operativos, algunos de los cuales requieren licencia para poder utilizarlos, mientras que otros son de código abierto. Cada sistema operativo utiliza recursos de forma diferente, lo que puede influir de forma indirecta en el rendimiento de las aplicaciones. Entre los sistemas operativos más populares se encuentran Windows, macOS y Linux.

En el contexto del presente trabajo de grado, el entorno de pruebas se implementará en el software de simulación de redes GNS3, seleccionado en el apartado 3.5. Teniendo en cuenta esos datos, se descarta el uso del sistema operativo de Windows debido a su alto consumo de recursos, como la CPU y RAM. Además, debido a las limitaciones de uso de recursos con que se dispone dentro del servidor PROXMOX VE, lo cual dificultará poder utilizarlo de manera eficiente dentro de la simulación del entorno de pruebas.

(GNS3, 2023) sugiere usar distribuciones de código abierto, como Ubuntu y Debian, ya que son las opciones más recomendadas y que aseguran el correcto funcionamiento eficiente del software de simulación y además garantizan un uso eficiente de recursos de hardware. En este proyecto se ha seleccionado Debian como sistema operativo debido a que tiene estabilidad y brinda un enfoque operacional en el tema de servidores, esto lo hace que sea una excelente opción para implementar el trabajo de grado. En este sentido, se empleará la versión 12 de Debian con entorno de escritorio, ya que se requiere una interfaz gráfica para interactuar con GNS3 y los demás complementos necesarios para el despliegue del entorno de pruebas. Los recursos requeridos para esta configuración se detallan en la Tabla 6.

Tabla 6

Requisitos para instalación de Debian 12.

Tipo de instalación	RAM (mínimo)	RAM (recomendado)	Disco Duro
Sin escritorio	256 Megabytes	512 Megabytes	2 Gigabytes
Con escritorio	512 Megabytes	2 Gigabytes	10 Gigabytes

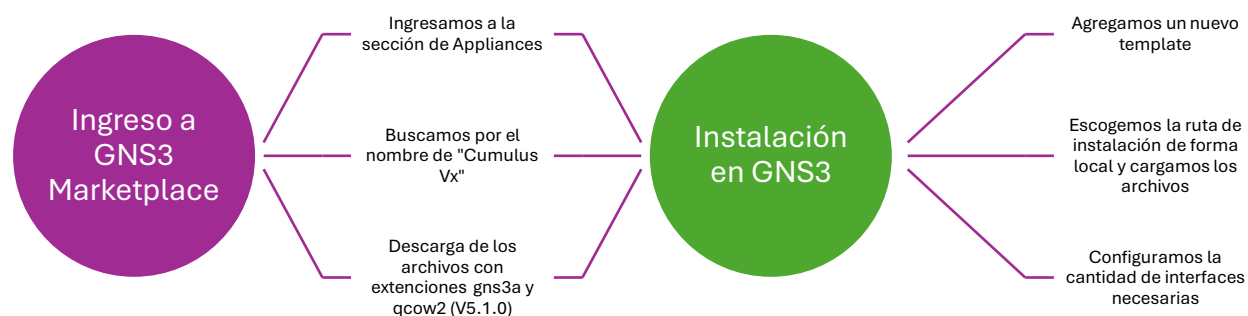
Fuente: Requerimientos tomados de (Debian, n.d.).

3.7 Despliegue De Dispositivos y Red En GNS3

Para poder desplegar la red en un entorno de GNS3, se utilizarán switches de Nvidia simulados con el sistema operativo de Cumulus Linux debido a que permite tener redes abiertas, ágiles y escalables, brindando una solución flexible para los centros de datos modernos, obteniendo un rendimiento excepcional y es por esa razón que se utilizan en arquitecturas de red como es la Spine-Leaf. En resumen, la Figura 7, presenta la fase de implementación y ajuste de este dispositivo, en donde, se evidencia como se realiza el proceso de obtención y de instalación del dispositivo en el software de GNS3, para luego utilizarlo en el despliegue del entorno de pruebas.

Figura 7

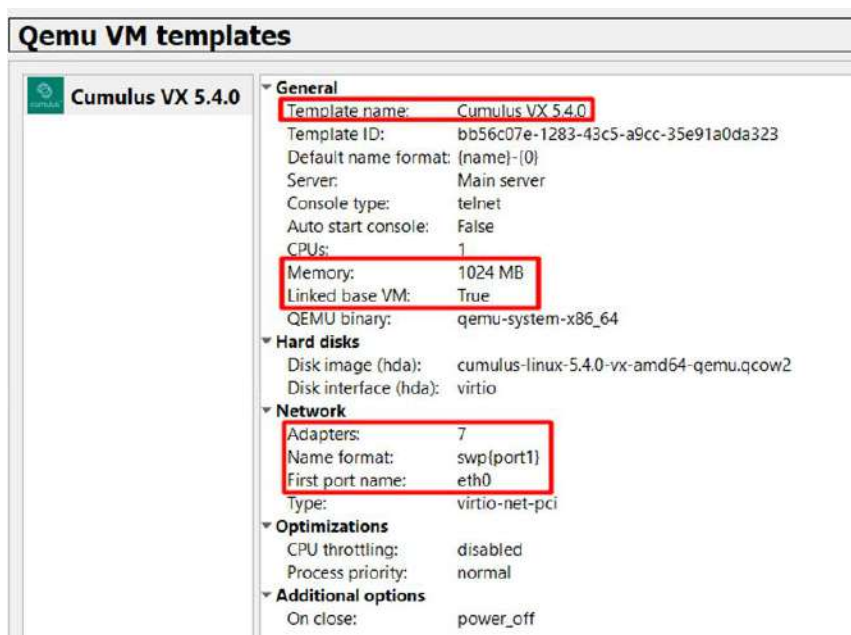
Proceso de instalación y configuración de equipos de red en GNS3.



En la Figura 8, se muestra la cantidad de recursos que se encuentra utilizando el switch virtual importado en GNS3, los valores mostrados son aquellos que ya vienen preconfigurados de manera predeterminada en donde se tiene 1024 MB de memoria RAM y 1 un procesador lógico, tal como recomienda (GNS3, 2024) para que el dispositivo pueda funcionar de una manera adecuada dentro del entorno de simulación se cambia el parámetro de memoria RAM a 1500 MB y se mantiene la cantidad de procesadores lógicos asignada por defecto. Además, se ha configurado el parámetro de red `Name format`, el cual define el formato de nombre de las interfaces de red de los switches; en este caso, las interfaces se denominan `swp{port1}` donde `port1` indica la posición de la interfaz.

Figura 8

Recursos por defecto de Cumulus VX.



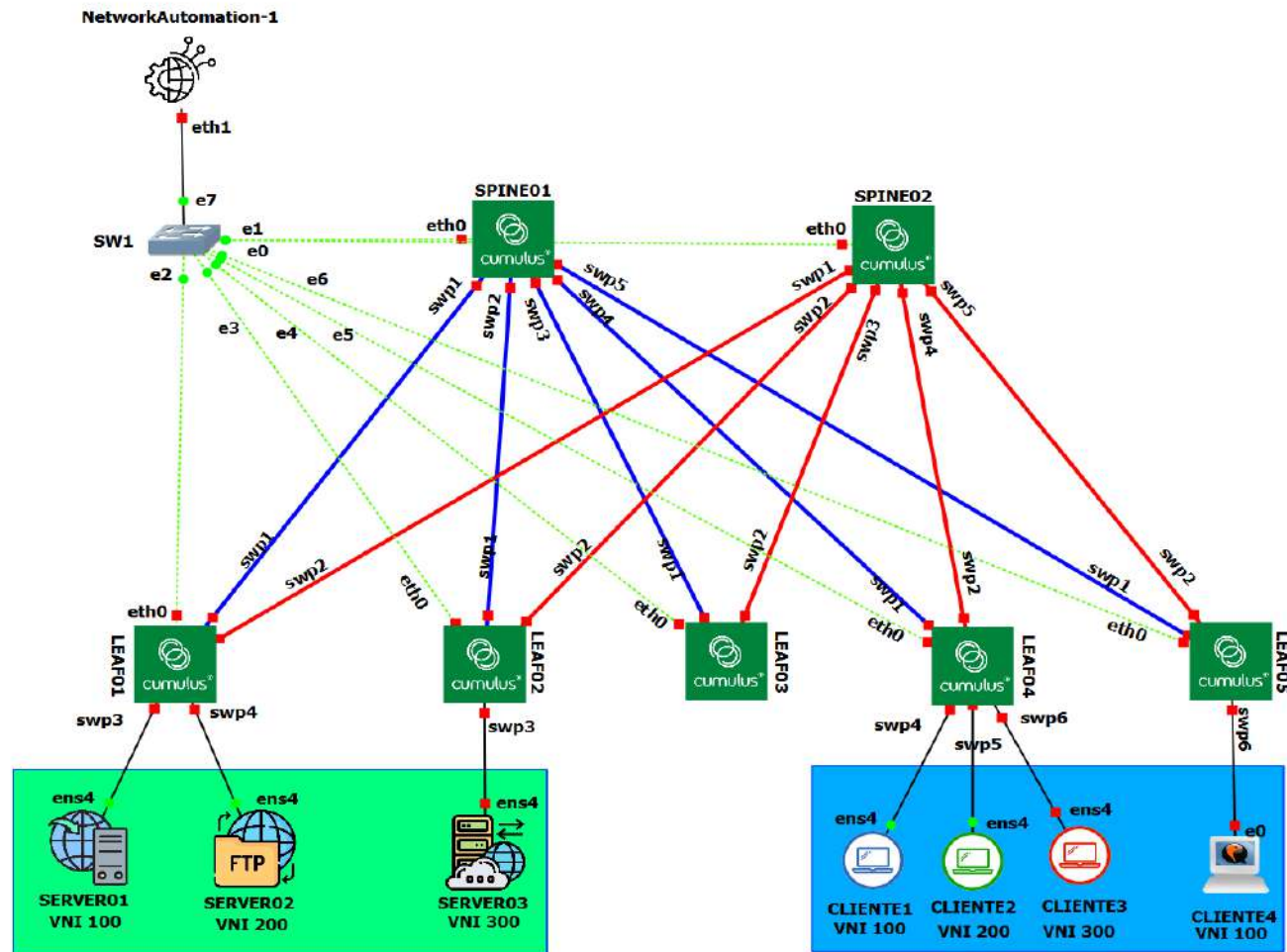
Antes de configurar una red Spine-Leaf en GNS3, es crucial planificar la topología física con detalle. Según lo planteado en la sección 3.5, se despliega dispositivos Cumulus VX, dos que

cumplirán el rol de SPINE y cinco de LEAF. También, se despliega un equipo de capa dos que se dedica a la gestión out of band con el fin de mantener el control centralizado y eficiente.

Para tener redundancia y rendimiento en la interconexión de los nodos SPINE y LEAF, cada enlace cuenta con una conexión doble, en donde se reduce los puntos de fallas por enlace. En la Figura 9, se muestra la topología implementada con los equipos de red e interfaces de red ya definidos, aunque en esta fase inicial no se ha aplicado ninguna configuración, sienta las bases para poder continuar con la implementación de VXLAN. Además, permite definir el direccionamiento IP para cada interfaz de red, asegurando tener una organización eficiente y correcta para evitar futuras equivocaciones y fallas en la red.

Figura 9

Entorno de pruebas implementado en GNS3.



Un aspecto importante dentro del despliegue de la red es la planificación del direccionamiento IP con direcciones IPv4 privadas, ya que esto permite que las conexiones entre equipos sean estable en toda la topología Spine-Leaf.

Dentro de este proceso, también se implementa el uso de Network Automation, el cual asigna direcciones IP de manera centralizada mediante el uso del protocolo DHCP, estas direcciones IP sirven para gestionar y configurar el equipo. En la Tabla 7, se incluyen las direcciones de Loopback para los dispositivos Spine-Leaf. Este tipo de direcciones se asocian de forma interna al dispositivo y no dependen de las interfaces físicas y esto hace que se tenga una mayor estabilidad dentro de la red. Las direcciones de Loopback son muy importantes dentro de VXLAN, ya que sirven como identificadores únicos que permiten la creación de los VTEP (VXLAN Tunnel End Points), garantizando puntos de origen y destino de forma confiable para dichos túneles.

Además, la Tabla 7 especifica la numeración de las VLAN que se utilizará dentro del entorno de pruebas junto con los VNI (VXLAN Network Identifiers) asignados a cada una. Esto permite visualizar de forma eficiente las VLAN locales a los VNI, facilitando el transporte de tráfico de capa 2 dentro de la red VXLAN. Además asegura una correcta segmentación del tráfico y una visualización clara entre las VLAN físicas y los identificadores lógicos que serán utilizados en el entorno de prueba.

Dentro de la tabla también se presenta una nomenclatura personalizada, la cual describe la configuración de conexión de los equipos. Por ejemplo, en el caso del LEAF 1, la columna "Interfaz" se indica el número de interfaz física que usará este dispositivo para conectar a un SPINE o a un dispositivo final. En la columna "Destino", se especifica el equipo al que está conectado el LEAF 1 y la interfaz que utiliza ese equipo para la conexión. Por ejemplo, en la primera fila se

muestra "Spine1 – Swp1", lo que significa que el LEAF 1 se conecta al equipo SPINE 1, y SPINE 1 utiliza su interfaz Swp1 para esta conexión.

Las columnas adicionales muestran las direcciones IP asignadas a cada interfaz, así como la dirección IP de la interfaz Loopback, que se utilizará en el entorno de pruebas. Esta estructura en la tabla facilita la comprensión de la topología de red, ya que brinda tener una visión clara de las conexiones y configuraciones realizadas en cada equipo.

Tabla 7

Direccionamiento IP para interfaces y equipos del entorno de pruebas.

DISPOSITIVO	INTERFAZ	IP DE LA INTERFAZ	DESTINO	LOOPBACK	VLAN	VNI
SPINE 01	Eth0	10.0.0.7/28	Switch1 – e0	Loopback 1 192.168.0.1/32	-	-
	Swp1	-	Leaf1 – Swp1		-	-
	Swp2	-	Leaf2 – Swp1		-	-
	Swp3	-	Leaf3 – Swp1		-	-
	Swp4	-	Leaf4 – Swp1		-	-
	Swp5	-	Leaf5 – Swp1		-	-
SPINE 02	Eth0	10.0.0.8/28	Switch1 – e1	Loopback 1 192.168.0.2/32	-	-
	Swp1	-	Leaf1 – Swp2		-	-
	Swp2	-	Leaf2 – Swp2		-	-
	Swp3	-	Leaf3 – Swp2		-	-
	Swp4	-	Leaf4 – Swp2		-	-
	Swp5	-	Leaf5 – Swp2		-	-
	Eth0	10.0.0.2/28	Switch1 – e2		-	-
	Swp1	-	Spine1 – Swp1		-	-

LEAF 01	Swp2	-	Spine2 – Swp1	Loopback 1 192.168.0.3/32	-	-
	Swp3	-	WEB – e0		<i>VLAN 10</i>	<i>100</i>
	Swp4	-	FTP – e0		<i>VLAN 20</i>	<i>200</i>
LEAF 02	Eth0	10.0.0.3/28	Switch1 – e3	Loopback 1 192.168.0.4/32	-	-
	Swp1	-	Spine1 – Swp2		-	-
	Swp2	-	Spine2 – Swp2		-	-
	Swp3	-	BDD – e0		<i>VLAN 30</i>	300
LEAF 03	Eth0	10.0.0.4/28	Switch1 – e4	Loopback 1 192.168.0.5/32	-	-
	Swp1	-	Spine1 – Swp3		-	-
	Swp2	-	Spine2 – Swp3		-	-
LEAF 04	Eth0	10.0.0.5/28	Switch1 – e5	Loopback 1 192.168.0.6/32	-	-
	Swp1	-	Spine1 – Swp4		-	-
	Swp2	-	Spine2 – Swp4		-	-
	Swp4	-	Cliente1 – e0		<i>VLAN 10</i>	<i>100</i>
	Swp5	-	Cliente2 – e0		<i>VLAN 20</i>	200
	Swp6	-	Cliente3 – e0		<i>VLAN 30</i>	300
LEAF 05	Eth0	10.0.0.6/28	Switch1 – e6	Loopback 1	-	-
	Swp1	-	Spine1 – Swp5		-	-

	Swp2	-	Spine2 – Swp5	192.168.0.7/32	-	-
	Swp6	-	Cliente4 – e0		<i>VLAN 10</i>	<i>100</i>
SERVIDORES	WEB - e0	10.10.10.3/24	Leaf1 – Swp3	-	<i>VLAN 10</i>	<i>100</i>
	FTP – e0	10.10.20.3/24	Leaf1 – Swp4		<i>VLAN 20</i>	<i>200</i>
	BDD – e0	10.10.30.3/24	Leaf2 – Swp3		<i>VLAN 30</i>	<i>300</i>
CLIENTES	CL1 – e0	10.10.10.10/24	Leaf4 – Swp4	-	<i>VLAN 10</i>	<i>100</i>
	CL2 – e0	10.10.20.10/24	Leaf4 – Swp5		<i>VLAN 20</i>	<i>200</i>
	CL3 – e0	10.10.30.10/24	Leaf4 – Swp6		<i>VLAN 30</i>	<i>300</i>
	CL4 – e0	10.10.10.50/24	Leaf5 – Swp6		<i>VLAN 10</i>	<i>100</i>
NetAutomation	Eth1	10.0.0.1/28	Sw1 – e7	-	-	-

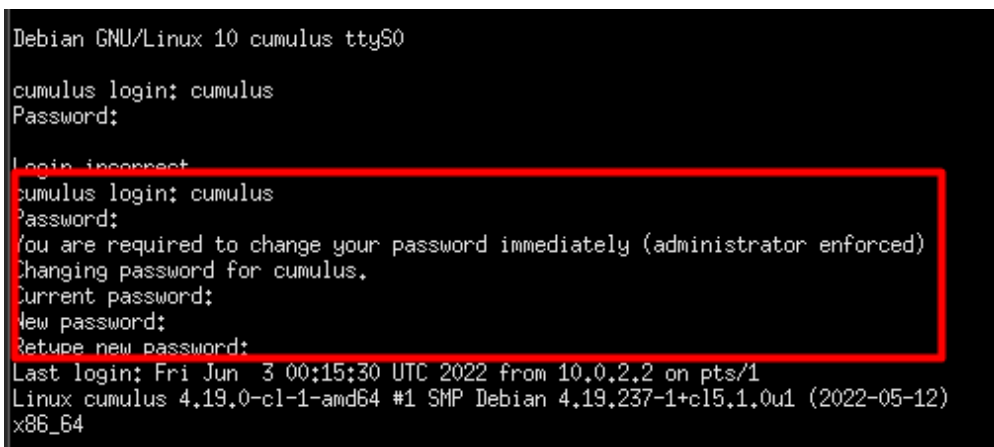
3.8 Configuraciones Iniciales De Dispositivos

Para proceder con la configuración lógica del equipo, se inicia los switches con el sistema operativo Cumulus Linux y se espera hasta que el sistema cargue de forma correcta mostrando inicialmente el ingreso de las credenciales por defecto, las cuales son “cumulus” tanto para usuario como para contraseña, caso contrario a pesar de tener corriendo el equipo y no agregar las credenciales correctas no se puede realizar la autenticación.

Una vez completado el ingreso de las credenciales, el sistema pedirá que se cambie la contraseña de administrador predeterminada por una nueva, de acuerdo con lo que se observa en la Figura 10. El procedimiento se realiza en todos los equipos de red que forman parte de la estructura SPINE y LEAF del entorno de prueba. Algo que se debe de recalcar es que las credenciales para el ingreso a los dispositivos son personalizadas dependiendo del administrador de red, ya que así se obtiene un acceso seguro y garantizado para la administración correcta de la red.

Figura 10

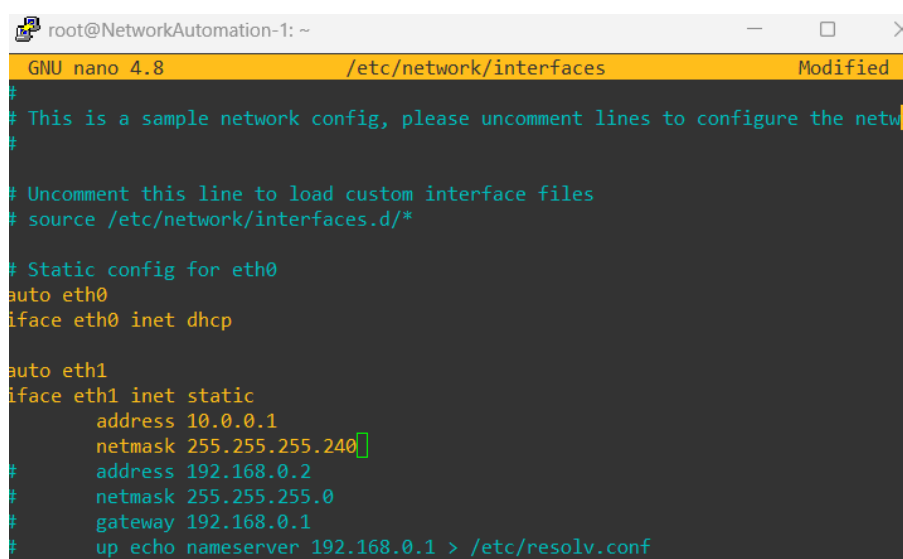
Configuración de nuevas credenciales de inicio de sesión.

A terminal window with a black background and white text. The text shows the login process for Cumulus Linux. It starts with 'Debian GNU/Linux 10 cumulus ttyS0'. Then it prompts for 'cumulus login: cumulus' and 'Password:'. After a failed login, it prompts for 'cumulus login: cumulus' and 'Password:'. This second attempt is enclosed in a red rectangular box. Below the box, the system prompts to change the password: 'You are required to change your password immediately (administrator enforced)', 'Changing password for cumulus.', 'Current password:', 'New password:', and 'Retype new password:'. The terminal then shows the last login information: 'Last login: Fri Jun 3 00:15:30 UTC 2022 from 10.0.2.2 on pts/1' and system details: 'Linux cumulus 4.19.0-cl-1-amd64 #1 SMP Debian 4.19.237-1+cl5.1.0u1 (2022-05-12) x86_64'.

El siguiente paso en el proceso es la automatización de la red (Network Automation). En la Figura 11, se detalla el proceso necesario para configurar las interfaces de red de este equipo. En este caso, se configura la interfaz `eth0`, que se encuentra conectada provisionalmente a una nube NAT. Esta conexión permite tener acceso a internet, lo cual es crucial para poder actualizar el equipo y garantizar el correcto funcionamiento. Adicionalmente, se configura la interfaz `eth1`, que se utilizará como Gateway para las interfaces de gestión de los equipos de red Cumulus Vx. Todo este proceso se lo lleva a cabo modificando el archivo de configuración ubicado dentro del directorio `/etc/network/interfaces`.

Figura 11

Configuración de interfaz de red en Network Automation.



```

root@NetworkAutomation-1: ~
GNU nano 4.8 /etc/network/interfaces Modified
#
# This is a sample network config, please uncomment lines to configure the network
#
# Uncomment this line to load custom interface files
# source /etc/network/interfaces.d/*
#
# Static config for eth0
auto eth0
iface eth0 inet dhcp
#
auto eth1
iface eth1 inet static
    address 10.0.0.1
    netmask 255.255.255.240
#
#    address 192.168.0.2
#    netmask 255.255.255.0
#    gateway 192.168.0.1
#
up echo nameserver 192.168.0.1 > /etc/resolv.conf

```

Una vez configuradas las interfaces de red, el siguiente paso es actualizar el sistema utilizando los comandos `apt-update` and `upgrade`. Tras completar la actualización, se procede a instalar `dnsmasq`, un servicio que ofrece diversas funcionalidades, destacando principalmente la gestión de DNS y DHCP. Para llevar a cabo la instalación de este complemento, se utiliza el comando `apt-get install dnsmasq`, como se indica en la Figura 12.

Figura 12

Instalación del complemento dnsmasq dentro de Network Automation.

```

root@NetworkAutomation-1:~# apt-get install dnsmasq
Reading package lists... Done
Building dependency tree
Reading state information... Done
The following additional packages will be installed:
  dns-root-data dnsmasq-base libidn11 libmn10 libnetfilter-conntrack3
  libnftnl0
Suggested packages:
  resolvconf
The following NEW packages will be installed:
  dns-root-data dnsmasq dnsmasq-base libidn11 libmn10 libnetfilter-conntrack3
  libnftnl0
0 upgraded, 7 newly installed, 0 to remove and 1 not upgraded.
Need to get 486 kB of archives.
After this operation, 1471 kB of additional disk space will be used.
Do you want to continue? [Y/n]

```

Una vez instalado el complemento, se procede a configurar las direcciones IP de gestión para cada equipo de red Cumulus VX, la cual es la interfaz `eth0`. En este caso, se utiliza el rango de red `10.0.0.0/28` para asignar las direcciones. En la Figura 13, se detalla el proceso de configuración para cada equipo. Para ello, se modifica el archivo `dnsmasq.conf`, ubicado en el directorio `/etc/`. Dentro de este archivo se configura el servicio DHCP especificando las direcciones IP correspondientes para cada host, las cuales se definen en función de la dirección MAC de cada equipo.

Figura 13

Configuración de interfaces de gestión.

```

GNU nano 4.8 /etc/dnsmasq.conf

# Include all files in a directory which end in .conf
#conf-dir=/etc/dnsmasq.d/*.conf

# If a DHCP client claims that its name is "wpad", ignore that.
# This fixes a security hole. see CERT Vulnerability VU#598349
#dhcp-name-match=set:wpad-ignore,wpad
#dhcp-ignore-names=tag:wpad-ignore

dhcp-range=10.0.0.2,10.0.0.14, 12h

dhcp-host=0c:bd:f2:11:00:00,leaf-01,10.0.0.2,12h
dhcp-host=0c:13:fb:ce:00:00,leaf-02,10.0.0.3,12h
dhcp-host=0c:e8:80:a5:00:00,leaf-03,10.0.0.4,12h
dhcp-host=0c:6f:c7:02:00:00,leaf-04,10.0.0.5,12h
dhcp-host=0c:7c:ef:92:00:00,leaf-05,10.0.0.6,12h

dhcp-host=0c:5d:46:b7:00:00,spine-01,10.0.0.7,12h
dhcp-host=0c:d8:bc:a1:00:00,spine-02,10.0.0.8,12h

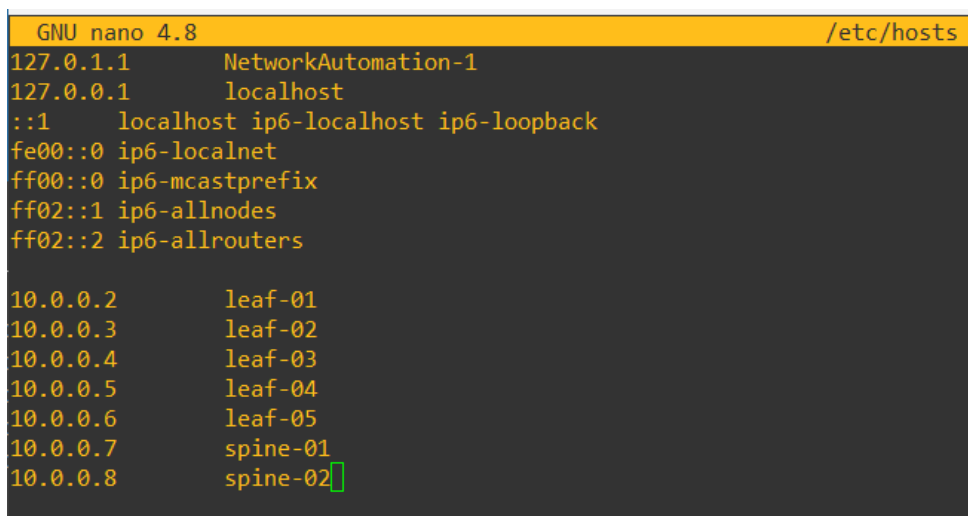
```

A continuación, se procede a configurar el hostname del equipo. Para ello, se realiza cambios en el archivo `/host`, ubicado en el directorio `/etc`. Este archivo permite asignar un nombre único a cada dispositivo, facilitando su identificación dentro de la red.

En la Figura 14, se muestra el proceso de edición del archivo, donde se asignan los hostnames a cada equipo Cumulus Vx de la red. Este paso es clave, ya que no solo organiza el sistema, sino que lo hace más comprensible y accesible para que cualquier administrador de red pueda comprender la topología.

Figura 14

Asignación de nombres a equipos Cumulus VX.



```
GNU nano 4.8 /etc/hosts
127.0.1.1    NetworkAutomation-1
127.0.0.1    localhost
::1         localhost ip6-localhost ip6-loopback
fe00::0     ip6-localnet
ff00::0     ip6-mcastprefix
ff02::1     ip6-allnodes
ff02::2     ip6-allrouters

10.0.0.2     leaf-01
10.0.0.3     leaf-02
10.0.0.4     leaf-03
10.0.0.5     leaf-04
10.0.0.6     leaf-05
10.0.0.7     spine-01
10.0.0.8     spine-02
```

Finalmente, se llega a un paso importante: iniciar los servicios de `dnsmasq`. Para ello, usamos el comando `/etc/init.d/dnsmasq start`. Esto activa todas las funcionalidades de DNS y DHCP, al hacerlo se permite que los equipos de la red se configuren según lo establecido anteriormente.

3.8.1 Asignación De Dirección Loopback En Equipos Cumulus VX

La adecuada configuración de los equipos de red garantiza tener una infraestructura eficiente y estable. Según (Nvidia, 2024), los equipos pueden configurarse utilizando comandos estándar de Linux o mediante comandos de NCLU¹, cada opción ofreciendo ventajas específicas para la administración y flexibilidad en la red.

Para empezar con la configuración de los dispositivos, lo primero que se realiza es la desactivación del modo ZTP (Zero Touch Provisioning), como se muestra en la Figura 15. Este modo ya viene configurado por defecto en el sistema operativo Cumulus Linux, este hace que un dispositivo se configure automáticamente desde el momento en que se conecta a la red. En las redes a gran escala, dicha característica resulta ser muy utilizada.

Figura 15

Supresión del modo ZTP.

```
cumulus@spine-01:mgmt:~$ sudo ztp -d  
[sudo] password for cumulus:
```

La Figura 16 muestra el proceso de asignación de direcciones IP a las interfaces Loopback en los dispositivos SPINE y LEAF. Estas direcciones son la base para la implementación de VXLAN para que funcione correctamente y también para que protocolos de enrutamiento funcionen de manera óptima dentro de la topología, ya que la Loopback es una IP fija y estable, fundamental para el establecimiento de rutas y túneles.

En el contexto de VXLAN, la interfaz Loopback juega un rol importante, actuando como el punto de terminación de los túneles en la red overlay. Gracias a esto asegura una transmisión

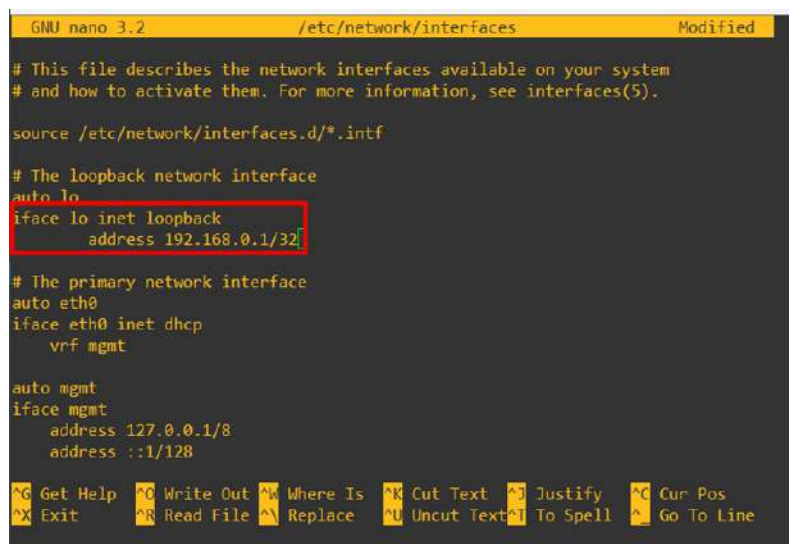
¹ Comandos NCLU: se trata de una interfaz de línea de comandos, especialmente diseñada para dispositivos de red de NVIDIA. Opera en el espacio de usuario de Linux y permite ejecutar comandos de red directamente desde bash, lo que la hace accesible para cualquier usuario que tenga conocimientos básicos en redes.

confiable del tráfico de la red virtualizada. Además, su uso aporta mayor resiliencia y escalabilidad, ya que esta interfaz no depende las físicas y siempre permanece operativa.

Por su parte, la Tabla 7 presenta cuales son las direcciones Loopback asignadas a los dispositivos de red SPINE y LEAF.

Figura 16

Establecimiento de direcciones de Loopback en Cumulus Vx.



```
GNU nano 3.2 /etc/network/interfaces Modified
# This file describes the network interfaces available on your system
# and how to activate them. For more information, see interfaces(5).

source /etc/network/interfaces.d/*.intf

# The loopback network interface
auto lo
iface lo inet loopback
       address 192.168.0.1/32

# The primary network interface
auto eth0
iface eth0 inet dhcp
       vrf mgmt

auto mgmt
iface mgmt
       address 127.0.0.1/8
       address ::1/128

^G Get Help  ^O Write Out ^W Where Is  ^K Cut Text  ^J Justify   ^C Cur Pos
^X Exit      ^R Read File ^N Replace   ^U Uncut Text ^I To Spell  ^_ Go To Line
```

3.8.2 Configuración De Interfaces De Acceso o Puente

El establecimiento de puentes (bridge) en una red es clave para segmentar y gestionar el tráfico de manera individual a través del protocolo 802.1Q o más conocido como VLAN. En este contexto, los puentes cumplen un rol de switch, pero de forma virtual permitiendo la conexión de interfaces físicas o lógicas dentro de un mismo dominio de broadcast.

En este caso, los puentes se implementan en todos los dispositivos de la red, tanto en los equipos con rol de SPINE como en los LEAF. Los bridges o puentes son especialmente relevantes en equipos de la red de acceso, ya que actúan como puntos de conexión para servidores y clientes

finales, y también en el SPINE2, que está conectado al router y facilita tener comunicación mediante enrutamiento inter-VLAN.

La Figura 17 muestra un ejemplo de configuración en el LEAF 1, donde las interfaces utilizadas como puente es la Swp3 y Swp4, cada una de estas asociadas a una VLAN en específico. Este diseño asegura que el tráfico de cada VLAN permanezca aislado, mejorando seguridad y rendimiento de la red.

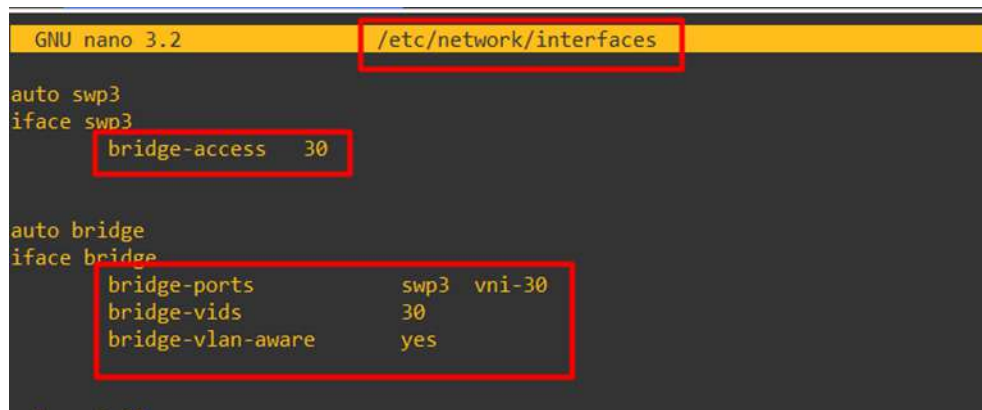
La configuración de bridges se lo realiza dentro del archivo `interfaces`, ubicado en el directorio `/etc/network/`. En este archivo se define la asociación de las interfaces físicas conectadas a servidores o clientes con los puentes, los parámetros configurados son los siguientes:

- Auto swp: especifica que la interfaz `swp3` se activara de forma automática al iniciar el equipo.
- Iface swp: define la configuración de la interfaz, en este caso se asocia a una VLAN.
- bridge-access: Configura la interfaz `swp4` como un puerto de acceso, el cual se asocia a la VLAN. Esto significa que todo tráfico que este en tránsito sobre la interfaz swp será procesado, ya que pertenece a la VLAN correspondiente.
- bridge-ports: especifica la interfaz física y lógica que se asocian a un puente.
- bridge-vids: Establece que el puente permite el tráfico etiquetado con el valor de la VLAN ID, asegurando que el puente filtre el tráfico ni etiquetado o de otras VLANs.
- Bridge-vlan-aware yes: Habilita el modo de reconocimiento de VLAN para el puente y permite conectar múltiples interfaces, separando y organizando el tráfico basado en VLANs.

Este mismo procedimiento se lo lleva a cabo de forma uniforme en los LEAF 1,2,4 y 5, ajustándose según las interfaces y las VLAN que necesite acceso, asegurando funcionalidad en toda la red.

Figura 17

Configuración de puentes en equipos LEAF.



```

GNU nano 3.2 /etc/network/interfaces

auto swp3
iface swp3
    bridge-access 30

auto bridge
iface bridge
    bridge-ports      swp3 vni-30
    bridge-vids        30
    bridge-vlan-aware  yes
  
```

3.9 Selección De Protocolo De Enrutamiento

La elección del protocolo de enrutamiento es clave dentro del diseño e implementación de redes basadas en VXLAN, sobre arquitecturas Spine-Leaf. Este protocolo no solo determina como fluye la información a través de la red, sino que también afecta a cosas como la velocidad con la que adapta a cambios y su capacidad para crecer. Como VXLAN está diseñado para facilitar la creación de redes de capa 2 sobre una infraestructura de capa 3, el protocolo de enrutamiento debe de ser capaz de manejar el tráfico encapsulado, asegurando la correcta transmisión de información entre los segmentos de red existentes.

Se considera algunos protocolos que ofrecen un balance entre sus características, ventajas y limitaciones. Dentro de las opciones se ha escogido 3 de los protocolos de enrutamiento más populares en este ámbito como es OSPF, BGP y RIP. Cada uno de estos protocolos tienen sus

propias características que los hacen más o menos adecuados dependiendo de la red, requisitos de convergencia, complejidad de configuración y capacidad de escalabilidad.

En la Tabla 8, se muestran las principales diferencias entre OSPF, BGP y RIP, comparando como se comportan en entornos VXLAN y también analizando sus ventajas y desventajas según las necesidades de la red.

Tabla 8

Comparación de protocolos de enrutamiento.

Criterio	RIP	BGP	OSPF
<i>Compatibilidad con VXLAN</i>	No es compatible con VXLAN, ineficiente para redes modernas.	Soporte nativo para VXLAN y ampliamente utilizado en redes Spine-Leaf.	Compatible con VXLAN, requiere configuraciones avanzadas y extensiones como MP-BGP.
<i>Facilidad de configuración</i>	Fácil de configurar, pero limitado a redes pequeñas.	Configuración promedio, con opciones intuitivas para definir áreas y rutas internas.	Configuraciones extensas.
<i>Consumo de recursos</i>	Bajo, pero no es eficiente en redes de centro de datos.	Optimizado para centro de datos, utiliza recursos moderados.	Tiene un alto consumo de CPU y memoria debido al manejo de grandes

			tablas de enrutamiento.
<i>Soporte Multicast</i>	Básico, pero no es compatible con VXLAN.	Soporte nativo para tráfico multicast, crucial para VXLAN.	No soporta multicast de manera nativa y requiere configuraciones adicionales para que funcione con VXLAN.

Tras analizar la información de la Tabla 8, la mejor opción para aplicar dentro del entorno de pruebas que se basa en VXLAN y con arquitectura Spine-Leaf es BGP. Este protocolo ofrece un buen equilibrio entre velocidad, rendimiento y escalabilidad, lo que permite tener una red confiable. En este tipo de topologías que son dinámicas y complejas el uso de BGP es ideal para este tipo de implementaciones.

3.9.1 Border Gateway Protocol

De acuerdo con (Juniper, 2025), el protocolo BGP (Border Gateway Protocol) es un protocolo de puerta de enlace exterior (EGP) utilizado para el intercambio información entre equipos de capa 3 de diferentes sistemas autónomos (AS). Este protocolo de enrutamiento mantiene una base de datos de información de accesibilidad de la red, además la intercambia con otros sistemas que utilizan el mismo protocolo de enrutamiento.

Una de las ventajas que BGP presenta es que elimina los bucles de los enrutadores y aplicar las políticas del sistema autónomo. BGP utiliza el protocolo TCP para el transporte de la información, utilizando el puerto 179 para poder establecer la conexión con otros equipos, es por esa razón que este protocolo es confiable y suprime la fragmentación de actualizaciones, retransmisiones, reconocimiento y secuenciación como lo hacen otros protocolos de enrutamiento.

3.10 Configuración De BGP En Cumulus VX

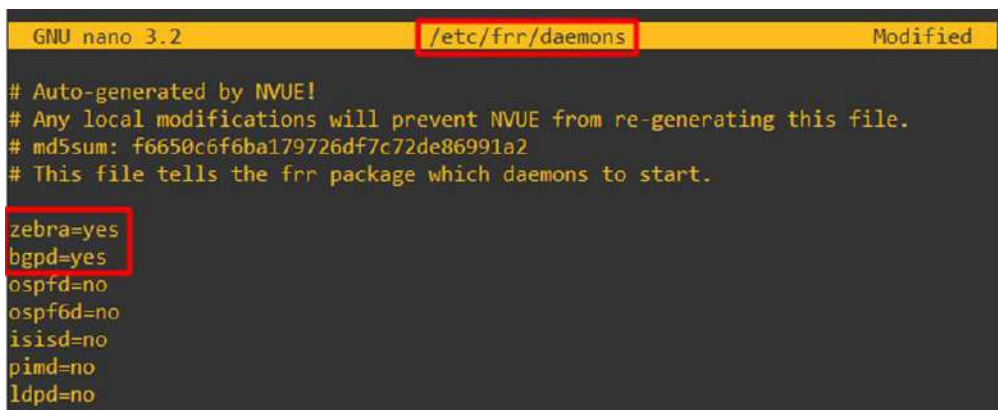
Dentro de la topología Spine-Leaf, se realizó la configuración en los equipos Cumulus VX, que comprenden los dos SPINE y los cinco LEAF. Para ello, lo primero es editar el archivo `daemons`, ubicado en la ruta `/etc/frr/`. Este archivo es importante porque permite habilitar o deshabilitar diferentes servicios, también llamados demonios, dentro del software FRRouting (FRR), usado en redes para implementar protocolos de enrutamiento dinámicos.

En la Figura 18 se muestra en cuadros rojos las configuraciones específicas hechas en este archivo. Mediante el uso de `nano`, editamos el archivo y lo primero que se realiza es agregar la línea `zebra=yes` para iniciar o habilitar el demonio Zebra, el cual se encarga de gestionar la tabla de enrutamiento de los equipos y también de coordinarse con todos los protocolos de enrutamiento dinámico, en caso de que se estará aplicando diferentes protocolos de enrutamiento.

Dentro de este mismo archivo, hay una línea de comando que permite activar BGP. Esta línea ya viene por defecto, por lo cual es necesario quitar el comentario y la línea debe mostrarse de la siguiente manera `bgpd=yes` para habilitar el protocolo de enrutamiento BGP. Con esta configuración se iniciará de forma automática cada vez que los equipos se enciendan, simplificando la operación en la red desplegada.

Figura 18

Habilitación de protocolos de enrutamiento en Cumulus Vx.



```
GNU nano 3.2 /etc/frr/daemons Modified
# Auto-generated by NVUE!
# Any local modifications will prevent NVUE from re-generating this file.
# md5sum: f6650c6f6ba179726df7c72de86991a2
# This file tells the frr package which daemons to start.
zebra=yes
bgpd=yes
ospfd=no
ospf6d=no
isisd=no
pimd=no
ldpd=no
```

En la Figura 19 se indica la configuración del protocolo de enrutamiento BGP en uno de los equipos de la topología implementada. Teniendo en cuenta a (Nvidia, 2024) existen dos formas de enrutamiento BGP en este tipo de equipos, la primera es la configuración de BGP enumerado en donde la distribución de información de enrutamiento es proporcionada por una dirección IP entre dos pares y la segunda opción es BGP sin enumerar, en donde el vecino en este caso ya no es necesario una dirección IP en la interfaz del dispositivo, sino el nombre de la interfaz a la cual está se conecta los vecinos. En este caso, se utilizó BGP sin enumerar para mayor facilidad y evitar confusiones, los comandos utilizados son los diseñados por NVIDIA para simplificar la configuración de BGP dentro de los equipos Cumulus VX.

Los comandos NCLU configuran las interfaces del switch y en este caso se asigna la dirección IP 192.168.0.1/32 para la interfaz de Loopback, la cual va a actuar como una dirección de identificación de router (Router ID) en el protocolo BGP. Además, se asigna el número de sistema autónomo 65199, el cual identifica al router dentro de la red BGP global.

Específicamente, las configuraciones incluyen la asignación de la interfaz de Loopback y las interfaces de los puertos de switch como vecinos BGP, con la designación 'remote-as external',

lo que indica que estos puertos están conectados a redes externas. Esto se complementa con la especificación de una red IPv4 unicast (192.168.0.1/32) que se anuncia a través del protocolo BGP, asegurando la propagación de esta red dentro del enrutamiento de BGP.

Este enfoque se aplica en toda la red, siguiendo una metodología estándar que mantiene la interoperabilidad entre los diferentes dispositivos de la topología y asegura que la red BGP sea confiable y eficaz.

Figura 19

Configuración de BGP sin enumerar en equipos Cumulus Vx.

```
cumulus@cumulus:mgmt:~$ nv set interface lo ip address 192.168.0.1/32
cumulus@cumulus:mgmt:~$ nv set interface swp1-5
cumulus@cumulus:mgmt:~$ nv set router bgp autonomous-system 65199
cumulus@cumulus:mgmt:~$ nv set router bgp router-id 192.168.0.1
cumulus@cumulus:mgmt:~$ nv set vrf default router bgp neighbor swp1 remote-as external
cumulus@cumulus:mgmt:~$ nv set vrf default router bgp neighbor swp2 remote-as external
cumulus@cumulus:mgmt:~$ nv set vrf default router bgp neighbor swp3 remote-as external
cumulus@cumulus:mgmt:~$ nv set vrf default router bgp neighbor swp4 remote-as external
cumulus@cumulus:mgmt:~$ nv set vrf default router bgp neighbor swp5 remote-as external
cumulus@cumulus:mgmt:~$ nv set vrf default router bgp address-family ipv4-unicast network 192.168.0.1/32
cumulus@cumulus:mgmt:~$ nv config apply
Applied from id: 21
```

Finalmente, en la Figura 20 se valida a través de comandos, que los cambios realizados están correctamente configurados. Este paso permite asegurar que la comunicación y el intercambio de rutas dentro del Sistema Autónomo se realicen sin problemas y de manera adecuada, indicando la sincronización de los equipos.

Figura 20

Verificación de rutas BGP.

```
cumulus@cumulus:mgmt:~$ net show route bgp
RIB entry for bgp
=====
Codes: K - kernel route, C - connected, S - static, R - RIP,
       O - OSPF, I - IS-IS, B - BGP, E - EIGRP, N - NHRP,
       T - Table, v - VNC, V - VNC-Direct, A - Babel, D - SHARP,
       F - PBR, f - OpenFabric, Z - FRR,
       > - selected route, * - FIB route, q - queued, r - rejected, b - backup
       t - trapped, o - offload failure
B>* 192.168.0.2/32 [20/0] via fe80::ed0:4ff:feb2:1, swp1, weight 1, 00:00:03
B>* 192.168.0.3/32 [20/0] via fe80::ed0:4ff:feb2:1, swp1, weight 1, 00:00:04
B>* 192.168.0.4/32 [20/0] via fe80::ebf:54ff:fe69:1, swp2, weight 1, 00:00:03
B>* 192.168.0.5/32 [20/0] via fe80::e44:f8ff:fea6:1, swp3, weight 1, 00:00:03
B>* 192.168.0.6/32 [20/0] via fe80::ee8:a8ff:fee9:1, swp4, weight 1, 00:00:03
B>* 192.168.0.7/32 [20/0] via fe80::ef6:b4ff:fe35:1, swp5, weight 1, 00:00:03
```

3.11 Configuración De Virtual eXtensible Local Area Network

Se procede a configurar VXLAN dentro de la red Spine-Leaf, esta es la parte principal del trabajo de grado, en este caso se utiliza túneles VXLAN estáticos, teniendo en cuenta a (Nvidia, 2024) los túneles VXLAN estáticos son una solución práctica para conectar dos VTEP en un entorno dado. Al asignar los VTEP a una VNI en particular, se reduce significativamente el esfuerzo necesario para configurar conexiones individuales.

La configuración mostrada en la Figura 21 utiliza comandos ejecutados en Linux, dentro del archivo `interfaces` ubicado en el directorio `/etc/network/`.

Este ejemplo muestra como se configura el LEAF 1, donde se crea puentes y se definen las interfaces físicas y virtuales (`vni-10` y `vni-20`). En este caso, el `vni-10` se asigna a la VLAN 10 haciendo uso del comando `bridge-access`, lo que asegura que este asociado correctamente con el segmento de red.

Para proteger la infraestructura contra bucles y posibles ataques relacionados con BPDU, se habilitan dos configuraciones clave de seguridad: `mstpctl-bpduguard`, que bloquea puertos en caso de recibir tramas BPDU, y `mstpctl-portbpdufilter`, que filtra dichas tramas por puertos específicos.

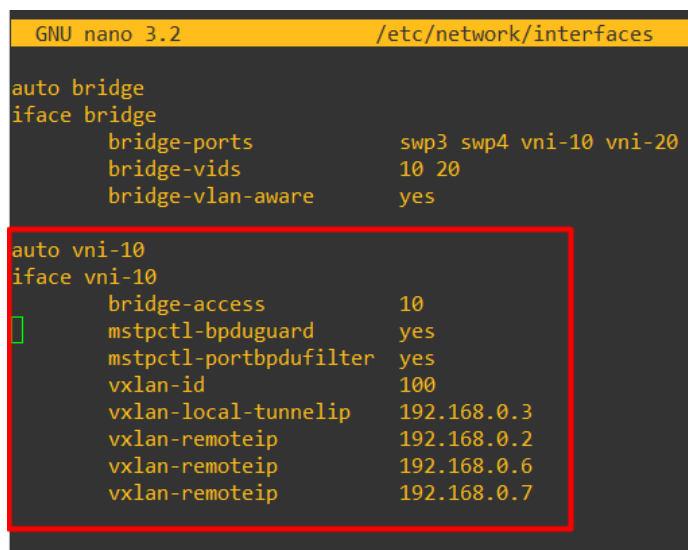
El `vxlan-id` sirve como el identificador para el dominio de encapsulación de tráfico de VXLAN, garantiza que el tráfico se encamine con precisión dentro de la red virtual. Por otro lado, el parámetro `vxlan-local-tunnelip` define el túnel de VXLAN y es la dirección IP de Loopback de cada equipo, mejorando la conectividad y redundancia dentro de la red.

Finalmente, el parámetro `vxlan-remoteip` que denota los nodos remotos involucrados en la comunicación. Lo que garantiza que la interconexión de los dispositivos sea correcta, además proporciona una infraestructura eficiente y confiable. Este proceso se replica en los demás equipos

que forman parte de la estructura LEAF, la única diferencia son las direcciones IP que utiliza cada equipo.

Figura 21

Configuración de VXLAN estático en switches LEAF.



```
GNU nano 3.2 /etc/network/interfaces
auto bridge
iface bridge
    bridge-ports      swp3 swp4 vni-10 vni-20
    bridge-vids       10 20
    bridge-vlan-aware yes

auto vni-10
iface vni-10
    bridge-access      10
    mstpctl-bpdguard   yes
    mstpctl-portbpdfilter yes
    vxlan-id           100
    vxlan-local-tunnelip 192.168.0.3
    vxlan-remoteip     192.168.0.2
    vxlan-remoteip     192.168.0.6
    vxlan-remoteip     192.168.0.7
```

3.12 Configuración De Servicios Para Pruebas De Operación De Red

Para llevar a cabo las pruebas del entorno, se instala y configura los servicios necesarios. Esto permite verificar que todos los elementos operen de manera correcta y que el sistema responde la manera óptima.

En este caso, se implementan tres servidores: WEB, FTP y DNS. Estos servicios fueron elegidos porque se utilizan en centro de datos reales y ayudan a simular un escenario cercano al que se encontraría en una red de producción. Los servidores se ejecutan en máquinas virtuales que se basan en Debian 12, ya que es una distribución de Linux que es ligera en el consumo de recursos y compatibilidad.

Con esta infraestructura, se prevé evaluar el rendimiento y el comportamiento de VXLAN en el entorno controlado para verificar los beneficios y desventajas de esta tecnología emergente y así poder aplicarlo en una red real.

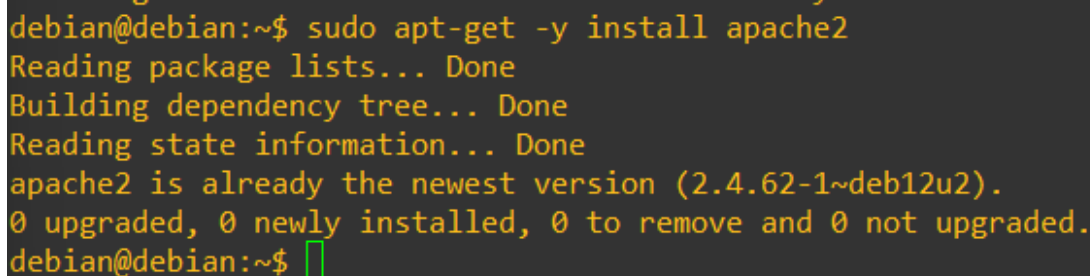
3.12.1 Servidor Web

Para configurar un servidor web en la distribución de Debian 12, lo primero que se recomienda es actualizar el sistema y para eso se hace uso de los siguientes comandos: `sudo apt-get update && upgrade -y`, este proceso se debe de realizar para asegurarse que los paquetes se encuentren en su versión más reciente y así se evita errores de funcionamiento en un futuro.

Una vez que el sistema se encuentre actualizado, se instala el paquete necesario. En este caso, se utiliza Apache2 que es uno de los servidores web más conocidos y confiable en este ámbito. La función principal es alojar y entregar páginas web a los usuarios de manera correcta y en la Figura 22 se puede observar el comando de instalación de Apache2.

Figura 22

Instalación de Apache2.



```
debian@debian:~$ sudo apt-get -y install apache2
Reading package lists... Done
Building dependency tree... Done
Reading state information... Done
apache2 is already the newest version (2.4.62-1~deb12u2).
0 upgraded, 0 newly installed, 0 to remove and 0 not upgraded.
debian@debian:~$
```

Una vez que se instala el Apache2, se configura reglas de firewall en el sistema para permitir el acceso al servicio web. En la línea de código que aparece en la Figura 23, se establece una regla que permite el tráfico entrante en el puerto 80, utilizando el protocolo TCP.

El puerto 80 es el que se encarga de gestionar las solicitudes HTTP, lo que permite que otros dispositivos accedan al servicio y visualicen la página web que se encuentra alojada en el servidor.

Figura 23

Reglas de firewall que permiten el tráfico HTTP.

```
root@debian:/home/debian# ufw allow 80/tcp
Skipping adding existing rule
Skipping adding existing rule (v6)
root@debian:/home/debian# ufw status
Status: active

To Action From
--
80/tcp ALLOW Anywhere
80/tcp (v6) ALLOW Anywhere (v6)
root@debian:/home/debian#
```

Finalmente, ajustamos la directiva de ServerName en el archivo de configuración principal ubicado en `/etc/apache2/apache2.conf`, tal y como se indica en la Figura 24. Esta directiva se emplea para definir el nombre del host o dominio que el servidor debe gestionar y en este caso, el dominio que se asigna es `www.tesisvxlan.com`.

Figura 24

Determinación del nombre del host para el servicio web.

```
GNU nano 7.2 /etc/apache2/apache2.conf *
63 # mounted filesystem then please read the Mutex documentation (available
64 # at <URL:http://httpd.apache.org/docs/2.4/mod/core.html#mutex>);
65 # you will save yourself a lot of trouble.
66 #
67 # Do NOT add a slash at the end of the directory path.
68 #
69 #ServerRoot "/etc/apache2"
70 #ServerName www.tesisvxlan.com
71 #
72 # The accept serialization lock file MUST BE STORED ON A LOCAL DISK.
73 #
74 #Mutex file:${APACHE_LOCK_DIR} default
75 #
76 #
77 # The directory where shm and other runtime files will be stored.
78 #
79
80 DefaultRuntimeDir ${APACHE_RUN_DIR}
81
82 #

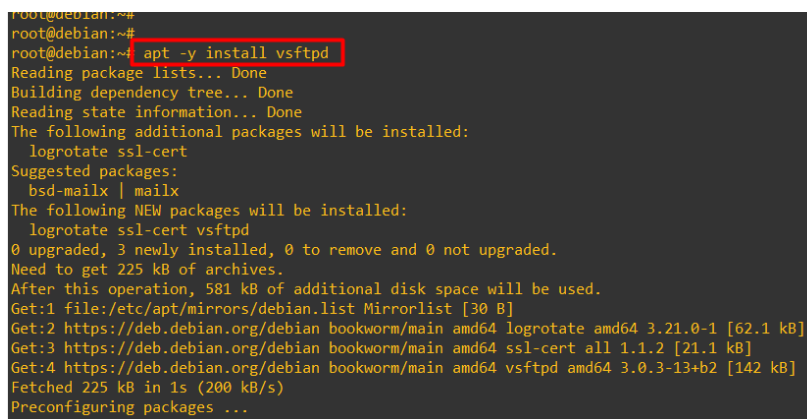
^G Help ^O Write Out ^W Where Is ^K Cut ^T Execute ^C Location
^X Exit ^R Read File ^\ Replace ^U Paste ^J Justify ^_ Go To Line
```

3.12.2 Servidor FTP

File Transfer Protocol (FTP), es un servicio que se utiliza para la transferencia y obtención de archivos entre dispositivos de forma remota, los puertos usados por este protocolo es el 20 y el 21 para la mayoría de los casos. En sistemas operativos basados en Debian 12, la instalación de este servicio se realiza de manera sencilla mediante el comando `apt -y install vsftpd`, como se ilustra en la Figura 25.

Figura 25

Comando usado para la instalación del servicio FTP.



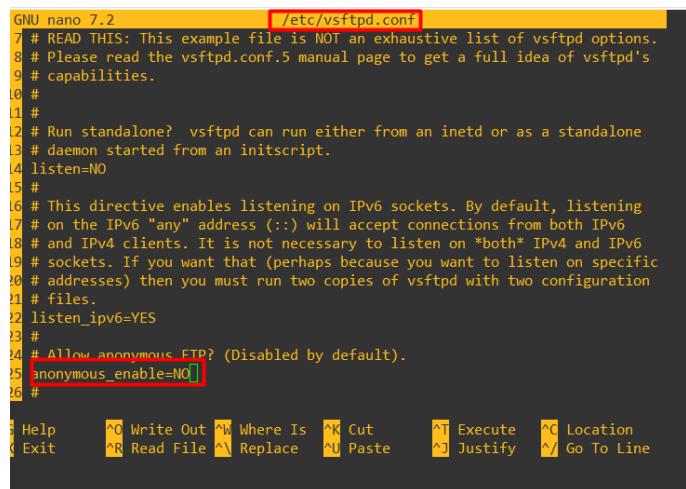
```
root@debian:~#  
root@debian:~#  
root@debian:~# apt -y install vsftpd  
Reading package lists... Done  
Building dependency tree... Done  
Reading state information... Done  
The following additional packages will be installed:  
  logrotate ssl-cert  
Suggested packages:  
  bsd-mailx | mailx  
The following NEW packages will be installed:  
  logrotate ssl-cert vsftpd  
0 upgraded, 3 newly installed, 0 to remove and 0 not upgraded.  
Need to get 225 kB of archives.  
After this operation, 581 kB of additional disk space will be used.  
Get:1 file:/etc/apt/mirrors/debian.list Mirrorlist [30 B]  
Get:2 https://deb.debian.org/debian bookworm/main amd64 logrotate amd64 3.21.0-1 [62.1 kB]  
Get:3 https://deb.debian.org/debian bookworm/main amd64 ssl-cert all 1.1.2 [21.1 kB]  
Get:4 https://deb.debian.org/debian bookworm/main amd64 vsftpd amd64 3.0.3-13+b2 [142 kB]  
Fetched 225 kB in 1s (200 kB/s)  
Preconfiguring packages ...
```

Después de que termine la instalación del servicio, el siguiente paso es configurar su archivo principal con el uso del comando `sudo nano /etc/vsftpd.conf`. En la Figura 26 se muestra una línea de comandos que por defecto esta comentada, en este caso se quita el comentario y se le configura como `Anonymous_enable=No`, el objetivo de esto es que se desactiva el acceso anónimo al servicio FTP.

Esta configuración es muy importante, ya que mantiene la seguridad del servicio y ayuda a evitar accesos no deseados, asegurando que solo los usuarios autorizados puedan hacer el uso del servicio.

Figura 26

Configuración para desactivar el acceso anónimo al servidor FTP.



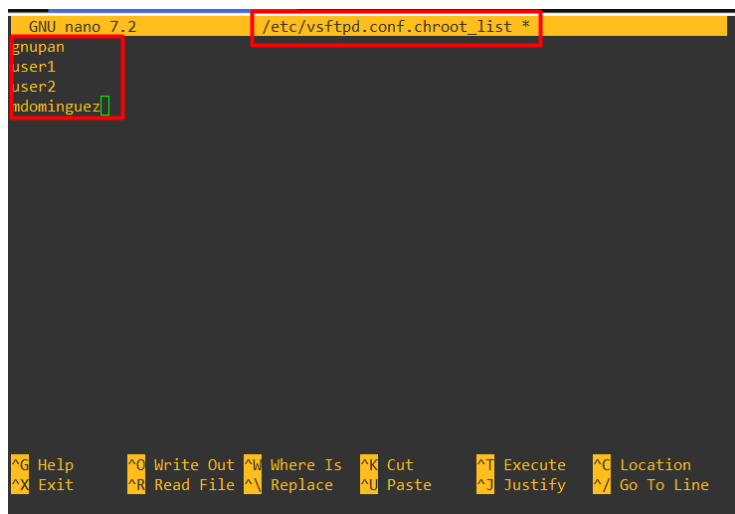
```

GNU nano 7.2 /etc/vsftpd.conf
7 # READ THIS: This example file is NOT an exhaustive list of vsftpd options.
8 # Please read the vsftpd.conf.5 manual page to get a full idea of vsftpd's
9 # capabilities.
10 #
11 #
12 # Run standalone? vsftpd can run either from an inetd or as a standalone
13 # daemon started from an initscript.
14 listen=NO
15 #
16 # This directive enables listening on IPv6 sockets. By default, listening
17 # on the IPv6 "any" address (::) will accept connections from both IPv6
18 # and IPv4 clients. It is not necessary to listen on *both* IPv4 and IPv6
19 # sockets. If you want that (perhaps because you want to listen on specific
20 # addresses) then you must run two copies of vsftpd with two configuration
21 # files.
22 listen_ipv6=YES
23 #
24 # Allow anonymous FTP? (Disabled by default).
25 anonymous_enable=NO
26 #
27
28 Help      ^O Write Out  ^W Where Is   ^K Cut        ^T Execute    ^C Location
29 Exit      ^R Read File  ^\ Replace    ^U Paste      ^J Justify    ^_ Go To Line
  
```

Finalmente, es necesario definir los usuarios que podrán acceder al servicio FTP, lo cual se configura el archivo principal `/etc/vsftpd.conf`, tal y como se indica en la Figura 27. Esta configuración permite incluir una lista específica de usuarios que tendrán permisos para salir de su directorio raíz y acceder a otros recursos del sistema, permitiendo a que los usuarios tengan mayor libertad de navegación para realizar tareas administrativas o de mantenimiento.

Figura 27

Configuración de usuarios para acceso al servicio FTP.



```

GNU nano 7.2 /etc/vsftpd.conf.chroot_list *
gnupan
user1
user2
ndominguez
  
```

3.12.3 Servidor De Base De Datos

Para el presente trabajo se toma en consideración un servidor de base datos porque son sistemas tradicionales que se utiliza dentro de un centro de datos. Según (Pilicita Garrido et al., 2020), MariaDB se considera como una versión optimizada de MySQL, que ofrece alta compatibilidad con sistemas operativos de Linux y ofrece una mejor seguridad al momento de implementar el servicio de base de datos.

Para la implementación de este servidor, se utiliza la distribución de Linux Debian 12. Para la instalación de este servicio se utiliza el comando `apt install mariab-server`, luego de que se haya terminado de instalar verificamos si ya se encuentra en funcionamiento, en la Figura 28 se puede observar los comandos para inicializar el servicio como también el estado en el cual se encuentra.

Figura 28

Servicio de MariaDB funcionando.

```
root@debian:/home/debian# systemctl start mysql
root@debian:/home/debian# systemctl status mysql
● mariadb.service - MariaDB 10.11.11 database server
   Loaded: loaded (/lib/systemd/system/mariadb.service; enabled; preset: enab
   Active: active (running) since Thu 2025-04-03 20:54:34 UTC; 34min ago
     Docs: man:mariadb(8)
           https://mariadb.com/kb/en/library/systemd/
   Main PID: 8490 (mariabdd)
    Status: "Taking your SQL requests now..."
     Tasks: 9 (limit: 3548)
    Memory: 82.8M
       CPU: 1.271s
    CGroup: /system.slice/mariadb.service
            └─8490 /usr/sbin/mariabdd
```

Al verificar que está funcionando, se procede a ingresar dentro de este servicio y para eso se hace uso del comando `mysql -u -p`, al momento de ingresar nos pedirá realizar ciertas configuraciones primarias antes de poder hacer uso del servicio. En la Figura 29, se indica las configuraciones iniciales como el establecimiento de una nueva contraseña de root y también la eliminación de permisos para usuario anónimos.

Figura 29*Configuración inicial de MariaDB.*

```

Change the root password? [Y/n] y
New password:
Re-enter new password:
Password updated successfully!
Reloading privilege tables..
... Success!

By default, a MariaDB installation has an anonymous user, allowing anyone
to log into MariaDB without having to have a user account created for
them. This is intended only for testing, and to make the installation
go a bit smoother. You should remove them before moving into a
production environment.

Remove anonymous users? [Y/n] y

```

Al realizar las respectivas configuraciones iniciales. Se procede a crear una nueva base de datos, en este caso se denomina `test_vxlan_gnupan` para luego poder crear dentro de esta base de datos una nueva tabla la cual permite almacenar información en donde la llave primaria de esta va a ser un ID y los comandos usados para lo indicado se muestra en la Figura 30.

Figura 30*Creación de base de datos.*

```

Database changed
MariaDB [mysql]> create database test_vxlan_gnupan;
Query OK, 1 row affected (0.001 sec)

MariaDB [test_vxlan_gnupan]> Create table test_table(id int, nombre varchar(50), empresa varchar(50), primary key (id));
Query OK, 0 rows affected (0.013 sec)

MariaDB [test_vxlan_gnupan]>

```

Una vez que la base de datos haya sido creada, el siguiente paso consiste en agregar datos dentro de esta. Para ello, se utilizan los comandos que se muestran en la Figura 31, donde se incorporan el ID, nombre de la empresa y nombre de una persona, como parte de proceso de prueba.

Figura 31

Agregación de información en la Tabla de la Base de datos.

```
MariaDB [test_vxlan_gnupan]> insert into test_table(id, nombre, empresa) values
(1, 'gnupan', 'utn');
Query OK, 1 row affected (0.002 sec)

MariaDB [test_vxlan_gnupan]> []
```

Finalmente, en la Figura 32 se muestra como agregar un nuevo usuario a la base de datos. A este usuario se le dan todos los privilegios necesarios para poder leer y escribir en la base de datos, lo que significa que tiene la capacidad de crear y modificar la información que se encuentra dentro de esta.

Figura 32

Creación de usuarios para ingreso a la base de datos.

```
MariaDB [(none)]> CREATE USER 'gnupan'@'%' IDENTIFIED BY 'root';
Query OK, 0 rows affected (0.006 sec)

MariaDB [(none)]> GRANT ALL PRIVILEGES ON *.* TO 'gnupan'@'%' IDENTIFIED BY 'contraseña' WITH GRANT OPTION;
Query OK, 0 rows a
```

4 CAPÍTULO IV – EVALUACIÓN DE RESULTADOS

En este capítulo, se exponen y examina los resultados obtenidos en el trabajo de grado, cuyo objetivo es evaluar el comportamiento de VXLAN en centro de datos que usan arquitectura Spine-Leaf. El propósito es evidencia el desempeño de VXLAN a través de importantes pruebas que ayudan a evaluar métricas como ancho de banda, latencia, escalabilidad y más aspectos que sean importantes que permitan obtener resultados claves.

A lo largo de este análisis, se busca determinar con el entorno de pruebas planteado, cuáles son los beneficios y posibles limitaciones de VXLAN. Además, al realizar el análisis se observa patrones y tendencias que permiten sacar conclusiones sólidas que respaldan la viabilidad de la solución propuesta, sino que también dan una idea de cómo el ambiente de red propuesto pueda comportarse en un ambiente de producción real.

4.1 Validación De Funcionamiento Del Protocolo VXLAN En El Entorno De Pruebas

Para garantizar la operatividad del entorno de pruebas propuesto, se realiza pruebas funcionales que se enfocan en verificar la configuración de los nodos y el establecimiento de lo que vendría siendo los túneles VXLAN. Estas pruebas permiten validar que la topología implementada en el software de GNS3, se encuentre funcionando de la manera esperada para luego proceder con la evaluación del desempeño de VXLAN.

En primer lugar, se comprueba la conectividad de capa 3 existente entre los dispositivos SPINE y LEAF mediante el protocolo de enrutamiento BGP. En la Figura 33, indica el comando NCLU que se utiliza para la verificación de este caso, lo que permite observar las múltiples rutas dando a conocer que BGP está funcionando.

Cada entrada de la tabla de enrutamiento indica información importante como:

- BGP router identifier: es la dirección IP que identifica al router en el proceso. En este caso es la dirección de Loopback del SPINE 01.
- Local AS number: es el número de sistema autónomo local del equipo, en este caso se configuro el 65199. Los AS se usan para identificar un grupo de equipos que tienen en común una política de enrutamiento.
- Vrf-id 0: es el identificador de la VRF (Virtual Routing Forwarding). En este caso el valor de 0 es porque se usa una tabla de enrutamiento de forma global.
- RIB entries: indica el número de entradas en la base de datos de rutas BGP, en el caso puesto de ejemplo son 13 veces.
- 2600 bytes of memory: indica la memoria utilizada para almacenar las entradas a BGP.
- Peers: indica el número de vecinos BGP con los que el equipo se encuentra intercambiando información. En este caso, la cantidad de vecinos es de 5.
- Neighbor: es el nombre del vecino BGP. En este caso tenemos el hostname del equipo vecino y la interfaz por el cual se comunican.
- V: indica la versión del protocolo BGP que está siendo usada. Conmunmente es usada la versión 4.
- MsgRcvd: indica el número de mensajes BGP recibidos del vecino.
- MsgSent: indica cuantos mensajes es enviado hacia el vecino.
- State/PfxRcd: muestra el estado de la sesión BGP, seguido del número de rutas recibidas desde ese vecino, en algunos casos se tiene 2 o 3 rutas aprendidas.

PfxSnt: indica el número de prefijos enviados al vecino. El valor es de 7 para todos los vecinos, es decir 7 rutas a cada uno de ellos.

Figura 33

Comprobación de vecinos BGP en el entorno de pruebas.

```
cumulus@spine-01:mgmt:~$ net show bgp summary
show bgp ipv4 unicast summary
=====
BGP router identifier 192.168.0.1, local AS number 65199 vrf-id 0
BGP table version 7
RIB entries 13, using 2600 bytes of memory
Peers 5, using 114 KiB of memory

Neighbor      V      AS  MsgRcvd  MsgSent  TblVer  InQ  OutQ  Up/Down  State/PfxRcd  PfxSnt
leaf-01(swp1) 4    65101    174     175      0    0    0 00:07:40      2         7
leaf-02(swp2) 4    65102    174     173      0    0    0 00:07:40      3         7
leaf-03(swp3) 4    65103    177     181      0    0    0 00:07:34      2         7
leaf-04(swp4) 4    65104    167     167      0    0    0 00:07:30      2         7
leaf-05(swp5) 4    65105    167     167      0    0    0 00:07:30      2         7

Total number of neighbors 5
```

Para verificar las interfaces VXLAN creadas, se utiliza el comando `ip -d link show | grep vxlan`. En la Figura 34 se muestra la información detallada del LEAF 4, que se lo utilizo como ejemplo. En ella, se puede visualizar detalles como el ID del VNI, la dirección IP utilizada para encapsular en la red Underlay (Corresponde a la dirección de Loopback), el puerto UDP utilizado para encapsular VXLAN (definido por la IANA) y finalmente, el tiempo en segundos para la caducidad de las direcciones MAC dentro del túnel VXLAN.

Figura 34

Información de túnel VXLAN en LEAF 4.

```
cumulus@leaf-04:mgmt:~$ ip -d link show | grep vxlan
vxlan id 300 local 192.168.0.6 srcport 0 0 dstport 4789 nolearning ttl 64 ageing 300 udpchecksum noudp6zerocsumtx
vxlan id 100 local 192.168.0.6 srcport 0 0 dstport 4789 nolearning ttl 64 ageing 300 udpchecksum noudp6zerocsumtx
vxlan id 200 local 192.168.0.6 srcport 0 0 dstport 4789 nolearning ttl 64 ageing 300 udpchecksum noudp6zerocsumtx
cumulus@leaf-04:mgmt:~$
```

Para verificar el funcionamiento de VXLAN, es necesario visualizar las direcciones MAC, como se muestra en la Figura 35. En ella, se puede observar que cada una de las VNI está asociado a su túnel de destino. En este caso, la VNI 10 y VNI 20, correspondiente a la VXLAN 100 y VXLAN 200, respectivamente, tienen como dirección MAC la de los servidores WEB para VNI

10 y FTP para VNI 20. Esta visualización es útil, ya que permite asegurar que las direcciones MAC correctas están siendo aprendidas y propagadas a través de VXLAN.

Figura 35

Información detallada de puente MAC en LEAF 1.

```
cumulus@leaf-01:mgmt:~$ net show bridge macs
```

VLAN	Master	Interface	MAC	TunnelDest	State	Flags	LastSeen
1	bridge	bridge	0c:d0:04:b2:00:03		permanent		00:24:33
untagged		vni-10	00:00:00:00:00:00	192.168.0.6	permanent	self	00:24:16
untagged		vni-10	00:00:00:00:00:00	192.168.0.7	permanent	self	00:24:16
untagged		vni-20	00:00:00:00:00:00	192.168.0.6	permanent	self	00:24:16
untagged	bridge	swp3	0c:d0:04:b2:00:03		permanent		00:24:33
untagged	bridge	swp4	0c:d0:04:b2:00:04		permanent		00:24:33
untagged	bridge	vni-10	96:51:c0:c2:4e:4e		permanent		00:24:33
untagged	bridge	vni-20	ce:87:03:03:51:f2		permanent		00:24:33

```
cumulus@leaf-01:mgmt:~$
cumulus@leaf-01:mgmt:~$
```

Como prueba final para verificar el correcto funcionamiento de VXLAN en la topología Spine-Leaf, se hace uso del comando `sudo tcpdump -i <interfaz> port 4789 -nn`. Este comando sirve para capturar y analizar el tráfico de red en una interfaz seleccionada, enfocándose en los paquetes que usan el puerto UDP 4789, que es el puerto utilizado por VXLAN para encapsular la información. Gracias a esta captura, se observa si los paquetes VXLAN están siendo transmitidos de forma correcta entre los nodos que conforman la red.

En la Figura 36, se muestra la salida obtenida tras la ejecución del comando, donde se puede mirar información importante como la dirección IP de origen y destino de los paquetes, el VNI al que pertenecen y el tráfico bidireccional entre los dos nodos. Usar esta herramienta permite confirmar que la red Spine-Leaf con VXLAN se encuentra funcionando correctamente, ya que los paquetes están siendo entregados a sus destinos finales sin ningún problema.

Figura 36

Captura de paquete VXLAN con tcpdump.

```
cumulus@leaf-04:mgmt:~$ sudo tcpdump -i swp1 port 4789 -nn
[sudo] password for cumulus:
Sorry, try again.
[sudo] password for cumulus:
tcpdump: verbose output suppressed, use -v or -vv for full protocol decode
listening on swp1, link type EN10MB (Ethernet), capture size 262144 bytes
23:57:38.380919 IP 192.168.0.6.53374 > 192.168.0.4.4789: VXLAN, flags [I] (0x08), vni 300
IP6 :: > ff02::16: HBH ICMP6, multicast listener report v2, 1 group record(s), length 28
23:57:38.759093 IP 192.168.0.6.53374 > 192.168.0.4.4789: VXLAN, flags [I] (0x08), vni 300
IP6 :: > ff02::16: HBH ICMP6, multicast listener report v2, 1 group record(s), length 28
23:57:39.464046 IP 192.168.0.6.35047 > 192.168.0.4.4789: VXLAN, flags [I] (0x08), vni 300
IP6 fe80::e93:c2ff:fe78:0 > ff02::16: HBH ICMP6, multicast listener report v2, 1 group record(s), length 28
23:57:39.464163 IP 192.168.0.6.40970 > 192.168.0.4.4789: VXLAN, flags [I] (0x08), vni 300
IP6 fe80::e93:c2ff:fe78:0 > ff02::2: ICMP6, router solicitation, length 16
```

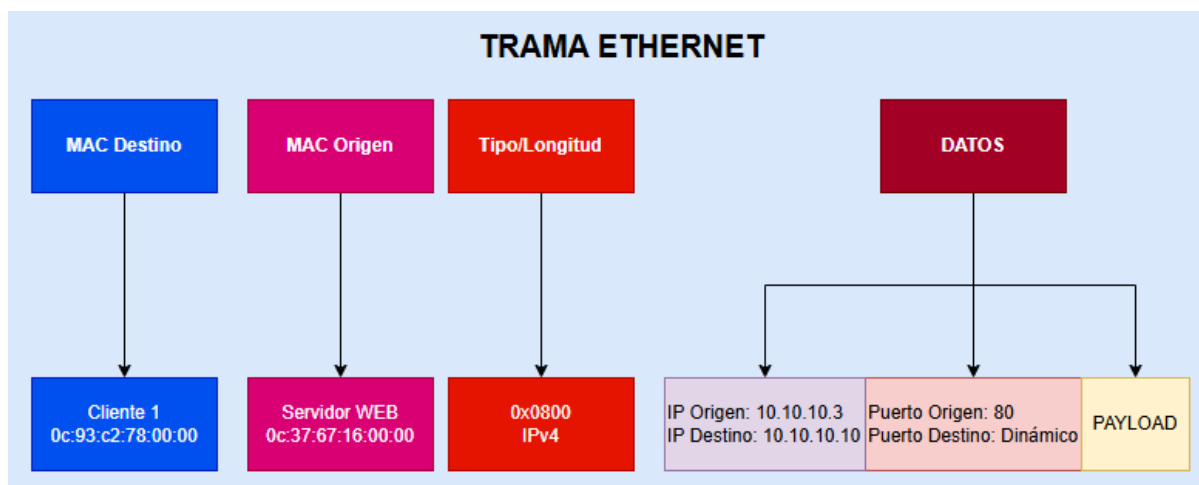
Para comprender el funcionamiento del entorno de pruebas de VXLAN, se realiza un ejemplo práctico. El servidor Web que se encuentra conectado al LEAF 1 desea comunicarse con el Cliente 1, el cual se encuentra conectado al LEAF 4. Este ejemplo permite comprender como es el proceso de encapsulación y desencapsulación de VXLAN en este entorno.

En primer lugar, se tiene el proceso de encapsulación. El servidor Web genera datos, comúnmente conocidos como el payload. Este payload se encapsula dentro de un segmento, al cual añade una cabecera TCP. Esta cabecera asegura una conexión confiable, ya que utiliza la comunicación por tres vías (Three-way Handshake). Durante este paso se asignan los puertos de origen y destino.

A continuación, se agrega una cabecera IP. La dirección IP de origen será la del servidor WEB, que se encuentra en la VLAN 10 y tiene la dirección IP 10.10.10.3/24. La dirección IP de destino es la del Cliente 1, la cual es la 10.10.10.10/24 indicando que es de la misma VLAN. Finalmente, se agrega la dirección MAC de origen (0c:37:67:16:00:00) y la dirección MAC de destino (0c:93:c2:78:00:00), estas direcciones MAC corresponden al servidor de base de datos y al Cliente 1. Así, la trama Ethernet original se forma tal y como se indica en la Figura 37.

Figura 37

Trama Ethernet originada por el servidor de base de datos.



Cuando la trama Ethernet generada por el servidor WEB sale y llega al switch de acceso, en este caso el LEAF 1, el equipo detecta que el tráfico está destinado a la red Overlay que sería la de VXLAN, esto lo hace basándose en las configuraciones previas realizadas en este equipo, como las asignaciones de VLAN y VXLAN. El LEAF 1 verifica la trama que llegó y determina que el tráfico debe de ser encapsulado en el formato de VXLAN.

En este punto empieza el proceso de encapsulación. El LEAF 1 agrega un VNI que permite identificar de manera exclusiva a la red VXLAN. El valor del VNI se basa en la tabla de mapeo que esta almacenada en el LEAF 1, esto quiere decir que el equipo consulta que VNI corresponde a que VLAN. Por ejemplo, en la Tabla 9 se muestra que VNI corresponde a la VLAN, según las configuraciones y determinaciones del entorno de prueba.

Tabla 9

Asociación de VLAN con VNI.

VLAN	10	20	30
VXLAN	100	200	300

A continuación, el LEAF 1 ya conoce que el tráfico proviene de un equipo que está conectado a la VLAN 10, entonces este se asigna el VNI 100, que se encuentra mapeado a dicha VLAN. Posteriormente, el LEAF 1 agrega el encabezado UDP, que es importante en el encapsulamiento VXLAN. En este encabezado, se asigna el puerto de origen de forma dinámica, mientras que el puerto de destino es el más importante, ya que este debe de asignarse el puerto estándar de VXLAN definido por la IANA que es el 4789.

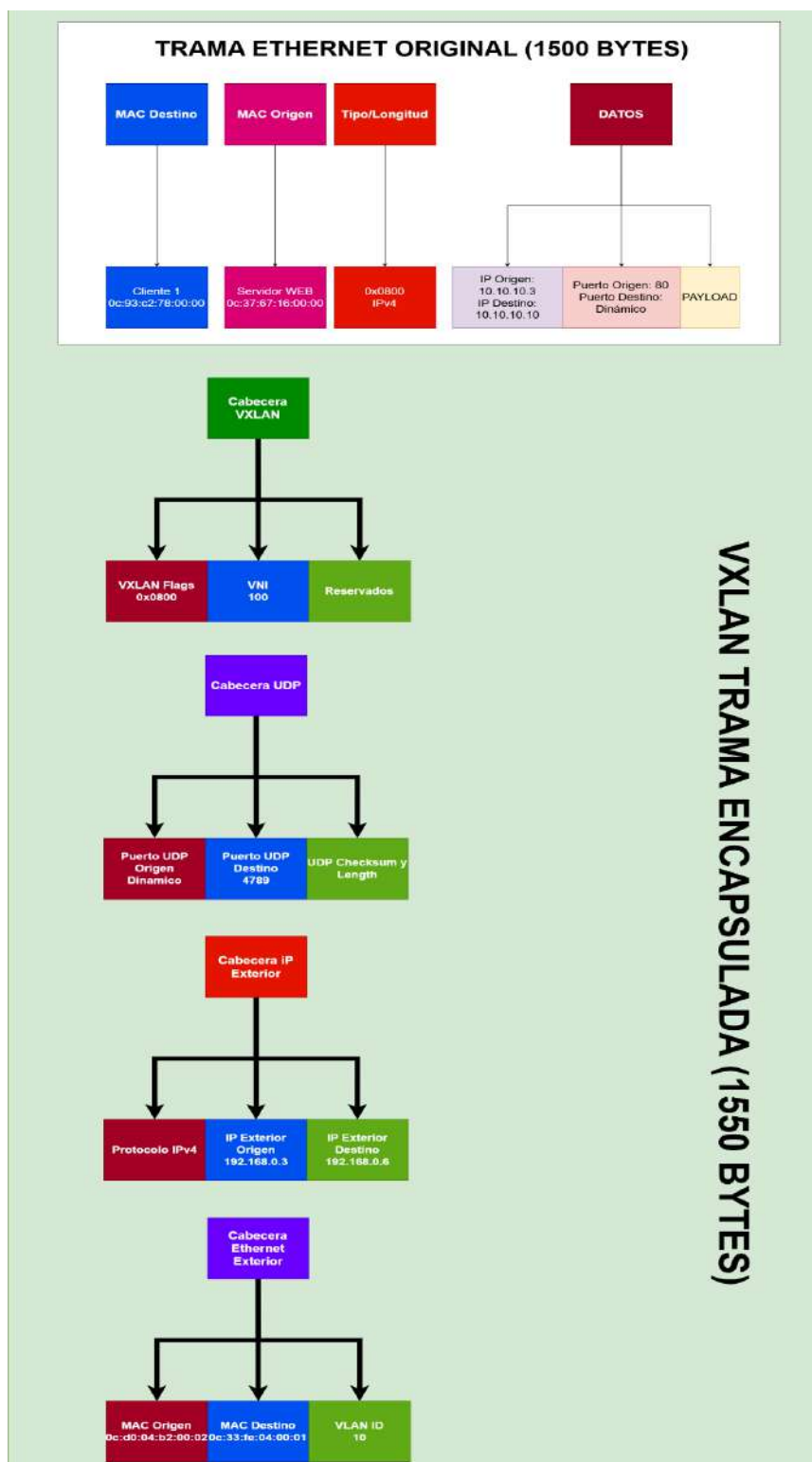
A continuación, se añade el encabezado IP. La dirección IP de origen es la dirección de Loopback del LEAF 1, la cual es la 192.168.0.3/32. Esta dirección IP de Loopback se utiliza como el punto de salida del tráfico VXLAN o comúnmente conocido como el VTEP. La dirección IP de destino será la dirección IP de Loopback del LEAF 4, que es la 192.168.0.6. Esta dirección IP corresponde al VTEP de este equipo.

Finalmente, se agrega una cabecera Ethernet la cual contiene la dirección MAC de origen y destino. Algo que se debe de recalcar es que este valor va a ir cambiando, dependiendo de los saltos que se haga. Entonces, en este caso la dirección MAC de origen sería la de la interfaz de red que conecta al LEAF 1 con el SPINE 2 (0c:d0:04:b2:00:02) y la dirección MAC de destino es la interfaz de red del SPINE 2 que conecta con el LEAF 1 (0c:33:fe:04:00:01). A medida que el paquete sigue su recorrido, comienza el proceso inverso: la desencapsulación del tráfico hasta llegar al Cliente 1, el cual es el destino final de la comunicación con el servidor WEB.

En la Figura 38, se puede ver el paquete completamente encapsulado en VXLAN, con toda la información relacionada a cada uno de los dispositivos involucrados en la comunicación.

Figura 38

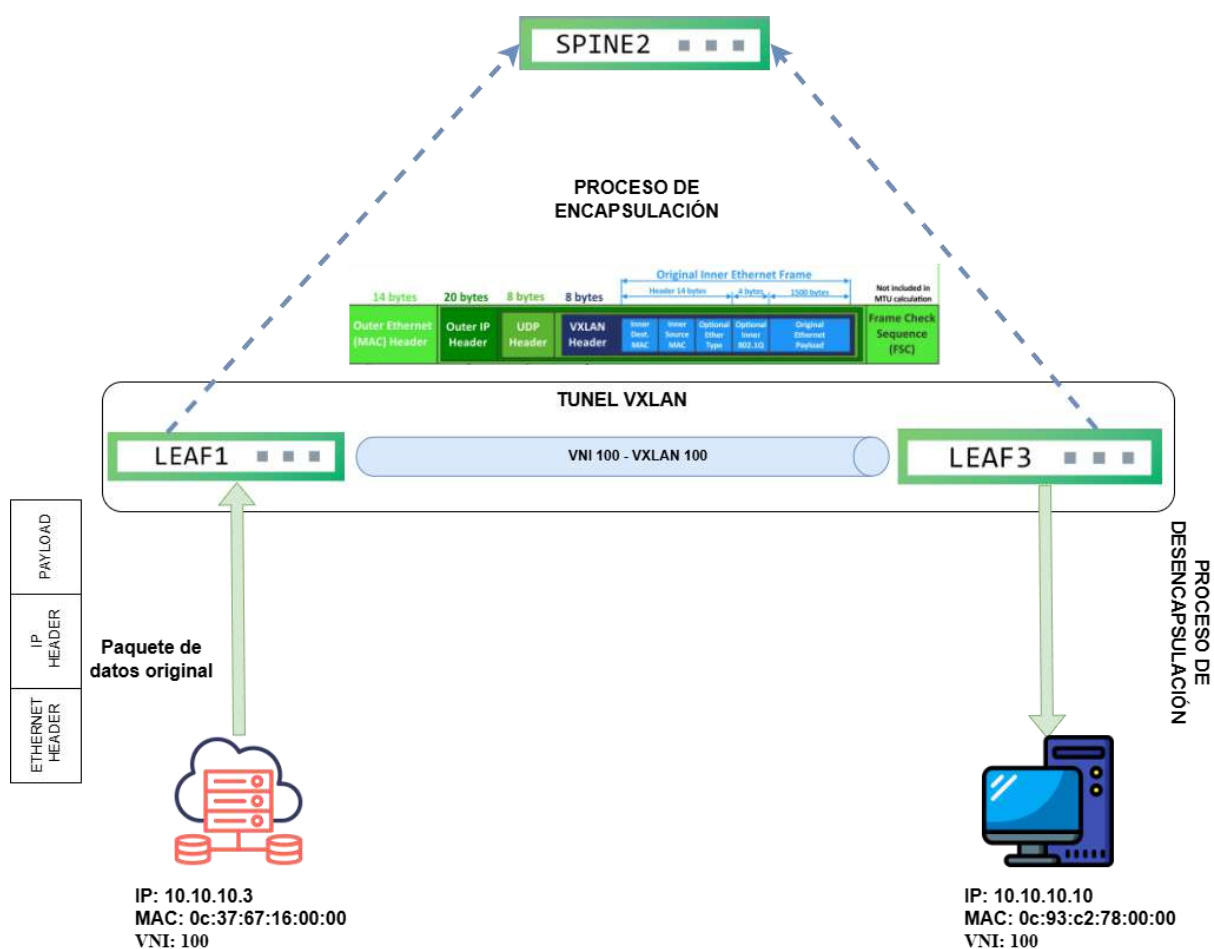
Encapsulación VXLAN en la comunicación de servidor WEB hasta el Cliente 1.



En la Figura 39, se muestra una representación gráfica de los dispositivos involucrados en la comunicación entre el servidor WEB y Cliente 1, en donde se detalla como es el inicio de la comunicación generando la trama Ethernet original. Posteriormente se inserta los encabezados necesarios para poder tener la encapsulación de VXLAN, además indica como es el proceso de tunelización en la red Overlay. Esta representación gráfica permite entender de manera visual como se transporta la información en una red VXLAN.

Figura 39

Proceso de comunicación de VXLAN.



4.2 Captura y Análisis De Tráfico VXLAN Usando Wireshark

Para analizar cómo se transmiten los paquetes VXLAN en la red, se hizo el uso del Sniffer Wireshark, ya que permite capturar y examinar los datos en una prueba de conectividad. La primera prueba para validar resultados sobre VXLAN es realizar un ping desde un cliente que se encuentra conectado al LEAF 4 hacia el servidor WEB conectado al LEAF 1, capturando paquetes de solicitud y respuesta. Se optó por utilizar el protocolo ICMP en lugar de otros protocolos para realizar la prueba porque la función específica de este es verificar la conectividad y diagnosticar si la red funciona. En este caso, permite identificar que el paquete que se está transmitiendo viaja correctamente, lo que permite validar la comunicación de equipos mediante VXLAN.

En la Figura 40 se muestra la captura de los paquetes que se están transmitiendo en la red, en este caso se va a analizar el paquete 95 correspondiente a una solicitud de ping request enviada desde la dirección IP 10.10.10.10 (Cliente1 – LEAF 4) hacia 10.10.10.3 (Servidor WEB – LEAF 1). En la inspección del paquete, se evidencia un tamaño de 148 bytes e incluye otra información como número de secuencia y el TTL (Time to Live).

Figura 40

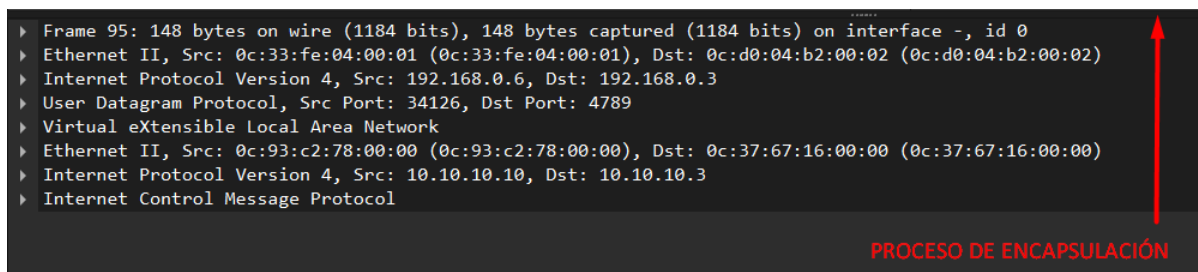
Paquete Capturado en Wireshark.

No.	Time	Source	Destination	Protocol	Length	Info
95	25.782558	10.10.10.10	10.10.10.3	ICMP	148	Echo (ping) request id=0x019e, seq=1/256, ttl=64 (reply in 96)
96	25.786280	10.10.10.3	10.10.10.10	ICMP	148	Echo (ping) reply id=0x019e, seq=1/256, ttl=64 (request in 95)

Para iniciar con el análisis del paquete, se identifica el tipo de proceso se va a realizar, que en este caso corresponde a la encapsulación. Esto permite comprobar lo señalado en la sección anterior. En la Figura 41 se presenta el procedimiento mediante el cual se analizará el paquete capturado en Wireshark.

Figura 41

Identificación de encapsulación de información.



```

▶ Frame 95: 148 bytes on wire (1184 bits), 148 bytes captured (1184 bits) on interface -, id 0
▶ Ethernet II, Src: 0c:33:fe:04:00:01 (0c:33:fe:04:00:01), Dst: 0c:d0:04:b2:00:02 (0c:d0:04:b2:00:02)
▶ Internet Protocol Version 4, Src: 192.168.0.6, Dst: 192.168.0.3
▶ User Datagram Protocol, Src Port: 34126, Dst Port: 4789
▶ Virtual eXtensible Local Area Network
▶ Ethernet II, Src: 0c:93:c2:78:00:00 (0c:93:c2:78:00:00), Dst: 0c:37:67:16:00:00 (0c:37:67:16:00:00)
▶ Internet Protocol Version 4, Src: 10.10.10.10, Dst: 10.10.10.3
▶ Internet Control Message Protocol

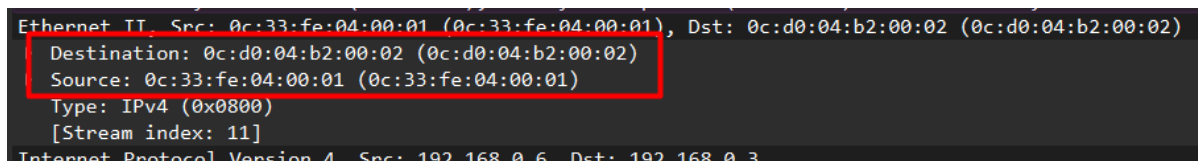
```

PROCESO DE ENCAPSULACIÓN

En la capa de enlace de datos externa se muestra las direcciones MAC, las cuales permiten identificar el origen y destino del tráfico, tal y como se puede apreciar en la Figura 42. La dirección 0c:33:fe:04:00:01 indica el equipo origen, mientras que 0c:d0:04:b2:00:02 es la dirección de destino.

Figura 42

Capa de enlace de datos externa.



```

Ethernet II, Src: 0c:33:fe:04:00:01 (0c:33:fe:04:00:01), Dst: 0c:d0:04:b2:00:02 (0c:d0:04:b2:00:02)
  Destination: 0c:d0:04:b2:00:02 (0c:d0:04:b2:00:02)
  Source: 0c:33:fe:04:00:01 (0c:33:fe:04:00:01)
  Type: IPv4 (0x0800)
  [Stream index: 11]
Internet Protocol Version 4, Src: 192.168.0.6, Dst: 192.168.0.3

```

En la capa de red externa que se muestra en la Figura 43, el paquete contiene direcciones IP privadas, donde 192.168.0.2 es la dirección de origen y 192.168.0.3 la de destino. De acuerdo con las configuraciones realizadas y la información en la Tabla 7, las direcciones IP corresponden a las interfaces de Loopback de los dispositivos LEAF 4 y LEAF 1, lo que indica que el tráfico se encuentra en el entorno controlado y no se expone a redes públicas.

Figura 43

Información de la capa de red externa.

```
Internet Protocol Version 4, Src: 192.168.0.6, Dst: 192.168.0.3
  0100 .... = Version: 4
  .... 0101 = Header Length: 20 bytes (5)
  ▶ Differentiated Services Field: 0x00 (DSCP: CS0, ECN: Not-ECT)
    Total Length: 134
    Identification: 0x16b8 (5816)
  ▶ 000. .... = Flags: 0x0
    ...0 0000 0000 0000 = Fragment Offset: 0
    Time to Live: 63
    Protocol: UDP (17)
    Header Checksum: 0xe355 [validation disabled]
    [Header checksum status: Unverified]
    Source Address: 192.168.0.6
    Destination Address: 192.168.0.3
    [Stream index: 0]
```

Finalmente, para terminar con las capas exteriores. La Figura 44, indica la capa de transporte, se observa el uso del protocolo UDP, el cual está utilizando el puerto estándar 4789, definido por la IANA. Esto confirma que el tráfico está encapsulado mediante VXLAN.

Figura 44

Capa transporte.

```
User Datagram Protocol, Src Port: 34126, Dst Port: 4789
  Source Port: 34126
  Destination Port: 4789
  Length: 114
  Checksum: 0x2f53 [unverified]
  [Checksum Status: Unverified]
  [Stream index: 2]
  [Stream Packet Number: 1]
  ▶ [Timestamps]
  UDP payload (106 bytes)
```

Dentro de la encapsulación VXLAN, se tiene el VXLAN Network Identifier (VNI) con un valor de 100, indicando la segmentación lógica del tráfico de este segmento, lo que hace la coexistencia entre diferentes redes virtuales sobre la misma infraestructura física. La Figura 45 indica esta información, además dentro de VXLAN, existe varias banderas que brindan

información relevante que da a conocer el comportamiento y procesamiento del paquete dentro del túnel. A continuación, se detallan sus funciones:

- **GBP Extension:** Indica si la Group-Based Policy (GBP) está habilitada. En este caso, no está definida, ya que no se aplican políticas basadas en grupos dentro del tráfico VXLAN.
- **VXLAN Network ID (VNI):** Es la bandera que indica que el VNI es correcto, permitiendo que continúe la transmisión.
- **Don't Learn:** En este caso, el valor es False, lo que permite el aprendizaje automático de direcciones MAC. Si fuera True, el switch o router no aprendería automáticamente las direcciones MAC.
- **Policy Applied:** Indica si se ha aplicado alguna política de control dentro de la transmisión del paquete. En este caso, muestra un valor de falso, ya que no hay políticas activas para un tipo de tráfico en particular.

Reserved 0x0000: Este campo, muestra que se encuentra reservado para que en caso de que exista una futura implementación específicas de proveedores. En este caso, no contiene información importante.

Figura 45

Información de encapsulamiento de VXLAN.

```
Virtual eXtensible Local Area Network
▼ Flags: 0x0800, VXLAN Network ID (VNI)
  0... .... = GBP Extension: Not defined
  .... 1... = VXLAN Network ID (VNI): True
  .... .... .0.. = Don't Learn: False
  .... .... .... 0... = Policy Applied: False
  .000 .000 0.00 .000 = Reserved(R): 0x0000
Group Policy ID: 0
VXLAN Network Identifier (VNI): 100
Reserved: 0
```

Finalmente, en la Figura 46 se indica la trama Ethernet Original, antes de ser encapsulada en VXLAN, en donde se tiene la dirección MAC origen como destino que en este caso es la del servidor WEB y la del Cliente 1, además de sus direcciones IPv4.

Figura 46

Trama Ethernet Original.

```
Ethernet II, Src: 0c:93:c2:78:00:00 (0c:93:c2:78:00:00), Dst: 0c:37:67:16:00:00 (0c:37:67:16:00:00)
Internet Protocol Version 4, Src: 10.10.10.10, Dst: 10.10.10.3
Internet Control Message Protocol
```

4.3 Pruebas De Rendimiento Del Entorno De Pruebas Basadas En La UIT-T Y.1540

Con el fin de asegurar el funcionamiento eficiente y de calidad de VXLAN en el entorno de pruebas propuesto, es importante evaluar el rendimiento de la red bajo ciertas condiciones operativas. En este contexto la (ITU, 2019) define la recomendación UIT-T Y.1540, denominada “Servicio de comunicaciones de datos con protocolo Internet – Parámetros de calidad de funcionamiento relativos a la disponibilidad y la transferencia de paquetes de protocolo Internet” y su objetivo es establecer los parámetros que se pueden establecer para poder comprobar la calidad del desempeño en términos de velocidad, precisión, fiabilidad operativa como también la disponibilidad en la transmisión de paquetes IP de los servicios de comunicación de datos regionales e internacionales con protocolo Internet.

En la presente investigación, se selecciona algunos parámetros que son: La variación del retardo de paquetes (IPDV), el ancho de banda (AB), la tasa de pérdida de paquetes (IPRL), disponibilidad de la red y la Utilización de los recursos.

La selección de los parámetros se rige en las siguientes razones:

1. **Tasa de Duplicación de paquetes (IPDR):** este parámetro tiene como propósito vincular con el número de segmentos TCP duplicados, ya que la presencia de este tipo de tráfico indica que hay problemas de transmisión. Estos problemas suelen causar congestión en la red, lo que afecta negativamente el rendimiento, especialmente en entornos virtualizados que utilizan túneles y capas de encapsulación como es VXLAN.
2. **Variación del retardo de los paquetes (IPDV):** también denominado como Jitter, es una medida importante en las redes, ya que garantiza un flujo constante de datos en aplicaciones que utilizan especialmente el protocolo UDP.
3. **Tasa de pérdida de paquetes (IPRL):** es un elemento importante en lo que es calidad de la red, ya que puede afectar gravemente al rendimiento de aplicaciones. En el tema de VXLAN, es importante monitorear este parámetro para identificar posibles problemas de la red, ya que los paquetes viajan a través de túneles virtuales.
4. **Disponibilidad de la red:** este parámetro se refiere a la capacidad que tiene una red de estar operativa y accesible en todo momento. En redes virtualizadas como es VXLAN, se requiere una infraestructura que sea robusta y confiable para asegurar una alta disponibilidad.
5. **Utilización de recursos:** nos indica la cantidad de uso de los recursos de la red usados como el procesador y memoria RAM, ya que dentro de la red virtualizada los recursos deben ser los adecuados para evitar la sobrecarga, congestión y fallas en la red.

- 6. Ancho de banda:** es un parámetro que mide la capacidad de transmisión de información de una red, en el contexto de red virtualizada, es esencial medir el ancho de banda para garantizar que la infraestructura esté transmitiendo datos de manera eficiente.

Estos parámetros de medición se eligieron por la importancia que tienen al momento de evaluar redes en donde se ha implementado tecnologías nuevas como VXLAN, ya que son factores importantes que permiten mantener la calidad de servicio y estabilidad de la red, ya que estos parámetros hacen que la investigación tenga un análisis detallado de las condiciones operativas de la red en diferentes escenarios.

4.3.1 Planteamiento De Pruebas

La recomendación UIT-T Y.1540 proporciona un marco estandarizado del proceso de pruebas que deben seguirse para determinar si un servicio IP cumple con los requisitos de disponibilidad y eficiencia esperados de acuerdo con los parámetros seleccionados. Los pasos que guían el despliegue de las pruebas son los siguientes:

A. Determinación del diagrama de referencia

Esta fase se centra en definir el diagrama de referencia en el esquema de pruebas. Para ello es importante establecer los conceptos y terminología que la recomendación de la UIT proporciona, a fin de definir de manera precisa los elementos que participaran de la prueba. En la Tabla 10, presenta la información correspondiente, brindando una visión clara y estructurada de la topología de red en donde se realiza la prueba.

Tabla 10

Terminología para establecer en el diagrama del entorno de pruebas.

Nombre	Abreviatura	Definición
Dispositivo de origen principal	SRC	Es un dispositivo que inicia el envío de los paquetes IP, actuando como punto de inicio.
Dispositivo de destino principal	DST	Es el dispositivo que recibe los paquetes IP, completando el trayecto de la transmisión.
Encaminador	ER	Se encarga de dirigir y reenviar los paquetes IP entre diferentes dispositivos.
Enlace de acceso central	EL	Este enlace conecta un dispositivo SRC o DST a un encaminador, también denominado enlace de acceso.
Conjunto de red	NS	Es el conjunto de dispositivos interconectados que, en conjunto, proporcionan servicios de comunicación basados en IP.
Punto de medición	MP	Es el límite entre un dispositivo y una conexión adyacente, donde se mide y se observa el funcionamiento de la red.

Nota. Información tomada y adaptada de (ITU, 2019).

Una vez clarificadas las terminologías esenciales para establecer el diagrama, se establecen dichas terminologías en la ilustración que se muestra en la Figura 47. Haciendo énfasis en los dispositivos que participaran en las pruebas de rendimiento. Los equipos de origen (SRC) y destino (DST) se determinan según la topología general de la red, representada en la Figura 9.

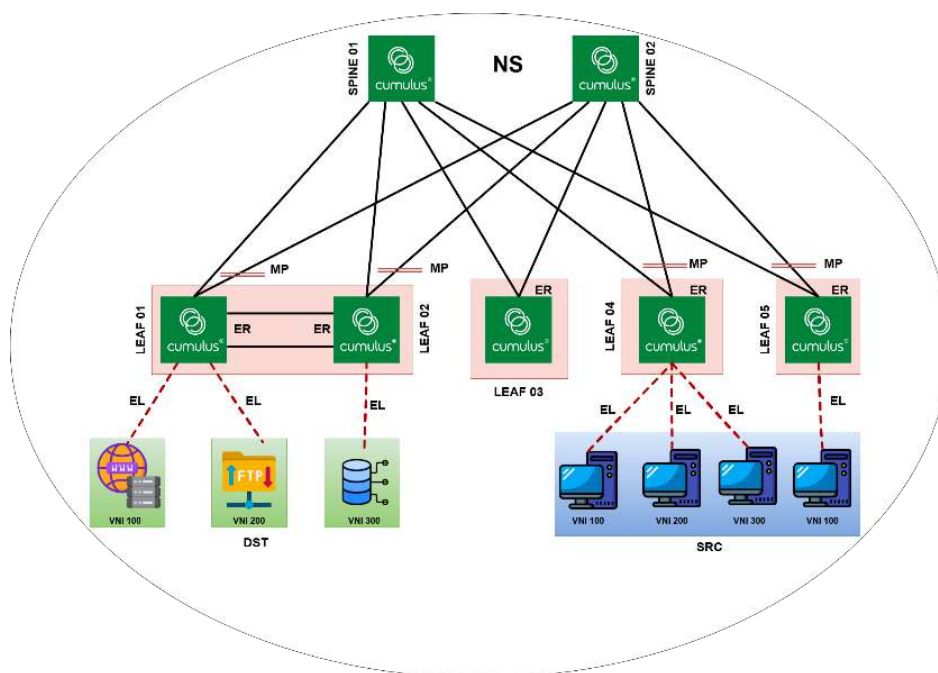
La asignación de los roles SRC y DST para la evaluación del rendimiento de la red se realiza a partir de los hosts y servidores ubicados en diferentes segmentos lógicos de la red. Los dispositivos con el rol de DST están marcados con un recuadro de color verde y corresponde a los servidores, estos dispositivos se les da ese rol porque tienen la capacidad de gestionar y responder al tráfico generado por los dispositivos SRC.

Por otro lado, los dispositivos que cumplen con el rol de SRC, están indicados con un recuadro de color azul y son los clientes. Estos dispositivos pueden generar varios tipos de tráfico y patrones de uso que simulan condiciones de un entorno de producción real.

El enlace central (EL) conecta los dispositivos de origen y destino a los equipos LEAF según corresponda. También los encaminadores (ER) que permiten la comunicación entre otros computadores, en este caso serían los LEAF y finalmente, los puntos de medición (MP) se ubican en los enlaces que conectan los LEAF con los SPINE, y en este caso, se prioriza los enlaces que conectan al SPINE 2, debido a la configuración de costo del enlace lo que permite obtener datos representativos del desempeño de VXLAN, mientras que los enlaces que conecta los LEAF al SPINE 1, no se toman como puntos de medición, ya que estos entran en funcionamiento siempre y cuando exista fallos en un enlace del SPINE 2.

Figura 47

Esquema de pruebas con nomenclatura según la UIT-T Y.1540.



B. Indicadores de rendimiento de transferencia de paquetes IP y herramientas de análisis

Aquí en este paso, se elige los diferentes softwares que permitan la evaluación de la red y que sean de código abierto, tales como iPerf3 o hping3. Estas herramientas son muy usadas en este ámbito, ya que permite realizar pruebas de rendimiento enviando diferente tráfico como TCP, UDP e ICMP, además de que ofrece una amplia posibilidad de adecuar parámetros de configuración y adaptarlos a las características del entorno de pruebas, según los requisitos específicos de cada evaluación.

Por otro lado, Wireshark se emplea para capturar y analizar los paquetes que circulan por la red, proporcionando información detallada de cada uno de estos, permitiendo inspeccionar a profundidad los datos que se encuentran en tránsito. En la Tabla 11, se presentan los parámetros específicos que se medirán con cada una de estas herramientas.

Tabla 11*Herramientas de medición.*

Herramienta	Parámetros
iPERf3 o Hping3	• Variación del retardo de los paquetes (IPDV) • Ancho de banda
Wireshark	• Tasa de pérdida de paquetes (IPLR) • Variación de duplicación de paquetes (IPDR)

***Nota.** La disponibilidad de la red no es una medición que se pueda realizar de forma directa, sino que se puede determinar mediante el valor de paquetes perdidos o de la topología.*

C. Configuración de pruebas de rendimiento para VXLAN

Siguiendo lo establecido en el paso anterior, en esta fase se enfoca en llevar a cabo mediciones a lo largo de un periodo específico de tiempo, utilizando paquetes de diferente tamaño. La finalidad de esto es obtener un conjunto de datos representativos que faciliten la comprensión adecuada de los resultados derivados de la evaluación.

La Tabla 12 indica una guía detallada para evaluar el funcionamiento de VXLAN dentro de la red, en donde se especifica principalmente el tiempo de las pruebas, tamaños de los paquetes para cada evaluación y el tipo de tráfico a enviar, lo que permite tener una visión clara de la red bajo diversas condiciones.

Tabla 12*Parámetros de evaluación.*

Parámetro	Tipo de trafico	Tiempo de prueba	Tamaño de paquete (Bytes)
Tasa de duplicación de paquetes (IPDR)	TCP	10 minutos	128, 384, 512, 896, 1024, 1280, 1514
Tasa de pérdida de paquetes (IPLR)	TCP	10 minutos	128, 384, 512, 896, 1024, 1280, 1514
Variación del retardo de paquetes (IPDV)	UDP	10 minutos	N/A
Ancho de banda (AB)	TCP	10 minutos	N/A

Nota. El ancho de banda se va a analizar mediante la utilización de los servicios WEB, FTP Y BDD.

Otros parámetros como es la disponibilidad de la red se definen a partir de la tasa de pérdida de paquetes. Este enfoque hace que se pueda medir indirectamente la disponibilidad de la red en función a la observación de la cantidad de paquetes perdidos durante las pruebas. Además, el uso de recursos se monitoriza a través del sistema operativo de los equipos de red permitiendo observar como los equipos usan la CPU y memoria RAM al momento que se realiza las pruebas para determinar los demás parámetros.

4.4 Pruebas De Rendimiento En El Entorno VXLAN

En esta parte se analiza el rendimiento de VXLAN a través de diversas pruebas. Dado que esta tecnología permite la segmentación de redes a gran escala, se prevé que es importante evaluar su impacto en parámetros importantes como ancho de banda, pérdida de paquetes, variación de

retardo de paquetes y uso de recursos. En el entorno de pruebas desarrollado, consta de tres segmentos de VXLAN, cada uno asignado a un servicio en específico:

- VNI 100: Servidor Web
- VNI 200: Servidor FTP
- VNI 300: Servidor Base de datos

Cada uno de estos servicios tienen características de tráfico y requisitos de rendimiento distinto, lo que permite analizar el comportamiento de VXLAN bajo varias condiciones de red.

Para evaluar el rendimiento, se utiliza la herramienta conocida como iPerf3, en donde se configura los parámetros para las pruebas dependiendo de lo que se desea analizar. Esto facilita la identificación de limitaciones y beneficios de VXLAN y la red.

4.4.1 Tasa De Pérdida De Paquetes (IPLR)

Los resultados obtenidos en cuanto a la tasa de pérdida de paquetes hacen referencia a mediciones efectuadas sobre el tráfico que se encuentra circulando por la capa de transporte, específicamente en la transmisión de información entre cliente y servidor. Este parámetro se obtiene mediante el cálculo tomando en cuenta el porcentaje de paquetes que no lograron llegar al destino correspondiente, comparado con el total de paquetes enviados, tal y como se establece en la recomendación UIT-T Y.1540. Para la obtención de estos resultados en el entorno de pruebas, se emplea el uso del software que permite analizar el tráfico, denominado Wireshark, el cual facilita la identificación y registro de los paquetes perdidos mediante sus contadores de paquetes especializados.

En este caso se está utilizando aplicaciones que utilizan protocolo TCP, este funciona en la capa de transporte y maneja el flujo de información a través de segmentos de datos, lo que permite aplicar estrategias de control de flujo y gestión de congestión. Estos mecanismos son importantes,

ya que aseguran que los segmentos que no fueron entregados correctamente sean enviados otra vez, manteniendo así la integridad y fiabilidad de la transmisión.

La evaluación del parámetro IPLR en TCP, se centra en identificar los segmentos que no pudieron llegar a su destino dentro del tiempo correspondiente. A diferencia de otros protocolos, esta medición se muestra en valor de porcentaje de segmentos perdidos, dado que el proceso captura las pérdidas a nivel de la capa transporte y no a nivel de paquetes de red. Este enfoque asegura que los resultados sean precisos y se puedan alinear a las características que presentan el protocolo TCP, lo cual hace que gestione las pérdidas de datos mediante la retransmisión de la información.

El proceso que se sigue para la medición del parámetro de IPLR esta mostrado en la Figura 48, en donde se muestra que acciones y filtros se utiliza para poder adquirir el valor del parámetro medido. Este proceso consta de diferentes etapas, la primera etapa es la captura del tráfico que está transitando en la red mediante el uso de Wireshark, para luego poder aplicar un filtro que permite enfocarse en paquetes específicos y en este caso en paquetes TCP que se asocian a puertos específicos (como 80=HTTP, 21= FTP y 3306=MySQL), de este modo el filtro que se utiliza en Wireshark es `tcp.port==##`, el cual permite aislar el tráfico de toda la red y observar el tráfico requerido, el # es remplazado por el número de puerto del servicio que se requiere analizar.

Al tener el tráfico filtrado, el proceso que sigue es poder identificar la cantidad de paquetes TCP que se han perdido dentro del tráfico capturado y para esto se hace la utilización de otro filtro de Wireshark conjuntamente con el anterior y queda formado de la siguiente manera `tcp.port==## && tcp.analysis.lost_segment`, el cual permite separar los segmentos de paquetes que no han logrado llegar a su destino, esto es importante porque permite determinar la cantidad de paquetes perdidos que se tiene en la comunicación TCP.

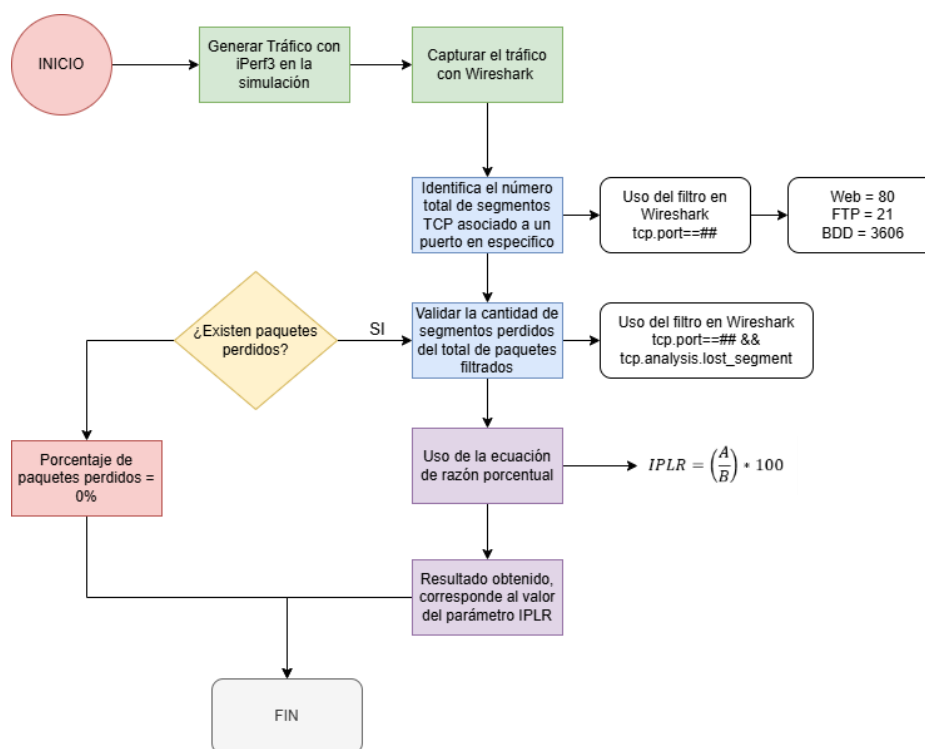
A continuación, al tener filtrado el tráfico el cual se va a analizar, el siguiente paso es verificar el número de paquetes perdidos, si en este caso el valor es mayor a cero, se procede a realizar el cálculo del porcentaje de pérdida de paquetes y para eso se aplica la Ec.1 (de razón porcentual). En esta ecuación matemática se presenta dos variables, la letra A denota la cantidad de paquetes TCP perdidos y la B indica el total de paquetes capturados en el análisis, finalmente se multiplica por 100 para tener el porcentaje de paquetes perdidos.

$$IPLR = \left(\frac{A}{B} \right) * 100$$

Ec. 1

Figura 48

Proceso para la obtención del parámetro IPLR.



A continuación, se presenta el despliegue de dos escenarios destinados a evaluar el rendimiento de la red. El primer escenario es el planteado en la sección 3.4, en donde esta implementado VXLAN y el segundo escenario es la topología similar a la propuesta con la diferencia de que se utiliza otro proceso de encapsulamiento denominado GRE.

Según (Ogudo, 2019), GRE (Generic Routing Encapsulation) es un protocolo de encapsulamiento de capa 3, definido en el RFC 2784. Su función principal es encapsular paquetes IP para permitir la transmisión a través de diversas redes. El proceso de encapsulación es sencillo, ya que únicamente agrega un encabezado de 24 bytes al paquete original. Debido a sus características, el protocolo GRE presenta similitudes con VXLAN, lo que lo convierte en un escenario útil para poder comparar el rendimiento y la operatividad de estos dos dentro de un entorno de pruebas que simula un centro de datos.

El entorno de pruebas cuenta con servidores Web, FTP y Base de datos, los cuales se utilizan para las pruebas correspondientes. En la Tabla 13, se especifican las direcciones IP y el número de puerto asignado a cada servicio y los cuales se utilizan para las pruebas. Con esta configuración, se genera tráfico específico para cada servicio, variando el tamaño de paquete, lo que facilita el análisis comparativo del impacto de las dos tecnologías.

Tabla 13

Dirección IP y puerto de servicios utilizados para las pruebas.

Servidor	Dirección IP	Puerto
WEB	10.10.10.3/24	80/TCP
FTP	10.10.20.3/24	21/TCP
Base de datos	10.10.30.3/24	3306/TCP

Ahora para obtener los paquetes perdidos, se inicializa utilizando la herramienta de iPerf3 dentro de las máquinas virtuales donde están alojados los servicios Web, FTP y de Base de Datos. Primeramente, se emplea el comando `sudo iperf3 -s -p #puerto`, donde el parámetro `#puerto` es el que indica el puerto utilizado por cada servidor. En este caso, se emplea los puertos indicados en la Tabla 13 dependiendo del servicio que se vaya a utilizar.

Los clientes son los que generan el tráfico hacia un servidor iPerf3, en donde se utiliza el siguiente comando `iperf3 -c <dirección IP del servidor> -p #puerto -l <Tamaño del paquete> -t 600 -i 1`. El parámetro `-l` define el tamaño del paquete en bytes, en este caso va a ir cambiando desde 128 hasta llegar a los 1500 bytes, otro parámetro importante es `-t 600` este indica el tiempo que va a durar la prueba y es expresada en segundos y finalmente el parámetro `-i 1` este establece el intervalo de tiempo en un segundo entre cada informe. Este proceso se vuelve repetitivo para todos los servicios, con la única diferencia de que cambia la dirección IP y el puerto.

Una vez realizada la prueba, se utiliza Wireshark para analizar los resultados, permitiendo obtener tanto el total de paquetes transmitidos como la cantidad de segmentos perdidos en la red. En la Figura 49 y 50, se muestra un ejemplo de cómo aplicar filtros en Wireshark, donde se puede visualizar claramente el número total de paquetes transmitidos y la cantidad de segmentos que no se lograron entregar correctamente, este proceso se lo aplica para todos los servicios y pruebas que se realizan.

El total de paquetes enviado para cada tamaño y tecnología refleja el volumen de datos transmitidos, mientras que los paquetes perdidos muestran la cantidad de paquetes que no lograron llegar al destinatario. El IPRL, que se calcula con la Ec.1 con respecto al total de paquetes enviados, lo que permite comparar la fiabilidad de las tecnologías de encapsulación.

Tabla 14

Resultados de IPRL del servicio WEB.

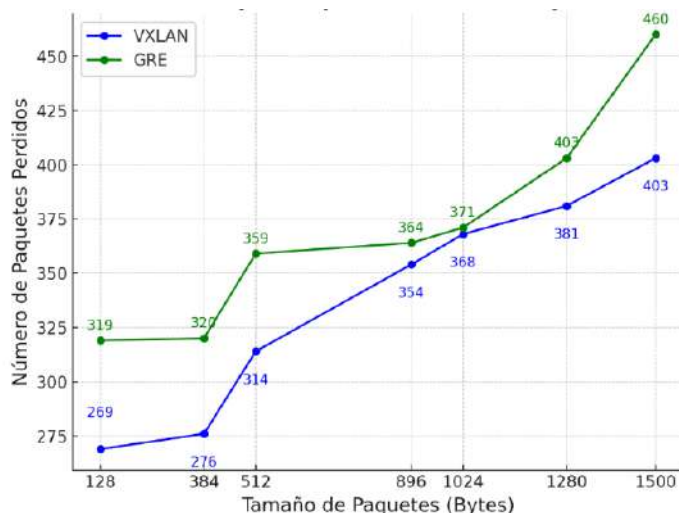
Tamaño del paquete (Bytes)	Tecnología	Total de paquetes	Paquetes Perdidos	IPRL (%)
128	VXLAN	1588436	269	0,017
	GRE		319	0,020
384	VXLAN	1477898	276	0,019
	GRE		320	0,022
512	VXLAN	1856393	314	0,017
	GRE		359	0,019
896	VXLAN	1931861	354	0,018
	GRE		364	0,019
1024	VXLAN	1663130	368	0,022
	GRE		371	0,022
1280	VXLAN	2104093	381	0,018
	GRE		403	0,019
1500	VXLAN	1718404	403	0,023
	GRE		460	0,027

En la Figura 51, se muestra la representación gráfica de la Tabla 14, en donde se revela diferencias importantes en el rendimiento de estas tecnologías frente a distintos tamaños de paquetes. Según los datos obtenidos, se observan que VXLAN presenta una menor tasa de pérdida de paquetes en comparación con GRE.

En términos generales, para paquetes pequeños como 128 bytes, la diferencia que existe entre VXLAN y GRE es significativa. Sin embargo, a medida que el tamaño del paquete va en aumento se puede observar que las pérdidas en paquetes de tamaños de 896 y 1024 bytes la pérdida de paquetes no es tan significativa entre las dos tecnologías. Sin embargo, para paquetes de 1500 bytes, la pérdida de paquetes GRE alcanza un valor de 460, mientras que VXLAN se mantiene en 403, demostrando la capacidad superior de VXLAN manejando el tráfico HTTP en este caso.

Figura 51

Paquetes perdidos haciendo uso del servicio WEB.



En la Figura 52, el análisis de la pérdida de paquetes en el servicio FTP, muestra un comportamiento similar en cuanto al servicio Web, ya que cada vez que aumenta el tamaño de paquete enviado, la cantidad de paquetes perdidos de igual forma aumenta.

Como se observa, VXLAN presenta una tasa de pérdida de paquetes inferior en comparación con GRE. En términos generales, GRE experimenta una pérdida aproximada de 100 paquetes más que VXLAN, otorgando una ventaja importante a VXLAN en cuanto a fiabilidad y desempeño. A pesar de que VXLAN utiliza una encapsulación de paquetes mayor, con una sobrecarga de 50 bytes adicionales, frente a los 24 bytes de encabezado que agrega GRE, esta diferencia en sobrecarga no parece afectar la eficiencia a de VXLAN de manera sustancial.

Este comportamiento puede atribuirse a que VXLAN está diseñado para optimizar redes de alto rendimiento y escalabilidad, como las topologías Spine-Leaf, donde la eficiencia en la transmisión de datos es crítica. En contraste, GRE se ve más impactado por la congestión de la red y posible sobrecarga al hacer uso del servicio FTP, esto resalta que VXLAN brinda un mejor rendimiento y mayor estabilidad frente a la congestión y pérdida de información.

Tabla 15

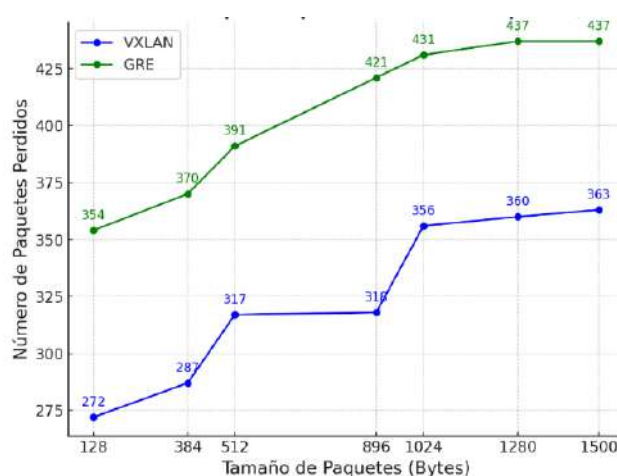
Resultados de IPRL del servicio FTP.

Tamaño del paquete (Bytes)	Tecnología	Total de paquetes	Paquetes Perdidos	IPRL (%)
128	VXLAN	1955506	272	0,014
	GRE		354	0,018
384	VXLAN	1686195	287	0,017
	GRE		370	0,022
512	VXLAN	2190131	317	0,014
	GRE		391	0,018
896	VXLAN	2055924	318	0,015
	GRE		421	0,020
	VXLAN		356	0,021

1024	GRE	1730054	431	0,025
1280	VXLAN	2183337	360	0,016
	GRE		437	0,020
1500	VXLAN	2311990	363	0,016

Figura 52

Paquetes perdidos haciendo uso del servicio FTP.



En la Figura 53, se observa que el comportamiento de VXLAN y GRE varía considerablemente según el tipo de servicio que se está evaluando. A partir de los resultados obtenidos anteriormente de los servicios WEB y FTP, se puede apreciar la diferencia notable de como cada tecnología actúa con los diferentes servicios. En primer lugar, se puede observar que la diferencia de paquetes perdidos de base de datos es mayor comparado a los anteriores servicios, indicando que este tipo de tráfico es más sensible a la pérdida de paquetes. Esto puede ser que el servicio WEB es más tolerante a fluctuaciones en la red. En comparación con el servicio FTP, la diferencia en la cantidad de paquetes perdidos de igual forma es superior a FTP, aunque los dos servicios transfieren grandes cantidades de datos, el servicio de base de datos es aún más sensible a la pérdida de paquetes. Además se puede observar que casi existe una similitud en la gráfica de

estos dos servicios pero sin embargo al hacer uso de VXLAN siguen mejorando y teniendo ventaja, además algo que recalcar de este resultado que con VXLAN en tamaños como 1280 y 1500 bytes se mantuvieron con la misma pérdida de paquetes, mientras que con GRE la brecha seguía creciendo, lo que se puede sacar que VXLAN ofrece una estabilidad y factibilidad mejor en comparación a GRE.

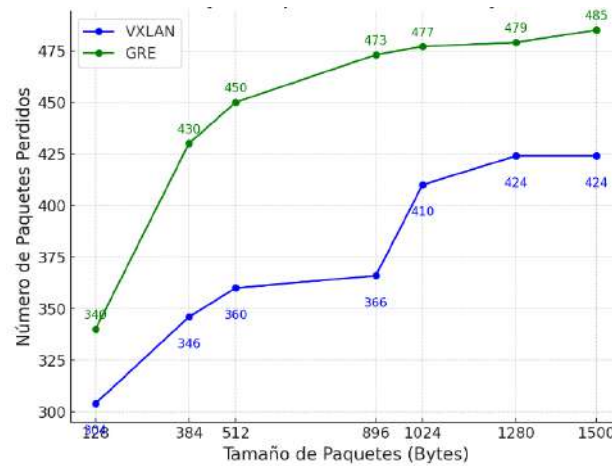
Tabla 16

Resultados de IPRL del servicio de Base de datos.

Tamaño del paquete (Bytes)	Tecnología	Total de paquetes	Paquetes Perdidos	IPRL (%)
128	VXLAN	1812808	304	0,017
	GRE		340	0,019
384	VXLAN	2239232	346	0,015
	GRE		430	0,019
512	VXLAN	2213180	360	0,016
	GRE		450	0,020
896	VXLAN	2367120	366	0,015
	GRE		473	0,020
1024	VXLAN	2598384	410	0,016
	GRE		477	0,018
1280	VXLAN	2990259	424	0,014
	GRE		479	0,016
1500	VXLAN	3199534	424	0,013
	GRE		485	0,015

Figura 53

Paquetes perdidos haciendo uso del servicio de Base de datos.



4.4.2 Disponibilidad De La Red

Según lo establecido por (ITU, 2019) la disponibilidad de la red se basa en función del umbral de la característica IPLR. El servicio IP se considera disponible de extremo a extremo si el valor de IPLR se encuentra por debajo del umbral crítico C_1 , el cual está definido con un valor de 0.20. Esto quiere decir que durante un periodo de observación determinado, la red se considera operativa si el criterio de interrupción $IPRL < C_1$ se cumple. En caso de que el IPRL sea superior al valor umbral, la red se califica como no disponible, dado que muchas aplicaciones críticas dejan de funcionar correctamente cuando la tasa de pérdida de paquetes supera el 20%.

En este contexto, en las Tablas 14, 15 y 16 presentan los valores correspondientes al IPRL, los cuales se calcularon aplicando la Ec.1. Dado que los servicios evaluados utilizan protocolo TCP, se ha previsto analizar un IPRL promedio para cada servicio, con el objetivo de determinar la disponibilidad de la red bajo las condiciones experimentales analizadas anteriormente. Este cálculo permite evaluar si los servicios mantienen un rendimiento adecuado y cumple con los estándares de disponibilidad establecidos.

En las Tablas 17 se resumen los valores promedios de IPLR obtenidos para cada servicio analizado. A partir de estos resultados, se puede decir que VXLAN y GRE, permiten que la red se mantenga disponible en todos los casos evaluados.

Tabla 17

Disponibilidad de la red por servicio en VXLAN y GRE.

Servicio	IPRL VXLAN	IPRL GRE	Disponibilidad
WEB	0,019219346	0,02116358	SI
FTP	0,01623403	0,02031501	SI
BDD	0.015	0,01825828	SI

4.4.3 Tasa De Duplicación De Paquetes (IPDR)

La Tasa de Duplicación de Paquetes (IPDR, acrónimo en inglés) es un parámetro crítico que refleja el número de segmentos TCP duplicados en la red. Este tipo de tráfico señala posibles fallas en la transmisión de datos, tales como congestión en la red o errores en las retransmisiones de los segmentos TCP, lo que da lugar a que los paquetes se dupliquen. La duplicación, a su vez, tiene un impacto negativo en el rendimiento de la red, ya que introduce tráfico redundante que genera retrasos en el procesamiento de información.

El IPDR se calcula mediante la relación que hay entre el número de paquetes retransmitidos y el total de paquetes transmitidos. El proceso de cálculo involucra el recuento de duplicaciones TCP identificadas, comparado con el total de paquetes TCP enviados. En la Figura 54 se describe el procedimiento para obtener este valor. El proceso inicia con el envío de tráfico desde el cliente mediante una petición al servidor, mientras que simultáneamente se inicia la captura del tráfico de red a través del software de Wireshark. Este análisis permite monitorear y registrar toda la

información durante la prueba. Al concluir la transmisión de información, se proceda a identificar la cantidad total de segmentos TCP vinculados a un puerto en específico (Web:80, FTP: 21 y BDD: 3306). En Wireshark se aplica el filtro `tcp.port==##` para contabilizar la cantidad de paquetes totales enviados en la red.

Los paquetes duplicados pueden ser identificados mediante el parámetro de retransmisión en los resultados obtenidos al final de la prueba realizada con iPerf3. En estos resultados, se refleja la cantidad de paquetes retransmitidos, lo cual indica que el emisor no ha recibido correctamente un segmento de datos y, por ende está retransmitiendo la información. La presencia de estos paquetes en la red puede ser un indicativo de que existe problemas en la comunicación, afectando el rendimiento de la red.

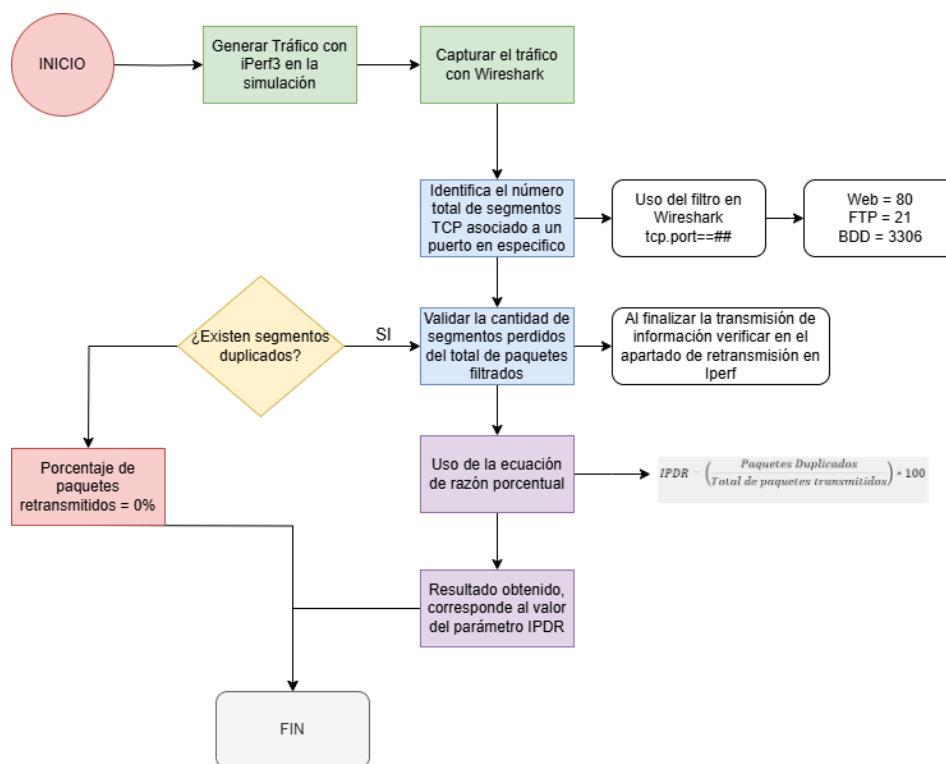
En situaciones donde no exista ningún segmento duplicado, se puede concluir que el porcentaje de duplicación es igual al 0%. Sin embargo, si se detecta la existencia de paquetes duplicados, el cálculo del parámetro IPDR comienza utilizando la Ec.2 (de razón porcentual) lo que permite determinar el porcentaje de segmentos TCP duplicados.

$$IPDR = \left(\frac{\text{Paquetes Duplicados}}{\text{Total de paquetes transmitidos}} \right) * 100 \quad \text{Ec. 2}$$

Finalmente, los valores obtenidos son asociados a los parámetros correspondientes de la Ec.2, y el resultado refleja el valor de IPDR, proporcionando una medición porcentual del grado de duplicación en la transmisión de datos, permitiendo evaluar el comportamiento de la red.

Figura 54

Proceso para la medición de IPDR en tráfico TCP.



Para obtener el IPDR, se utiliza la herramienta de iPerf3 dentro de las máquinas virtuales que se encuentran en el entorno de pruebas. Inicialmente, se configura el servidor iPerf3, utilizando el comando `iperf3 -s -p #puerto`. Los clientes, por su parte, generan el tráfico hacia el servidor iPerf3 utilizando el siguiente comando `iperf3 -c <dirección IP del servidor> -p #puerto -l <tamaño del paquete> -t 600`, tal y como se muestra en la Figura 55. En este comando, el parámetro `-l` define el tamaño del paquete en bytes, el cual variará desde 128 hasta 1500 bytes durante las pruebas. El parámetro `-t 600` especifica la duración de la prueba, que se expresa en segundos (600 segundos, o 10 minutos). Este proceso se repite para cada uno de los servicios (Web, FTP, Base de Datos), variando únicamente la dirección IP y el puerto utilizado.

Figura 55

Comando de iPerf3 para generar tráfico.

```

debian@debian:~$ iperf3 -c 10.10.20.3 -p 21 -l 1024 -t 600
Connecting to host 10.10.20.3, port 21
[ 5] local 10.10.20.10 port 49994 connected to 10.10.20.3 port 21
[ ID] Interval          Transfer    Bitrate      Retr  Cwnd
[ 5]  0.00-1.00    sec   16.6 MBytes  139 Mbits/sec    5   177 KBytes
[ 5]  1.00-2.00    sec   15.9 MBytes  133 Mbits/sec   12   151 KBytes
[ 5]  2.00-3.00    sec   15.4 MBytes  130 Mbits/sec    8   103 KBytes

```

Una vez que se haya completado la prueba de iPerf3, se obtiene un resultado como se muestra en la Figura 56, en donde se obtiene información importante de cómo se estableció la conexión pero el enfoque principal es en el apartado de retransmisión, lo cual nos indica la cantidad de paquetes duplicados que hubo en el proceso, en este caso 6898 paquetes retransmitidos.

Figura 56

Cantidad de paquetes retransmitidos o duplicados.

```

[ 5] 596.00-597.00 sec  13.7 MBytes  115 Mbits/sec    25   133 KBytes
[ 5] 597.00-598.00 sec  10.3 MBytes  86.2 Mbits/sec    5   116 KBytes
[ 5] 598.00-599.00 sec  13.4 MBytes  112 Mbits/sec    9   143 KBytes
[ 5] 599.00-600.00 sec  13.4 MBytes  112 Mbits/sec   27   115 KBytes
-----
[ ID] Interval          Transfer    Bitrate      Retr
[ 5]  0.00-600.00 sec  6.99 GBytes  100 Mbits/sec  6898
[ 5]  0.00-600.01 sec  6.99 GBytes  100 Mbits/sec

```

sender
receiver

Ahora, en la Figura 57 se emplea el filtro en Wireshark para analizar los resultados y obtener el total de paquetes transmitidos, utilizando el filtro mencionado anteriormente, lo que ayuda a obtener una visualización clara del número total de segmentos TCP transmitidos. Este proceso se repite en cada uno de los servicios y pruebas.

Figura 57

Total de paquetes TCP transmitidos.

tcp.port==21

No.	Time	Source	Destination	Protocol	Length	Info
9	3.987603	10.10.20.10	10.10.20.3	TCP	98	51624 → 21 [SYN] Seq=0 Win=64240 Len=0 MSS=1460
10	3.990258	10.10.20.3	10.10.20.10	TCP	98	21 → 51624 [SYN, ACK] Seq=0 Ack=1 Win=65160 Len=0
11	3.992412	10.10.20.10	10.10.20.3	TCP	90	51624 → 21 [ACK] Seq=1 Ack=1 Win=64256 Len=0
12	3.992444	10.10.20.10	10.10.20.3	FTP	127	Request: 77c2yri3qoxf5iuvb6pcsjzjemrm3zep1g5x'
13	3.994378	10.10.20.3	10.10.20.10	TCP	90	21 → 51624 [ACK] Seq=1 Ack=38 Win=65152 Len=0
14	3.994672	10.10.20.3	10.10.20.10	FTP	91	Response: \t
15	4.007740	10.10.20.10	10.10.20.3	TCP	90	51624 → 21 [ACK] Seq=38 Ack=2 Win=64256 Len=0

Frame 9: 98 bytes on wire (784 bits), 98 bytes captured (784 bits) on interface -, id 0

Ethernet II, Src: 0c:b2:a6:d2:00:01 (0c:b2:a6:d2:00:01), Dst: 0c:dd:ed:e5:00:01 (0c:dd:ed:e5:00:01)

Internet Protocol Version 4, Src: 192.168.0.6, Dst: 192.168.0.3

Generic Routing Encapsulation (IP)

Internet Protocol Version 4, Src: 10.10.20.10, Dst: 10.10.20.3

Transmission Control Protocol, Src Port: 51624, Dst Port: 21, Seq: 0, Len: 0

1024.pcapng

Paquetes: 2173643 - Displayed: 2172640 (100.0%)

En las Tablas 18, 19 y 20 se muestran los resultados que se obtienen luego de realizar las pruebas correspondientes que permitieron medir el IPDR. Cada prueba fue repetida tres veces para poder asegurar la fiabilidad y precisión de la información para así tener un valor promedio representativo. Los resultados de las pruebas han variado debido a factores como fluctuaciones en la medición de paquetes duplicados, causados por congestión momentánea de la red, procesos en segundo plano de los equipos de red como el servidor en el cual se esta alojado y también por la limitación de recursos. Al promediar los resultados, se obtiene un valor estable y preciso que muestra como es el comportamiento de cada tecnología bajo las condiciones de pruebas.

Además, el promedio del total de paquetes transmitidos es importante para realizar una comparativa entre las dos tecnologías que es VXLAN y GRE, cada una de estas tiene características diferentes que influyen en la cantidad de paquetes enviados en cada prueba. De este modo, promediando ambos parámetros (paquetes totales transmitidos y paquetes duplicados) se tiene una comparación clara y justa permitiendo realizar el análisis correctamente en términos de eficiencia y fiabilidad dentro del entorno de pruebas.

Tabla 18

Resultados del servicio WEB para IPDR.

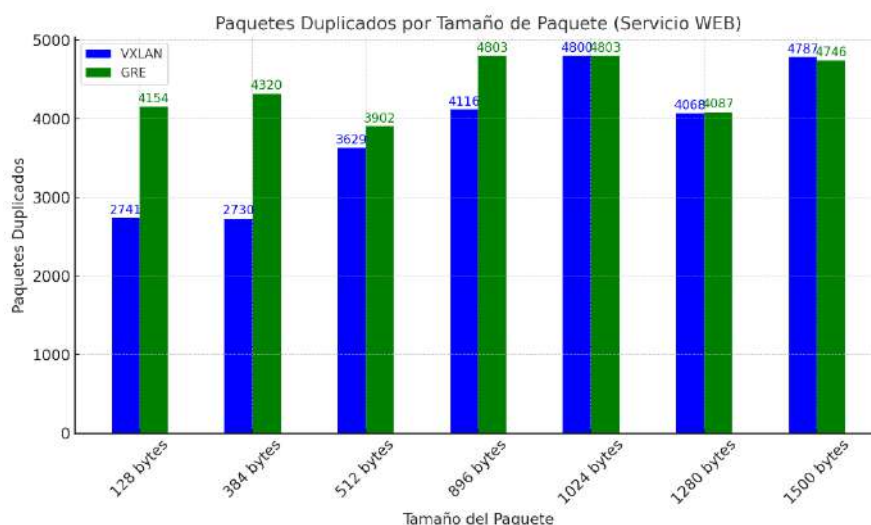
Tamaño del paquete (Bytes)	Tecnología	Total de paquetes	Paquetes Duplicados	IPDR (%)
128	VXLAN	1749622	2741	0.157
	GRE		4154	0.237
384	VXLAN	1757757	2730	0.155
	GRE		4320	0.246
512	VXLAN	1795234	3629	0.202
	GRE		3902	0.217
896	VXLAN	1902923	4116	0.216
	GRE		4803	0.252
1024	VXLAN	1945746	4800	0.247
	GRE		4803	0.247
1280	VXLAN	2054530	4068	0.198
	GRE		4087	0.199
1500	VXLAN	2254179	4787	0.212
	GRE		4746	0.211

En la Figura 58, se muestra los resultados de la Tabla 19 en forma gráfica para una mejor comprensión, en donde se evidencia una tendencia general en la que VXLAN presenta una menor cantidad de paquetes duplicados en comparación con GRE. A medida que aumenta el tamaño del paquete, ambos protocolos indican un aumento en los paquetes duplicados, lo cual es una tendencia muy común debido a que existe una mayor carga de información en la red cuando el tamaño del paquete aumenta. Sin embargo, la diferencia en los valores de duplicación es evidente en todos los

tamaños del paquete. Esto se debe a que VXLAN está diseñado para entornos como centro de datos, optimiza la transmisión de datos de una mejor manera, obteniendo una mayor eficiencia. Mientras que GRE, al ser un protocolo de red utilizado para redes pequeñas y en este caso el entorno de pruebas simula un centro de datos, presenta una mayor tendencia a la duplicación debido a su simplicidad en la gestión de tráfico. Posiblemente en redes pequeñas pueda ofrecer un mejor rendimiento GRE que VXLAN pero en este caso fue diferente.

Figura 58

Paquetes duplicados entre VXLAN y GRE al usar el servicio WEB.



En la Tabla 19, se presenta los resultados obtenidos durante las pruebas de paquetes duplicados en el servicio FTP, cada fila representa un tamaño de paquetes específico y muestra los valores correspondientes a los paquetes duplicados para cada tecnología. A medida que aumenta el tamaño del paquete tanto VXLAN y GRE tiene la misma tendencia de un incremento de paquetes duplicados como el ejemplo anterior.

Tabla 19*Resultados del servicio FTP para IPDR.*

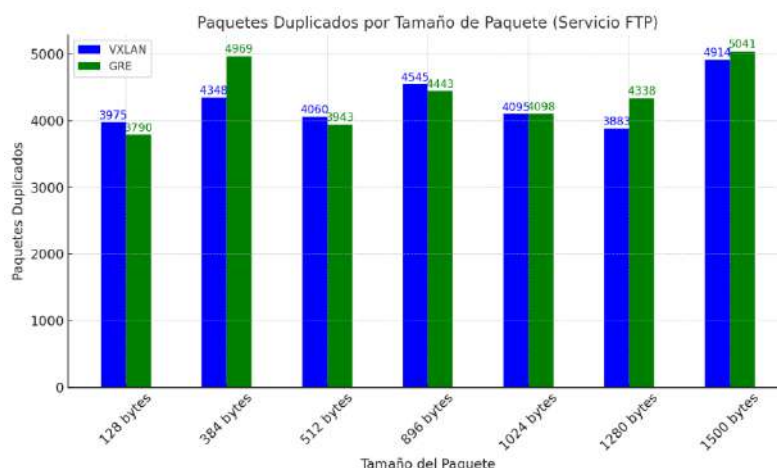
Tamaño del paquete (Bytes)	Tecnología	Total de paquetes	Paquetes Duplicados	IPDR (%)
128	VXLAN	1411433	3975	0.282
	GRE		3790	0.269
384	VXLAN	1990718	4348	0.218
	GRE		4969	0.25
512	VXLAN	2085414	4060	0.195
	GRE		3943	0.189
896	VXLAN	2180736	4545	0.208
	GRE		4443	0.203
1024	VXLAN	2211444	4095	0.185
	GRE		4098	0.185
1280	VXLAN	2949244	3883	0.132
	GRE		4338	0.147
1500	VXLAN	3054542	4914	0.161
	GRE		5041	0.165

En la Figura 59, se representa gráficamente los resultados obtenidos. En el caso de FTP, comúnmente se transmite información de gran tamaño y la tasa de duplicación puede afectar directamente al rendimiento del servicio. VXLAN, al manejar menos duplicaciones, puede presentar una eficiencia mayor comparado a GRE. Sin embargo, en estos resultados se observa que VXLAN no le saca una ventaja a GRE, los resultados son casi similares, como se evidencia en la columna de IPDR en los primeros tamaños de paquetes la diferencia es mínima y se mantiene a lo largo de estos tamaños, hasta llegar al tamaño de paquete de 1024 bytes, punto en el que

VXLAN y GRE tienen el mismo porcentaje de IPDR de 0.185%, esto puede deberse a problemas de recursos en los equipos o posibles eventos que este pasando en la red. No obstante, en entornos donde la transmisión de datos es crítica, como en FTP, VXLAN conserva una mínima ventaja sobre GRE.

Figura 59

Paquetes duplicados en tráfico FTP.



En la Tabla 20, se tiene los resultados que se obtuvieron en la prueba de paquetes duplicados para el servicio de Base de Datos. Al comparar estos resultados con los servicios Web y FTP, se observan algunas diferencias en el comportamiento de VXLAN y GRE.

En este sentido, para tamaños de paquetes pequeños (128 bytes a 896 bytes), GRE demuestra ser más eficiente que VXLAN, ya que nos muestra una cantidad de paquetes duplicados menor y un IPDR más bajo, especialmente en el tamaño de 512 bytes y 896 bytes, donde GRE tiene valores de duplicaciones considerablemente bajos en comparación con VXLAN. Sin embargo, en tamaños de paquetes grandes desde 1024 a 1500 bytes, VXLAN supera a GRE con un valor porcentual de IPDR de 0.145%, frente al 0.199% de GRE. Esta diferencia se vuelve significativa, lo que indica que VXLAN puede manejar paquetes más grandes de manera más eficiente, reduciendo la duplicación y optimizando el rendimiento de la red.

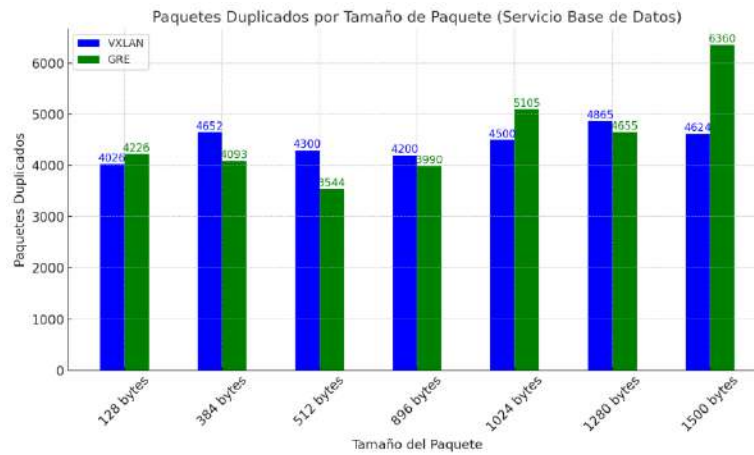
Tabla 20*Resultados del servicio de Base de Datos para IPDR.*

Tamaño del paquete (Bytes)	Tecnología	Total de paquetes	Paquetes Duplicados	IPDR (%)
128	VXLAN	1848772	4026	0.218
	GRE		4226	0.229
384	VXLAN	2487574	4652	0.187
	GRE		4093	0.165
512	VXLAN	2714463	4300	0.158
	GRE		3544	0.131
896	VXLAN	2766141	4200	0.152
	GRE		3990	0.144
1024	VXLAN	2990259	4500	0.150
	GRE		5105	0.170
1280	VXLAN	3099534	4865	0.157
	GRE		4655	0.150
1500	VXLAN	3199534	4624	0.145
	GRE		6360	0.199

En la Figura 60, se ilustran estos resultados de manera gráfica, lo que facilita la comparación visual de cómo cada tecnología se comporta según el tamaño del paquete. Mientras que GRE es más eficiente en paquetes pequeños, VXLAN se destaca en el manejo de paquetes más grandes, lo que hace que VXLAN sea la opción más adecuada para redes de alto rendimiento como las que se utilizan en Data Centers y entornos de virtualización, donde la gestión eficiente de grandes volúmenes de datos es esencial.

Figura 60

Paquetes duplicados en el servicio de Base de Datos.



4.4.4 Variación Del Retardo De Paquetes (IPDV)

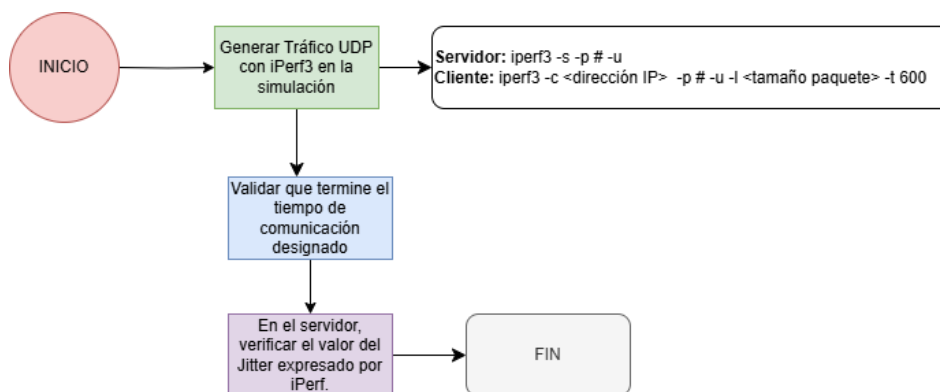
La variabilidad en el retardo de paquetes más conocida en el mundo de las redes como Jitter, este es un parámetro importante en redes de comunicación, permitiendo analizar las fluctuaciones temporales en la entrega de datos durante la transmisión de información en una red. Este parámetro es muy importante en aplicaciones que requieren una alta precisión como son servicios IPTV, VoIP, juegos en línea y videos en tiempo real.

La variación del retardo de paquetes (IPDV) permite tener una medición estadística de la variabilidad en el retardo de paquetes, este valor hace referencia al Jitter máximo que se registra en una transmisión de información dentro de una red. Para la medición de IPDV se lleva a cabo simulando tráfico UDP mediante la herramienta iPerf. El proceso para la obtención del Jitter se muestra en la Figura 61, donde se ilustra la metodología seguida para obtener el resultado. El procedimiento inicia con la generación de tráfico UDP desde un cliente hacia un servidor. Posteriormente, se espera a que la comunicación entre ambos termine. Finalmente, se verifica el valor máximo del Jitter registrado al concluir la simulación, lo cual proporciona el retardo de

transmisión de los paquetes UDP en milisegundos, reflejando la variación temporal en el retardo de la información transmitida.

Figura 61

Proceso para obtener el parámetro de IPDV.



Para evaluar el rendimiento de la red, en términos de variabilidad en el retardo de paquetes, se realiza pruebas de Jitter utilizando iPerf3. Las pruebas se llevan a cabo configurando primero el servidor con el comando `iperf3 -s -p #puerto -u`, donde `-p #` corresponde al valor del puerto, siendo el 5060, un puerto estándar que permite simular el servicio VoIP. El indicador `-u` es para poder recibir paquetes que utilizan protocolo UDP, se realiza eso ya que es el protocolo más adecuado para medir el Jitter.

Por otro lado, en el cliente, el comando para generar tráfico UDP fue: `iperf3 -c <dirección ip del servidor> -p #puerto -l <Tamaño del paquete> -t <Tiempo de pruebas> -u`. Se emplea el mismo valor del puerto que se asigna al servidor, el tamaño del paquete varió entre 128 y 1500 bytes, y el tiempo de las pruebas se establece en 600 segundos, lo que permite obtener una medición adecuada del Jitter en condiciones controladas. En la Figura 62 se muestra una captura de pantalla que indica como se obtiene el valor del Jitter al finalizar el envío de los paquetes.

Figura 62

Obtención del valor del Jitter con iPerf3.

[ID]	Interval	Transfer	Bitrate	Jitter
[5]	0.00-600.00 sec	75.0 MBytes	1.05 Mbits/sec	0.000 ms
[5]	0.00-600.01 sec	74.6 MBytes	1.04 Mbits/sec	0.389 ms

Los resultados que se muestran en la Tabla 21, indican claramente las diferencias en tiempo del retardo de paquete entre GRE y VXLAN. A lo largo de las pruebas, VXLAN demuestra un rendimiento correcto en el tiempo que se realiza la prueba dentro del entorno, manteniendo un Jitter bajo en comparación con GRE para cada tamaño de paquete probado. Esto muestra que VXLAN es un protocolo más eficiente en la entrega consistente de los paquetes, lo cual hace que sea un protocolo adecuado para el uso en centro de datos que tenga aplicaciones en tiempo real.

Tabla 21

Resultados obtenidos de Jitter de VXLAN y GRE.

Tamaño del paquete (bytes)	VxLAN (ms)	GRE (ms)
128	0.588	1.162
384	0.386	0.740
512	0.760	1.895
896	1.560	2.118
1024	1.780	2.394
1280	2.011	2.556
1500	2.352	2.728

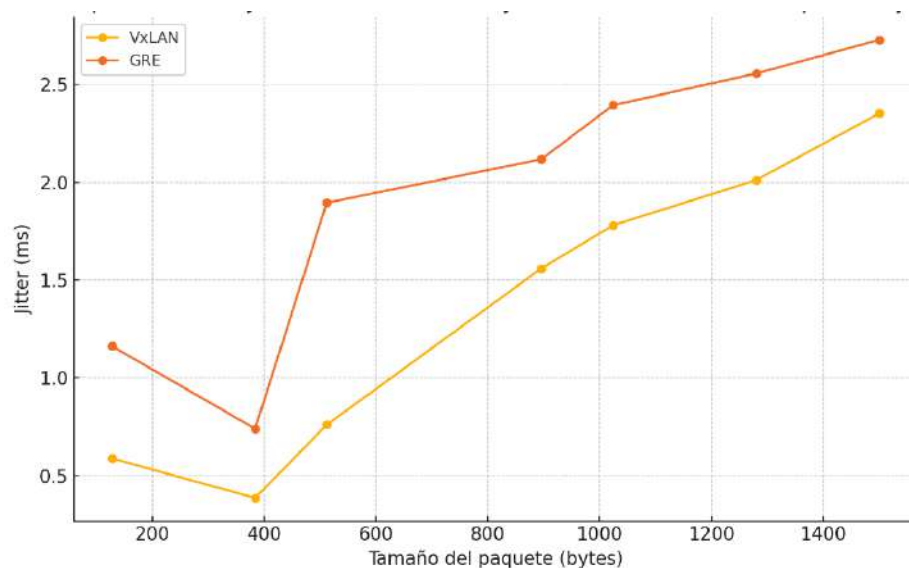
En la Figura 63, se presenta la comparación de resultados entre los dos protocolos de forma gráfica. En donde se puede decir que el comportamiento del protocolo VXLAN, se puede observar que el Jitter aumenta a medida que crece el tamaño del paquete. Este comportamiento es esperado debido a que, en redes con mayor volumen de datos y mayor tamaño de paquetes, la sobrecarga asociada con el encapsulamiento del paquete genera un incremento en la latencia. A partir de un tamaño de paquete de 896 bytes, el Jitter en VXLAN experimenta un aumento considerable, alcanzando 2.352 ms a 1500 bytes. Este resultado refleja el impacto de los protocolos de encapsulamiento en el desempeño de la red, donde la sobrecarga se incrementa con el tamaño del paquete.

Por otro lado, el protocolo GRE muestra tener un comportamiento similar al de VXLAN en cuanto al aumento del Jitter con el incremento del tamaño del paquete. Sin embargo, el Jitter que tiene GRE es un poco más alto que en VXLAN. Por ejemplo, en el caso de un paquete de 128 bytes, el Jitter en GRE es de 1.162 ms, mientras que en VXLAN es de 0.588 ms. Este tiene un patrón que se mantiene a lo largo de todos los tamaños de paquete, con el Jitter en GRE alcanzando los 2.728 ms a 1500 bytes, lo que refleja una mayor sobrecarga de encapsulamiento en este protocolo en comparación con VXLAN.

Estos resultados demuestran que el rendimiento máximo de retardo no supera los 3 milisegundos, por lo que el desempeño del retardo en estos dos protocolos de túnel puestos a prueba dentro de un entorno virtualizado es confiable para el uso, ya que cumple con el estándar ITU G.1010 el cual se relaciona directamente con el acceso a servidores, en donde establece un retardo aceptable siempre y cuando sea menor a los 3 milisegundos.

Figura 63

Comparación del Jitter entre VXLAN y GRE.



4.4.5 Uso De Recursos

En sistemas operativos de red basados en Linux, como en este caso Cumulus Linux, el monitoreo en tiempo real del uso de recursos es esencial para comprender el rendimiento de las tecnologías de red, como lo es VXLAN y GRE. El uso del comando `htop`, permite visualizar el consumo del vCPU, y en este estudio se observa que el uso de la vCPU alcanzó su máximo valor durante la transmisión y recepción de información al momento de realizar pruebas de tráfico de red. En las Figuras 64 y 65 se muestra como el vCPU llega al máximo en el switch virtualizado donde tiene un uso de 81% en VXLAN y 69% en GRE, como también el de la máquina virtual que es GNS3 VM llegando a un valor de 92% en VXLAN y 88% en GRE respectivamente.

Las simulaciones realizadas con la herramienta de iPerf, mostraron variabilidad de resultados. Esto se debe, en parte, al hecho de que los recursos del sistema, como la CPU y la memoria, no solo se miran afectados por el tráfico de la red en sí, sino también por otros procesos que operan en segundo plano en la maquina física, máquina virtual y en los equipos virtualizados como son los switches y servidores que se utilizan en la simulación. Dado también que iPerf genera una carga significativa en la red y en la CPU, cualquier alteración en los recursos puede modificar el comportamiento de la herramienta y generar variaciones en las métricas obtenidas.

En las Tablas 22 y 23, se registra los valores de pruebas realizadas en las simulaciones para comprobar el comportamiento de la red de los resultados mencionados anteriormente. Se observa variaciones en todos los resultados obtenidos, y su análisis es importante para comprender el rendimiento de la red bajo ciertas condiciones y con recursos limitados. Estas fluctuaciones en los resultados se deben a ciertas razones como es la limitación de recursos en la red, ya que los equipos dependen de una gran medida de recursos disponibles del sistema, como CPU y memoria RAM. En este caso son limitados y sobrecargan los recursos lo que conlleva a que los resultados cambien tanto en VXLAN como en GRE. También se debe a la variabilidad inherente de la herramienta de iPerf por sus procesos en segundo plano, además es importante destacar que los resultados obtenidos, aunque muestran ciertas fluctuaciones, son consistentes con lo esperado en redes con recursos limitados y bajo condiciones de tráfico fluctuante.

Tabla 22*Base de datos de resultados de pruebas de VXLAN.*

VXLAN												
	Paquetes totales				Paquetes perdidos				Paquetes duplicados			
<i>Tamaño del paquete (bytes)</i>	<i>Prueba 1</i>	<i>Prueba 2</i>	<i>Prueba 3</i>	<i>Promedio</i>	<i>Prueba 1</i>	<i>Prueba 2</i>	<i>Prueba 3</i>	<i>Promedio</i>	<i>Prueba 1</i>	<i>Prueba 2</i>	<i>Prueba 3</i>	<i>Promedio</i>
128	1800453	2273836	1793229	1955506	245	278	293	272	3511	4211	4203	3975
384	1872300	1698465	1489970	1686195	250	280	331	287	4459	4701	3884	4348
512	2023793	2405128	2145312	2190131	299	330	321	317	3607	4257	4317	4060
896	2123378	1807551	2230893	2055924	282	317	355	318	3822	5005	4808	4545
1024	1789125	2287145	1115892	1730054	340	368	360	356	3275	4654	4357	4095
1280	1768543	2349586	2435922	2183337	370	330	380	360	2821	4065	4764	3883
1500	1965373	1987812	2988785	2311990	345	355	389	363	4647	3701	6394	4914

Tabla 23*Base de datos de resultados de pruebas de GRE.*

GRE												
	Paquetes totales				Paquetes perdidos				Paquetes duplicados			
<i>Tamaño del paquete (bytes)</i>	<i>Prueba 1</i>	<i>Prueba 2</i>	<i>Prueba 3</i>	<i>Promedio</i>	<i>Prueba 1</i>	<i>Prueba 2</i>	<i>Prueba 3</i>	<i>Promedio</i>	<i>Prueba 1</i>	<i>Prueba 2</i>	<i>Prueba 3</i>	<i>Promedio</i>
128	2198345	1647521	2023371	1955506	312	238	412	354	3350	4300	3820	3790
384	1860712	1467098	1720776	1686195	274	395	338	370	5700	3850	4357	4969
512	2405752	1910216	2254425	2190131	379	252	405	391	4420	3300	4109	3943
896	1779028	2053802	2334943	2055924	238	398	457	421	3650	4930	4749	4443
1024	1178473	2027345	2004345	1730054	322	240	550	431	4300	3700	4095	4098
1280	1871911	2289543	2388560	2183337	300	388	420	437	4500	4900	3600	4338
1500	2091780	2357218	2286973	2311990	275	400	436	437	3880	6200	5043	5041

4.4.6 Medición Del Ancho De Banda

La medición del ancho de banda es una de las tareas importantes que permite evaluar el comportamiento de la red de comunicaciones. Este parámetro hace referencia a la cantidad de datos que se transmiten a través de una red en un período de tiempo determinado, generalmente se lo mide en Megabits por segundo (Mbps/sec) o en otros casos en Gigabits por segundo (Gbps/sec).

En este caso para obtener este parámetro se hace uso de iPerf, las pruebas se llevan a cabo configurando primeramente el comando `iperf -s -p #puerto`, donde `-p` define el valor del puerto que se utiliza para el servicio WEB, FTP y Base de Datos. Por otro lado, en el cliente, el comando el cual permite generar tráfico es `iperf -c <dirección ip del servidor> -p #puerto -t 600 -i 10`. Se emplea el mismo valor del puerto que se asigna al servidor, el tiempo de las pruebas se establece en 600 segundos y en intervalo de 10 segundos. Al momento que se realiza la prueba se mira como la Figura 66, en donde tenemos el intervalo del tiempo en el cual se está enviando la información la transferencia de datos y el ancho de banda expresado en Mbps/sec.

Figura 66

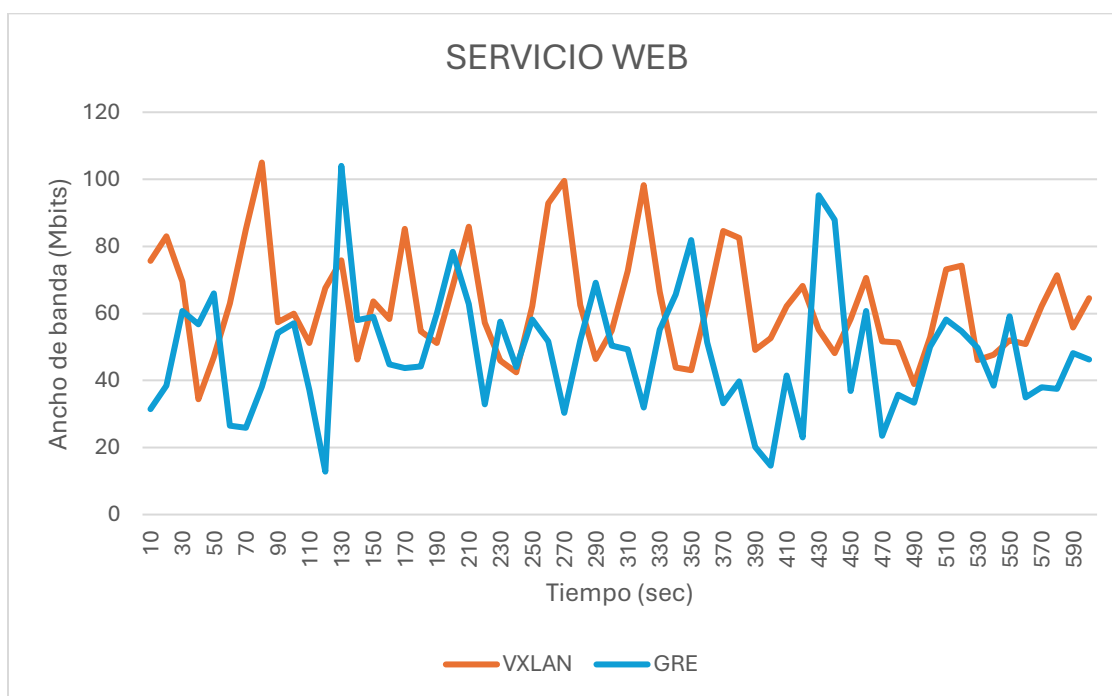
Prueba de ancho de banda con iPerf.

ID]	Interval	Transfer	Bandwidth
[1]	0.0000-10.0000 sec	37.5 MBytes	31.5 Mbps/sec
[1]	10.0000-20.0000 sec	45.8 MBytes	38.4 Mbps/sec
[1]	20.0000-30.0000 sec	72.4 MBytes	60.7 Mbps/sec
[1]	30.0000-40.0000 sec	67.8 MBytes	56.8 Mbps/sec
[1]	40.0000-50.0000 sec	78.5 MBytes	65.9 Mbps/sec
[1]	50.0000-60.0000 sec	31.6 MBytes	26.5 Mbps/sec
[1]	60.0000-70.0000 sec	30.9 MBytes	25.9 Mbps/sec
[1]	70.0000-80.0000 sec	45.5 MBytes	38.2 Mbps/sec
[1]	80.0000-90.0000 sec	64.6 MBytes	54.2 Mbps/sec

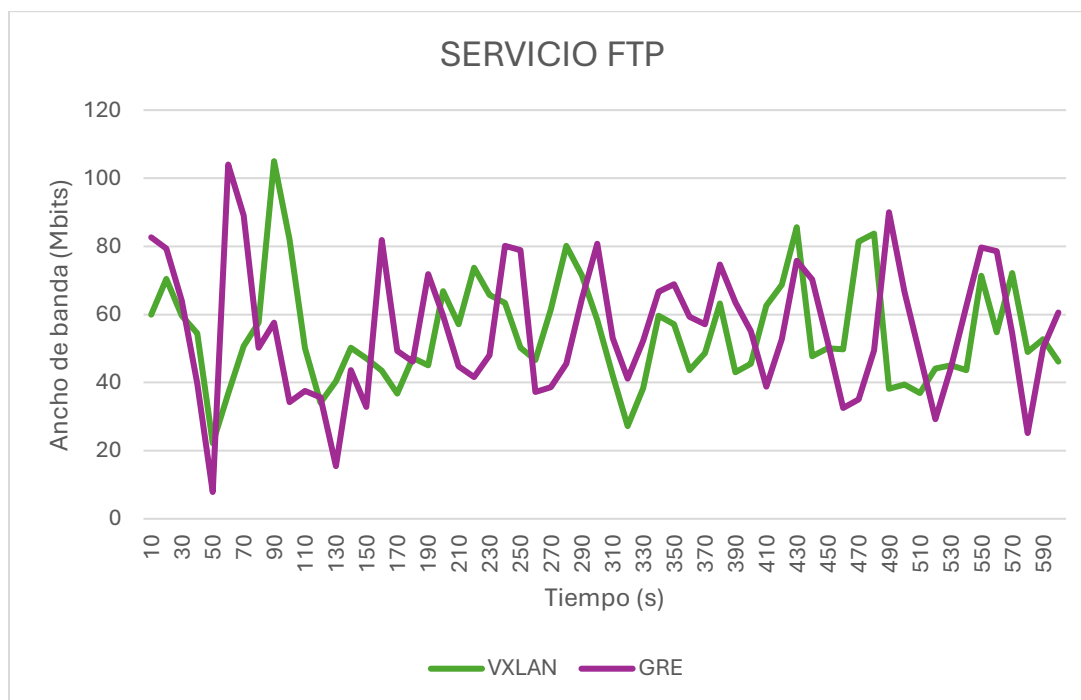
En la Figura 67, se muestra el ancho de banda que utiliza el servicio WEB tanto en VXLAN como en GRE. En donde, muestra un uso de ancho de banda inestable en VXLAN con picos más altos en comparación con GRE. En términos de datos, para VXLAN en los intervalos de 10 a 80, se observa un incremento hasta llegar al pico más alto llegando aproximadamente a más de 100 Mbits/sec, mientras que GRE tiene un valor máximo aproximado al de VXLAN a los 130 segundos, dentro de esta gráfica también indica que GRE utiliza un valor menor comparado a VXLAN para este caso.

Figura 67

Ancho de banda servicio WEB.



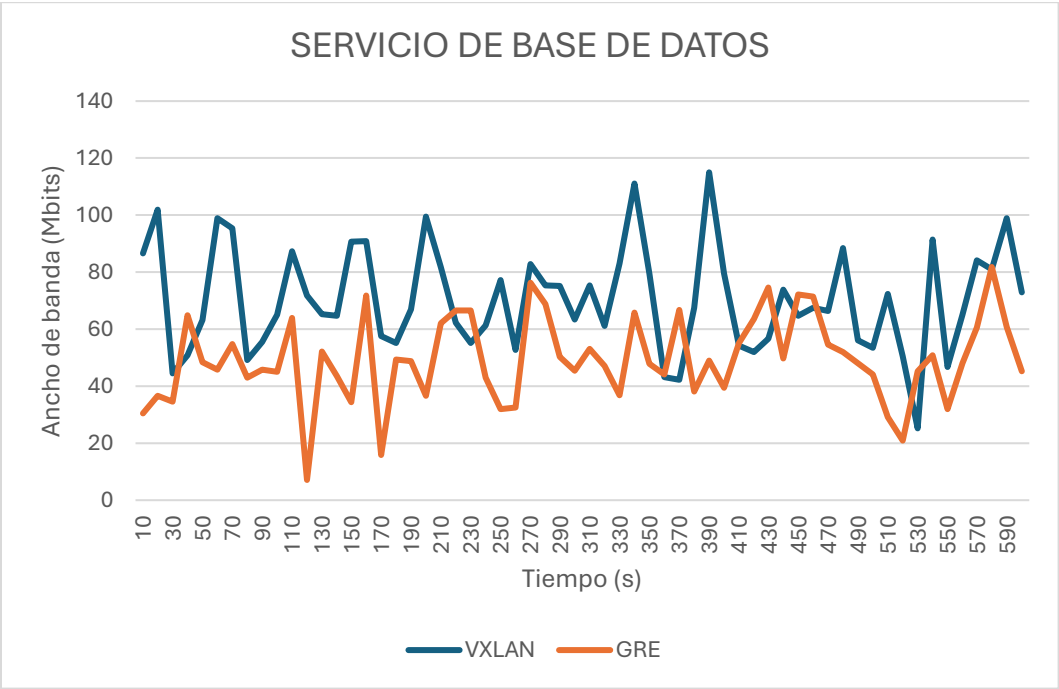
En la Figura 68, muestra el ancho de banda que se obtiene del servicio FTP, los valores de ancho de banda para VXLAN muestran una mayor variabilidad, aunque los picos de VXLAN no son tan elevados como en el Servicio Web. Sin embargo, en este caso el ancho de banda tiene una tendencia igualitaria en VXLAN como en GRE.

Figura 68*Ancho de banda Servicio FTP.*

En la Figura 69, se indica el ancho de banda utilizado por el servicio de Base de datos, en donde muestra que VXLAN tiene algunas fluctuaciones en el uso del ancho de banda, aunque generalmente este se mantiene en un rango alto. Se tiene que hay picos que logran superar los 100 Mbits/sec en algunos puntos de la transmisión y en este caso los valores de VXLAN son más estables en comparación con GRE, manteniéndose por encima de los 60 Mbits/sec en la mayoría de los puntos. Mientras que GRE presenta más variabilidad en su uso de ancho de banda, lo cual podría ser indicativo de una mayor ineficiencia o sobrecarga en el manejo del tráfico comparado con VXLAN.

Figura 69

Ancho de banda Servicio de Base de Datos.



5 CAPITULO V – CONCLUSIONES Y RECOMENDACIONES

5.1 Conclusiones

El desarrollo de este trabajo confirma que la tecnología VXLAN representa una evolución importante en la virtualización de redes, especialmente para centro de datos y entornos en la nube. Su capacidad de ofrecer una segmentación lógica más amplia y escalable, al utilizar encapsulamiento de tramas ethernet sobre redes IP mediante UDP, va más allá de las limitaciones que tienen las VLAN tradicionales, especialmente en redes de gran tamaño.

En arquitecturas de red Spine-Leaf, VXLAN demuestra ser la opción más eficiente, destacándose por su estabilidad y capacidad que tiene para manejar el tráfico encapsulado de manera correcta. Esto hace que se tenga un uso óptimo de los recursos de la red, además de que permite tener una reducción notable en lo que es pérdida y duplicación de paquetes en comparación con otras tecnologías como es GRE. Además, VXLAN ofrece tener un rendimiento superior en cuanto a término de Jitter, este es un factor muy importante dentro de lo que es aplicaciones que dependen de tiempos de respuesta rápido, como son aplicaciones en tiempo real como transmisiones en vivo o VoIP.

VXLAN, por su eficiencia en el manejo de tráfico, su estabilidad, control de jitter y optimización del uso de recursos, se posiciona como una solución ideal para redes de gran escala, garantizando un alto rendimiento y una escalabilidad perfecta para los centros de datos y entornos de nube modernos.

5.2 Recomendaciones

- Se sugiere que las organizaciones que se basan en redes como tecnología tradicional investiguen el impacto que tiene actualizar la infraestructura que implemente VXLAN. Esto debería incluir un análisis de como la transmisión de los datos influye en la continuidad del servicio, sincronización de la migración y principalmente en la gestión eficiente que brinda esta infraestructura actualizada.
- Aunque VXLAN puede demandar el uso de mayor cantidad de recursos que otras tecnologías de virtualización como GRE, se recomienda invertir en equipos de alto rendimiento que permitan manejar la carga adicional de procesamiento y más que todo que sean compatibles con el protocolo VXLAN. Esta inversión asegurara una mayor eficiencia operativa y mejorara el rendimiento del centro de datos.
- Considerar a VXLAN como una inversión a largo plazo que no solo cubra las necesidades actuales, sino que también permita un crecimiento sostenible de la infraestructura de red.
- Se recomienda realizar pruebas adicionales sobre la escalabilidad de VXLAN en redes extremadamente grandes, con miles de nodos conectados. Esto permitirá entender mejor cómo se comporta la tecnología a medida que la red crece, identificando posibles cuellos de botella y proporcionando recomendaciones sobre cómo mitigar estos problemas.

6 BIBLIOGRAFÍA

- Arbeláez Cardona, L. G. (2013). *Diagnóstico de la red de comunicaciones de la universidad católica de Pereira*. <http://repositorio.ucp.edu.co/handle/10785/1736>
- Arista. (2024). *AWE-7200R Routers*.
- Arregoces, Mauricio., & Portolani, Maurizio. (2004). Data Center Security Overview . *Data Center Fundamentals*, 159–202.
https://books.google.com/books/about/Data_Center_Fundamentals.html?hl=es&id=DRlryrLoxKkC
- Aruba Networking. (n.d.). *What is Spine-Leaf Architecture and How Does It Work?* | HPE Aruba Networking. Retrieved January 2, 2024, from <https://www.arubanetworks.com/faq/what-is-spine-leaf-architecture/>
- Capella Hernández, nombre, & Vicente, J. (2012). *Características y configuración básica de VLANs*. <https://riunet.upv.es/handle/10251/16310>
- Chavarro, D., Vélez, I., Tovar, G., Montenegro, I., Hernández, A., & Olaya, A. (2017). *Los Objetivos de Desarrollo Sostenible en Colombia y el aporte de la ciencia, la tecnología y la innovación*.
- Chen, K., Singla, A., Singh, A., Ramachandran, K., Xu, L., Zhang, Y., Wen, X., & Chen, Y. (2014). OSA: An optical switching architecture for data center networks with unprecedented flexibility. *IEEE/ACM Transactions on Networking*, 22(2), 498–511.
<https://doi.org/10.1109/TNET.2013.2253120>

- Chkirbene, Z., Hadjidj, R., Fougou, S., & Hamila, R. (2020). Lascada: A novel scalable topology for data center network. *IEEE/ACM Transactions on Networking*, 28(5), 2051–2064. <https://doi.org/10.1109/TNET.2020.3008512>
- Cisco. (2020, August 16). *Cisco Nexus 9300-EX Series Switches Data Sheet - Cisco*. <https://www.cisco.com/c/en/us/products/collateral/switches/nexus-9000-series-switches/datasheet-c78-742283.html>
- Darquea Endara, P. E. (2017). *Rediseño de la red del centro de datos de la empresa Andean Trade S.A.* <http://dspace.udla.edu.ec/handle/33000/8091>
- Davoli, R., & Goldweber, M. (2015). VXVDE: A switch-free VXLAN replacement. *2015 IEEE Globecom Workshops, GC Wkshps 2015 - Proceedings*. <https://doi.org/10.1109/GLOCOMW.2015.7414109>
- Debian. (n.d.). 3.4. *Cumplir los requisitos mínimos de hardware*. Retrieved April 20, 2024, from <https://www.debian.org/releases/buster/i386/ch03s04.es.html>
- Dutt, D., Networks, C., Duda, K., Agarwal, A. P., Kreeger, B. L., Sridhar, C. T., Bursell, V. M., & Wright, I. C. (2014). *Independent Submission M. Mahalingam Request for Comments: 7348 Storvisor Category: Informational Virtual eXtensible Local Area Network (VXLAN): A Framework for Overlaying Virtualized Layer 2 Networks over Layer 3 Networks*. <http://www.rfc-editor.org/info/rfc7348>.
- Franco, H. (2020). *Overlay and Underlay – Héctor Franco*. <https://hectorfranco.work/overlay-networks/>
- Gartner Peer Insights. (2024). *Best Data Center and Cloud Networking Reviews 2024 | Gartner Peer Insights*. <https://www.gartner.com/reviews/market/data-center-and-cloud-networking>

- GNS3. (2023, September). *Getting Started with GNS3 | GNS3 Documentation*.
<https://docs.gns3.com/docs/>
- GNS3. (2024). *Introducción a GNS3 | Documentación de GNS3*. <https://docs.gns3.com/docs/>
- Huawei. (2021, November 18). *What Is Overlay Network? Overlay Network vs. Underlay Network* - Huawei. <https://info.support.huawei.com/info-finder/encyclopedia/en/Overlay+network.html>
- Huawei. (2022, December 27). *Conoce sobre VETP en una red con VXLAN*.
<https://forum.huawei.com/enterprise/es/Conoce-sobre-VETP-en-una-red-con-VXLAN/thread/667235030830825472-667212883219591168>
- Interpolados. (2017, May). *Etiquetado de tramas de Ethernet para la identificación de VLAN*.
<https://interpolados.wordpress.com/2017/05/01/etiquetado-de-tramas-de-ethernet-para-la-identificacion-de-vlan/>
- ITU. (2019, December). *Y.1540 : Internet protocol data communication service - IP packet transfer and availability performance parameters*. <https://www.itu.int/rec/T-REC-Y.1540-201912-I>
- Juniper. (2024). *EX4400 LINE OF ETHERNET SWITCHES DATASHEET Product description*.
- Juniper. (2025, January 17). *Descripción general de BGP | Junos OS | Juniper Networks*.
<https://www.juniper.net/documentation/mx/es/software/junos/bgp/topics/topic-map/bgp-overview.html>
- Lagla Gallardo, C. P. (2023). *Prueba de concepto de VXLAN-EVPN en infraestructura OpenNetworking*. <http://repositorio.puce.edu.ec:80/handle/22000/21243>
- Lebiednik, B., Mangal, A., & Tiwari, N. (2016). *A Survey and Evaluation of Data Center Network Topologies*. <https://arxiv.org/abs/1605.01701v1>

- Liu, Y., Muppala, J. K., Veeraraghavan, M., Lin, D., & Hamdi, M. (2013). *Data Center Networks*.
<https://doi.org/10.1007/978-3-319-01949-9>
- Luo, L., Guo, D., Li, W., Zhang, T., Xie, J., & Zhou, X. (2015). Compound graph based hybrid data center topologies. *Frontiers of Computer Science*, 9(6), 860–874.
<https://doi.org/10.1007/S11704-015-4483-5/METRICS>
- Mitkov Angelov, I. (2021). *Análisis de paquetes con Wireshark, estudio de vulnerabilidades*.
<https://riunet.upv.es/handle/10251/174236>
- Naranjo, E. F., & Salazar Ch, G. D. (2018). Underlay and overlay networks: The approach to solve addressing and segmentation problems in the new networking era: VXLAN encapsulation with Cisco and open source networks. *2017 IEEE 2nd Ecuador Technical Chapters Meeting, ETCM 2017, 2017-January*, 1–6. <https://doi.org/10.1109/ETCM.2017.8247505>
- Nvidia. (2024, September 16). *Specifications - NVIDIA Docs*.
<https://docs.nvidia.com/networking/display/sn2000pub/specifications>
- Ogudo, K. A. (2019). Analyzing generic routing encapsulation (GRE) and IP Security (IPSec) tunneling protocols for secured communication over public networks. *IcABCD 2019 - 2nd International Conference on Advances in Big Data, Computing and Data Communication Systems*. <https://doi.org/10.1109/ICABCD.2019.8851004>
- Pilicita Garrido, A., Borja López, Y., Gutiérrez Constante, G., Pilicita Garrido, A., Borja López, Y., & Gutiérrez Constante, G. (2020). Rendimiento de MariaDB y PostgreSQL. *Revista Científica y Tecnológica UPSE (RCTU)*, 7(2), 9–16.
<https://doi.org/10.26423/RCTU.V7I2.538>

- Razo Achig, A. P. (2022). *Automatización de redes utilizadas para EOT : análisis conceptual del protocolo VXLAN para la virtualización de centro de datos para EOT*. <http://bibdigital.epn.edu.ec/handle/15000/23199>
- Salazar-Chacón, G., & Marrone, L. (2022). Open Networking for Modern Data Centers Infrastructures: VXLAN Proof-of-Concept Emulation using LNV and EVPN under Cumulus Linux. *6th IEEE Ecuador Technical Chapters Meeting, ETCM 2022*. <https://doi.org/10.1109/ETCM56276.2022.9935681>
- Shazni Mursaleen. (n.d.). *Los 5 problemas de rendimiento de red más comunes - CSG Technologies*. Retrieved June 3, 2025, from <https://csgtechnologies.net/the-5-most-common-network-performance-issues-and-how-to-fix-it/>
- Tobón, S. (2013). *Formación integral y competencias Pensamiento complejo, currículo, didáctica y evaluación*.
- Tripathi, S., Chickering, R., & Gainsley, J. (2015). Distributed control plane for high performance switchbased VXLAN overlays. *ANCS 2015 - 11th 2015 ACM/IEEE Symposium on Architectures for Networking and Communications Systems*, 185–186. <https://doi.org/10.1109/ANCS.2015.7110132>
- You, L., Tang, H., Zhang, J., & Li, X. (2020). Fast Configuration Change Impact Analysis for Network Overlay Data Center Networks. *ACM International Conference Proceeding Series*, 8–15. <https://doi.org/10.1145/3411029.3411031>

ANEXOS

ANEXO 1: Instalación de disco SSD en servidor DELL POWEREDGE R750xs

Objetivo:

- El propósito de este anexo es proporcionar una guía detallada para la instalación física y configuración de un nuevo disco SSD en el servidor DELL PowerEdge R750xs, ubicado en el laboratorio de Fibra óptica de la carrera de Telecomunicaciones. Este procedimiento está diseñado para asegurar la correcta integración al sistema de almacenamiento del servidor.

El motivo para agregar un nuevo disco de estado sólido (SSD) al servidor que aloja el sistema Proxmox VE es la falta de espacio disponible, lo que impide la carga y gestión de la máquina virtual de este trabajo de grado en donde se encuentra implementado el entorno de pruebas. El disco SSD, detallado en la Figura 70, es de la marca Intel y tiene un tamaño de 2.5 pulgadas, con su respectiva bandeja de 3.5 pulgadas, lo que lo hace compatible para la integración en servidores DELL. Este tipo de disco tiene una alta capacidad para soportar temperaturas altas y un excelente rendimiento, características que lo hacen ideal para entornos de alto rendimiento como los servidores.

Figura 70

Disco SSD listo para acoplar al servidor Dell PowerEdge R750xs.

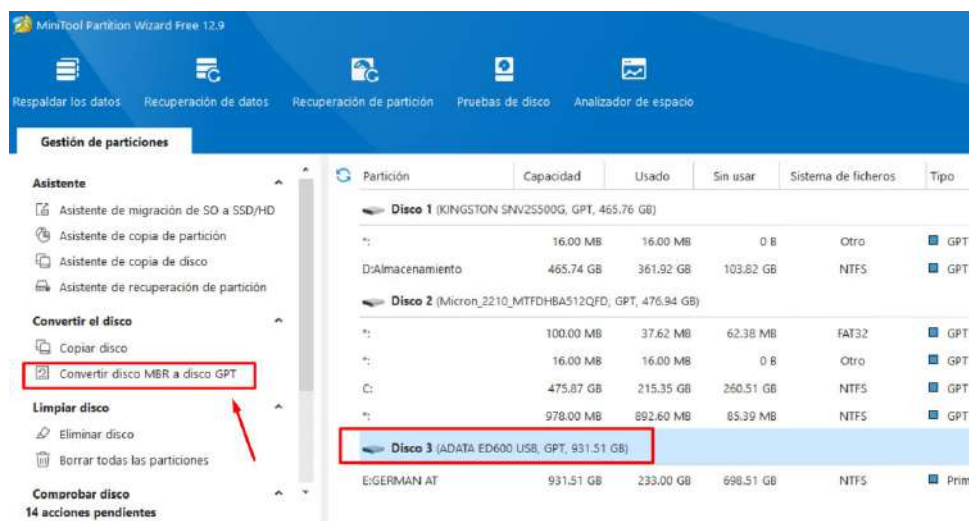


Antes de montar el disco SSD al servidor, es importante observar el formato de particionamiento en el que encuentra el disco. Por defecto, este disco viene configurado en el formato MBR (Master Boot Record), el cual es un esquema de particionamiento antiguo y es por esa razón que al momento de instalarle no es reconocido por el equipo DELL.

Debido a esas limitaciones, se cambia el disco a formato GPT (GUID Partition Table), que es un esquema de particionamiento más moderno y robusto. Para realizar este cambio de formato de MBR a GPT. En la Figura 71, se observa el software especializado como MiniTool Partition Wizard, el cual permite realizar esta conversión de formato de manera sencilla.

Figura 71

Cambio de formato de disco de MBR a GPT.



Al tener realizado el cambio de formato de particionamiento al disco SSD, se agrega al servidor DELL en una de las bandejas disponibles, tal y como se indica en la Figura 72, el equipo puede estar apagado o encendido. Es recomendable tenerle apagado para luego ingresar a la BIOS y realizar el siguiente procedimiento

Figura 72

Montaje de disco SSD en servidor DELL PowerEdge R750xs.

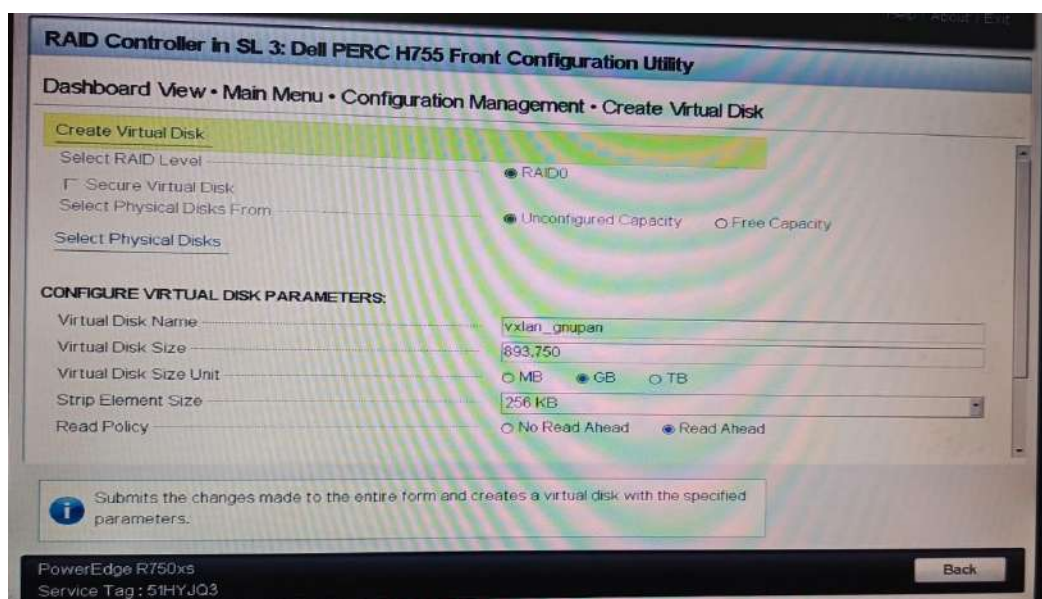


El siguiente paso consiste en acceder a la iDRAC (Integrated Dell Remote Access Controller) del servidor DELL, presionando la tecla especial F2 durante el arranque. Al ingresar, se presentará una pantalla con diversas opciones. En este caso, seleccionamos la opción DELL PERC H755, que corresponde a la controlado RAID del servidor.

Una vez dentro de la interfaz de la controladora, se mostrará el dashboard principal. Desde aquí, se accede a la opción Configure en la parte inferior. Dentro de este menú, se crea el nuevo disco virtual que permite gestionar el nuevo almacenamiento SSD. El proceso de configuración se realiza según los parámetros que se muestran en la Figura 73.

Figura 73

Creación del nuevo disco virtual.



Una vez creado el disco virtual, el siguiente paso es añadirlo a Proxmox VE. En la sección de Disks dentro de la interfaz de Proxmox, se puede visualizar un nuevo disco con el formato /dev/sda de tipo unknown. Para proceder, se inicia una terminal con privilegios de administrador para poder tener acceso por completo al sistema de Proxmox VE.

En la Figura 74, se detallan los comandos necesarios para la creación y configuración del nuevo disco. El primer comando para utilizar es `fdisk`, que permite gestionar una nueva partición de un disco. En este caso, se ejecuta sobre el disco `/dev/sda`. A continuación, se utiliza el comando `n` para crear una nueva partición. Al ejecutar ese comando, se despliegan varias opciones de configuración:

- ✓ **Número de particiones (Partition number):** se utiliza como “1” para crear una única partición.
- ✓ **Sector inicial de la partición (First sector):** se utiliza el valor predeterminado de 2048.
- ✓ **Espacio de la partición (Last Sector):** esta parte se debe de dejar por defecto, ya que esto indicaría que se va a utilizar todo el espacio disponible del disco para la nueva partición.

Finalmente, para guardar los cambios realizados con la letra `w`, lo que confirma la creación y escritura de la partición del disco.

Figura 74

Creación de una nueva partición en Proxmox VE.

```
root@pv5:~# fdisk /dev/sda

Welcome to fdisk (util-linux 2.36.1).
Changes will remain in memory only, until you decide to write them.
Be careful before using the write command.

Command (m for help): n
Partition number (1-128, default 1): 1
First sector (34-1874329566, default 2048): 2048
Last sector, +/-sectors or +/-size{K,M,G,T,P} (2048-1874329566, default 1874329566):

Created a new partition 1 of type 'Linux filesystem' and of size 893.7 GiB.

Command (m for help): w
The partition table has been altered.
Calling ioctl() to re-read partition table.
Syncing disks.
```

ANEXO 2: Importación y configuración de VMs en Proxmox VE

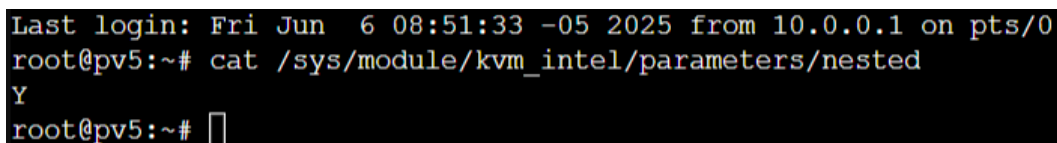
Objetivo:

- Este anexo tiene como objetivo proporcionar un procedimiento detallado para la importación y configuración de máquinas virtuales (VM) en un servidor Proxmox VE, con el fin de integrar eficientemente entornos virtualizados desde plataformas externas como VMware Workstation o VirtualBox.

El primer paso en el proceso de importación de máquinas virtuales en Proxmox VE es verificar que la característica de virtualización aninada se encuentre habilitada en el nodo físico que ejecuta Proxmox. Para ello, se utiliza el comando que se indica en la Figura 75, que permite comprobar si esa función esta habilitada. En este caso, aparece la letra Y, confirmando el estado activo de dicha característica.

Figura 75

Verificación de virtualización aninada de Proxmox.



```
Last login: Fri Jun 6 08:51:33 -05 2025 from 10.0.0.1 on pts/0
root@pv5:~# cat /sys/module/kvm_intel/parameters/nested
Y
root@pv5:~#
```

El siguiente paso exportar la máquina virtual, la cual se va a importar en Proxmox VE. Este proceso es muy importante, ya que evita la pérdida de información almacenada en la máquina virtual original. En este caso, se exporta la GNS3 VM en donde se encuentra el entorno de pruebas desarrollado para este trabajo de grado. Al exportar la máquina, se genera un archivo en formato OVA, el cual se lo va a comprimir en un archivo ZIP para luego transferir a lo que sería el servidor.

El archivo comprimido debe ser importado a Proxmox VE. La Figura 76 muestra la conexión por FTP entre el equipo donde se encuentra la máquina virtual exportada y el servidor

Proxmox, permitiendo así la transferencia segura de la información y la Figura 77, indica el proceso de transferencia de la información a una carpeta temporal de Proxmox VE.

Figura 76

Conexión de FTP entre cliente y servidor Proxmox.

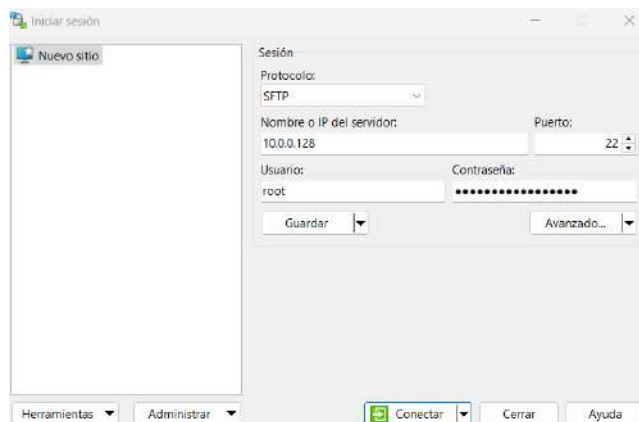
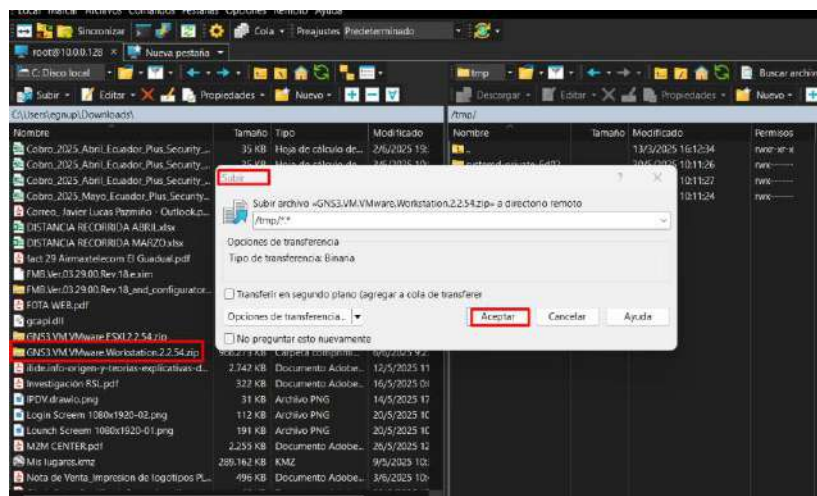


Figura 77

Transferencia de la máquina virtual a Proxmox VE.



Una vez finalizada la transferencia del archivo comprimido que contiene la máquina virtual, el siguiente paso es descomprimir el archivo ZIP. Para ello, se debe abrir una terminal en el servidor y ejecutar el comando correspondiente para extraer el contenido. El proceso de descompresión se muestra en la Figura 78.

Figura 78

Descomprimir archivo en Proxmox VE.

```
root@pv5:/tmp# unzip GNS3.VM.VMware.Workstation.2.2.54.zip
Archive:  GNS3.VM.VMware.Workstation.2.2.54.zip
  inflating: GNS3 VM.ova
root@pv5:/tmp#
```

Una vez descomprimido el archivo, se obtiene un archivo en formato OVA, el cual debe ser descomprimido utilizando el comando `tar -xvf 'GNS3 VM.ova'`. Este paso genera varios archivos OVF y VMDK, los cuales contienen la configuración de la máquina y los discos virtuales correspondientes.

En la Figura 79 se muestra el comando utilizado para importar la máquina virtual. En este comando, el valor 600 corresponde al ID asignado a la nueva máquina virtual, "GNS3 VM.ovf" es el archivo que describe la máquina virtual y se importa, y finalmente se especifica el nombre del almacenamiento local en el que se alojará la máquina virtual. Este proceso permite que la máquina virtual se agregue correctamente al entorno de Proxmox VE.

Figura 79

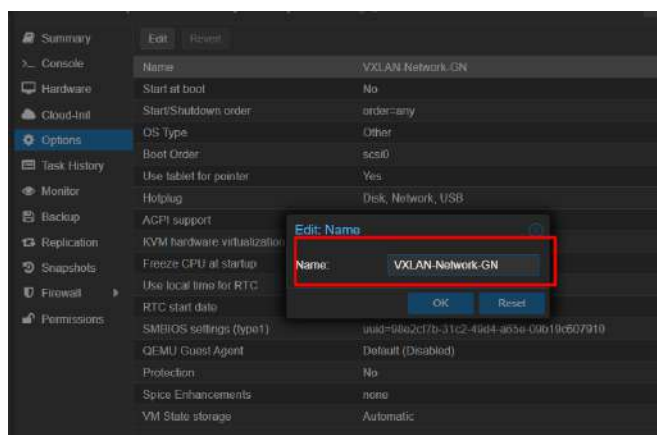
Importación de máquina virtual en Proxmox VE.

```
root@pv5:/tmp# qm importovf 600 'GNS3 VM.ovf' vxlan_GN
Logical volume "vm-600-disk-0" created.
transferred 0.0 B of 19.5 GiB (0.00%)
transferred 200.0 MiB of 19.5 GiB (1.00%)
transferred 400.0 MiB of 19.5 GiB (2.00%)
transferred 600.0 MiB of 19.5 GiB (3.00%)
transferred 800.0 MiB of 19.5 GiB (4.00%)
transferred 1000.0 MiB of 19.5 GiB (5.00%)
```

Finalmente, al completar la transferencia de la máquina virtual, esta aparecerá en el apartado del nodo como una máquina ya creada. El siguiente paso es realizar las configuraciones necesarias. En la Figura 80 se muestra el proceso para asignar un nombre a la máquina virtual, lo que permite distinguirla fácilmente de las demás en el entorno de Proxmox VE.

Figura 80

Asignación de nombre a la máquina virtual en Proxmox VE.



En las Figuras 81 y 82 se muestra cómo se han asignado los recursos a la máquina virtual dentro de Proxmox VE. En este caso, se configuró 20 GB de memoria RAM para la máquina virtual, lo que permite un rendimiento adecuado para las tareas que se realizarán. Además, se asignaron 8 procesadores virtuales, que provienen del servidor físico, el Dell PowerEdge R750xs. Es importante aclarar que estos procesadores virtuales no son recursos físicos dedicados, sino que el hipervisor, Proxmox VE, se encarga de gestionar y distribuir los recursos del servidor entre las distintas máquinas virtuales que están corriendo.

Figura 81

Asignación de memoria RAM.

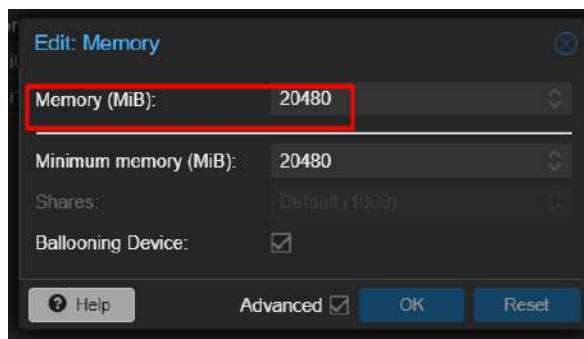
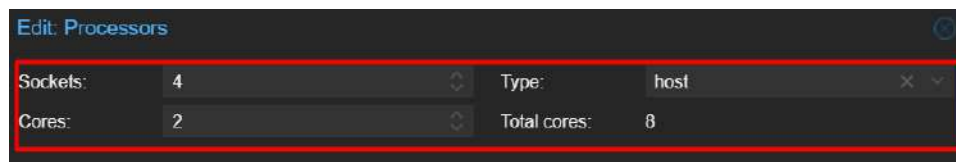


Figura 82

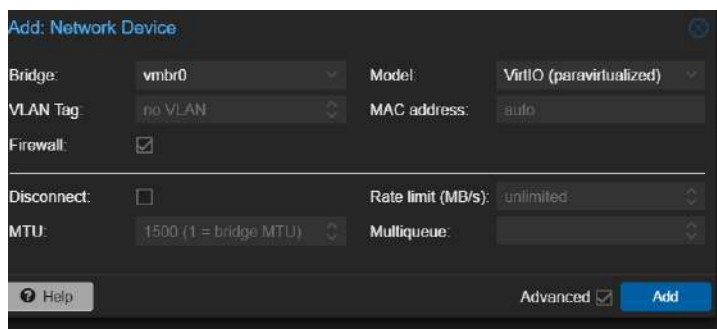
Asignación de números de procesadores.



El siguiente paso es agregar un dispositivo de red a la máquina virtual. Para hacerlo, se debe seleccionar la máquina virtual y, en la parte superior izquierda de la pantalla, dentro de un recuadro azul, hacer clic en la opción para agregar la red o network. En la Figura 83 se muestra cómo se configura el dispositivo de red, donde se elige el tipo VirtIO como controlador de red. Esta interfaz permite una comunicación más eficiente entre la máquina virtual y la red física, asegurando un rendimiento óptimo en la transmisión de datos.

Figura 83

Configuración de interfaz de red de la máquina virtual.

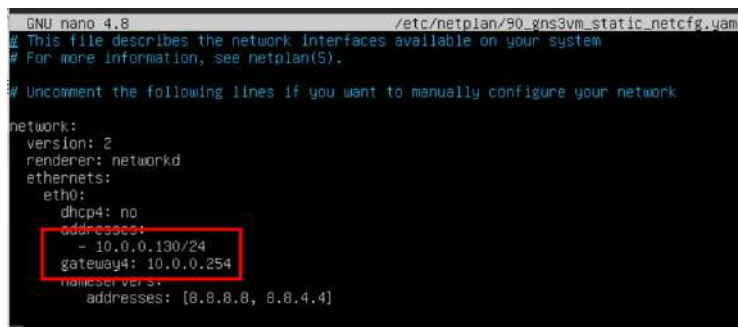


Una vez iniciada la máquina virtual, aparecerá la pantalla azul tradicional de GNS3 VM. El siguiente paso es configurar la dirección IP, para lo cual se selecciona una dirección IP estática dentro del rango de direcciones de la red 10.0.0.0/24, la cual es la misma red en la que se encuentra el controlador de red VirtIO. En la Figura 84 se muestra la dirección IP configurada para la máquina virtual. Esta configuración se realiza a través de la terminal de GNS3 VM, donde se ingresan los

comandos necesarios para asignar la IP estática, asegurando que la máquina virtual esté correctamente integrada en la red y pueda comunicarse con otros dispositivos en el entorno.

Figura 84

Configuración de dirección IP de GNS3 VM.



Finalmente, en una computadora con GNS3 instalado, se debe establecer la conexión con el servidor remoto, en este caso, la máquina virtual GNS3 VM. Esta conexión permite interactuar y gestionar la máquina virtual desde el entorno local de GNS3. En las Figuras 85 y 86 se muestra el proceso de conexión entre la computadora local y el servidor remoto, donde se configura el acceso adecuado y se verifica la conectividad. Este proceso asegura que ambos sistemas estén sincronizados y listos para realizar simulaciones y pruebas de redes de manera efectiva.

Figura 85

Configuración inicial de GNS3.

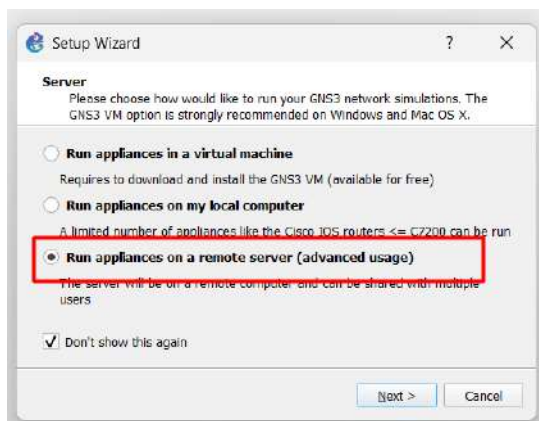


Figura 86

Conexión con el servidor remoto de GNS3 VM.

The screenshot shows the 'Setup Wizard' window for GNS3. The title bar says 'Setup Wizard'. The main heading is 'Remote server' with a subtext: 'Everything will run on a remote server. No data will be saved on this computer.' Below this, there are several input fields: 'Host' with the value '10.0.0.130', 'Port' with a dropdown menu showing '80 TCP', a checked checkbox for 'Enable authentication', 'User' with the value 'admin', and 'Password' with four dots. At the bottom right, there are three buttons: '< Back', 'Next >', and 'Cancel'.

ANEXO 3: Configuraciones del entorno de pruebas

Objetivo:

- Este anexo tiene como objetivo documentar detalladamente las configuraciones implementadas en el entorno de pruebas utilizado para el desarrollo del estudio.

----- CLAVES DE ACCESO ----- GNS3 VM:

- Usuario: root
- Contraseña: root

SWITCH CUMULUS:

- Usuario: cumulus
- Contraseña: cumulus1

MAQUINAS VIRTUALES DEBIAN:

- Usuario: debian
- Contraseña: debian

----- CONFIGURACION DE NETWORK MANAGER -----

1. Actualizacion de paquetes
`apt-get update`
`apt-get upgrade`
2. Instalar dnsmasq
`apt-get install dnsmasq`
3. Definir dirección IP para la interfaz que va al Switch que conecta a los equipos VXLAN
`nano /etc/network/interfaces`

```
auto eth1
iface eth1 inet static
    address 10.0.0.1
    netmask 255.255.255.240
```

4. Configurar el hostname de los equipos

```
nano /etc/hosts
10.0.0.2      leaf-01
10.0.0.3      leaf-02
10.0.0.4      leaf-03
10.0.0.5      leaf-04
10.0.0.6      leaf-05
10.0.0.7      spine-01
10.0.0.8      spine-02
```

5. Configurar las interfaces de gestión

```
dhcp-range=10.0.0.2,10.0.0.14, 12h

dhcp-host=0c:d0:04:b2:00:00,leaf-01,10.0.0.2,12h
dhcp-host=0c:bf:54:69:00:00,leaf-02,10.0.0.3,12h
dhcp-host=0c:44:f8:a6:00:00,leaf-03,10.0.0.4,12h
dhcp-host=0c:e8:a8:e9:00:00,leaf-04,10.0.0.5,12h
dhcp-host=0c:f6:b4:35:00:00,leaf-05,10.0.0.6,12h

dhcp-host=0c:33:e1:56:00:00,spine-01,10.0.0.7,12h
dhcp-host=0c:33:fe:04:00:00,spine-02,10.0.0.8,12h
```

6. Iniciamos el servicio

```
/etc/init.d/dnsmasq start
/etc/init.d/dnsmasq enable
/etc/init.d/dnsmasq status

/etc/init.d/dnsmasq restart ---- En caso de que se requiera reiniciar
```

----- CONFIGURACION DE INTERFACES DE LOOPBACK DE LOS EQUIPOS -----

```
nv set interface lo ip address 192.168.0.1/32 ----- SPINE 1
nv set interface lo ip address 192.168.0.2/32 ----- SPINE 2
nv set interface lo ip address 192.168.0.3/32 ----- LEAF 1
```

```

nv set interface lo ip address 192.168.0.4/32 ----- LEAF 2
nv set interface lo ip address 192.168.0.5/32 ----- LEAF 3
nv set interface lo ip address 192.168.0.6/32 ----- LEAF 4
nv set interface lo ip address 192.168.0.7/32 ----- LEAF 5

```

----- CONFIGURACION DE BORDER GATEWAY PROTOCOL -----

***** SPINE 1 *****

```

nv set interface lo ip address 192.168.0.1/32
nv set interface swp1-5
nv set router bgp autonomous-system 65199
nv set router bgp router-id 192.168.0.1
nv set vrf default router bgp neighbor swp1 remote-as external
nv set vrf default router bgp neighbor swp2 remote-as external
nv set vrf default router bgp neighbor swp3 remote-as external
nv set vrf default router bgp neighbor swp4 remote-as external
nv set vrf default router bgp neighbor swp5 remote-as external
nv set vrf default router bgp address-family ipv4-unicast network
192.168.0.1/32
nv config apply

```

***** SPINE 2 *****

```

nv set interface lo ip address 192.168.0.2/32
nv set interface swp1-6
nv set router bgp autonomous-system 65198
nv set router bgp router-id 192.168.0.2
nv set vrf default router bgp neighbor swp1 remote-as external
nv set vrf default router bgp neighbor swp2 remote-as external
nv set vrf default router bgp neighbor swp3 remote-as external
nv set vrf default router bgp neighbor swp4 remote-as external
nv set vrf default router bgp neighbor swp5 remote-as external
nv set vrf default router bgp neighbor swp6 remote-as external
nv set vrf default router bgp address-family ipv4-unicast network
192.168.0.2/32
nv config apply

```

***** LEAF 1 *****

```

nv set interface lo ip address 192.168.0.3/32
nv set router bgp autonomous-system 65101
nv set router bgp router-id 192.168.0.3
nv set vrf default router bgp neighbor swp1 remote-as external
nv set vrf default router bgp neighbor swp2 remote-as external
nv set vrf default router bgp neighbor swp5 remote-as external
nv set vrf default router bgp neighbor swp6 remote-as external
nv set vrf default router bgp address-family ipv4-unicast network
192.168.0.3/32
nv set vrf default router bgp address-family ipv4-unicast redistribute
connected
nv config apply

```

***** LEAF 2 *****

```
nv set interface lo ip address 192.168.0.4/32
nv set router bgp autonomous-system 65102
nv set router bgp router-id 192.168.0.4
nv set vrf default router bgp neighbor swp1 remote-as external
nv set vrf default router bgp neighbor swp2 remote-as external
nv set vrf default router bgp neighbor swp5 remote-as external
nv set vrf default router bgp neighbor swp6 remote-as external
nv set vrf default router bgp address-family ipv4-unicast network
192.168.0.4/32
nv set vrf default router bgp address-family ipv4-unicast redistribute
connected
nv config apply
```

***** LEAF 3 *****

```
nv set interface lo ip address 192.168.0.5/32
nv set router bgp autonomous-system 65103
nv set router bgp router-id 192.168.0.5
nv set vrf default router bgp neighbor swp1 remote-as external
nv set vrf default router bgp neighbor swp2 remote-as external
nv set vrf default router bgp address-family ipv4-unicast network
192.168.0.5/32
nv set vrf default router bgp address-family ipv4-unicast redistribute
connected
nv config apply
```

***** LEAF 4 *****

```
nv set interface lo ip address 192.168.0.6/32
nv set router bgp autonomous-system 65104
nv set router bgp router-id 192.168.0.6
nv set vrf default router bgp neighbor swp1 remote-as external
nv set vrf default router bgp neighbor swp2 remote-as external
nv set vrf default router bgp address-family ipv4-unicast network
192.168.0.6/32
nv set vrf default router bgp address-family ipv4-unicast redistribute
connected
nv config apply
```

***** LEAF 5 *****

```
nv set interface lo ip address 192.168.0.7/32
nv set router bgp autonomous-system 65105
nv set router bgp router-id 192.168.0.7
nv set vrf default router bgp neighbor swp1 remote-as external
nv set vrf default router bgp neighbor swp2 remote-as external
nv set vrf default router bgp address-family ipv4-unicast network
192.168.0.7/32
nv set vrf default router bgp address-family ipv4-unicast redistribute
connected
nv config apply
```

----- CONFIGURACION DE INTERFACES DE ACCESO -----

```
***** LEAF 1 *****
nv set interface swp3 bridge domain br_default access 10
nv set interface swp4 bridge domain br_default access 20

***** LEAF 2 *****
nv set interface swp3 bridge domain br_default access 30

***** LEAF 4 *****
nv set interface swp4 bridge domain br_default access 10
nv set interface swp5 bridge domain br_default access 20
nv set interface swp6 bridge domain br_default access 30

***** LEAF 5 *****
nv set interface swp6 bridge domain br_default access 10
```

----- CONFIGURACION DE VXLAN -----

```
sudo nano /etc/network/interfaces
```

```
***** LEAF 1 *****
auto br_default
iface br_default
bridge-ports swp3 swp4 vni10 vni20
bridge-vlan-aware yes
bridge-vids 10 20
bridge-pvid 1

auto vni10
iface vni10
    bridge-access 10
    mstpctl-bpduguard yes
    mstpctl-portbpdufilter yes
    vxlan-local-tunnelip 192.168.0.3
    vxlan-remoteip 192.168.0.6
    vxlan-remoteip 192.168.0.7
    vxlan-id 100

auto vni20
iface vni20
    bridge-access 20
    mstpctl-bpduguard yes
    mstpctl-portbpdufilter yes
    vxlan-local-tunnelip 192.168.0.3
    vxlan-remoteip 192.168.0.6
    vxlan-id 200
```

```

***** LEAF 2 *****
auto br_default
iface br_default
bridge-ports swp3 vni30
bridge-vlan-aware yes
bridge-vids 30
bridge-pvid 1

auto vni30
iface vni30
    bridge-access 30
    mstpctl-bpduguard yes
        mstpctl-portbpdufilter yes
    vxlan-local-tunnelip 192.168.0.4
    vxlan-remoteip 192.168.0.6
    vxlan-id 300

***** LEAF 4 *****
auto br_default
iface br_default
bridge-ports swp4 swp5 swp6 vni10 vni20 vni30
bridge-vlan-aware yes
bridge-vids 10 20 30
bridge-pvid 1

auto vni10
iface vni10
    bridge-access 10
    mstpctl-bpduguard yes
        mstpctl-portbpdufilter yes
    vxlan-local-tunnelip 192.168.0.6
    vxlan-remoteip 192.168.0.3
    vxlan-remoteip 192.168.0.7
    vxlan-id 100

auto vni20
iface vni20
    bridge-access 20
    mstpctl-bpduguard yes
        mstpctl-portbpdufilter yes
    vxlan-local-tunnelip 192.168.0.6
    vxlan-remoteip 192.168.0.3
    vxlan-id 200

auto vni30
iface vni30
    bridge-access 30
    mstpctl-bpduguard yes
        mstpctl-portbpdufilter yes
    vxlan-local-tunnelip 192.168.0.6
    vxlan-remoteip 192.168.0.3
    vxlan-id 300

```

```
***** LEAF 5 *****
auto br_default
iface br_default
bridge-ports swp6 vni10
bridge-vlan-aware yes
bridge-vids 10
bridge-pvid 1

auto vni10
iface vni10
    bridge-access 10
    mstpctl-bpduguard yes
        mstpctl-portbpdufilter yes
    vxlan-local-tunnelip 192.168.0.7
    vxlan-remoteip 192.168.0.6
    vxlan-remoteip 192.168.0.3
    vxlan-id 100
```

ANEXO 4: Configuraciones de acceso y conexión de equipos físicos a la red

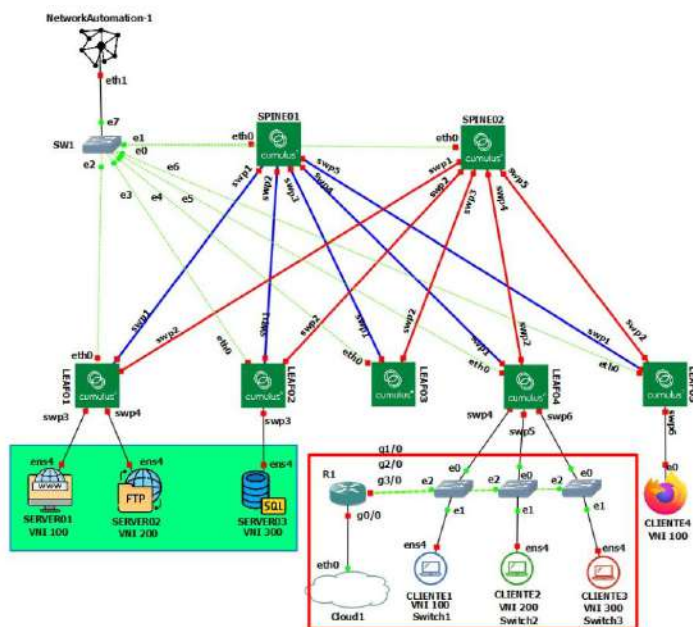
Objetivo:

- Proporcionar el procedimiento detallado para la integración de equipos físicos a la red, asegurando el correcto funcionamiento dentro de la infraestructura de red existente.

El primer paso para la implementación de este anexo consiste en modificar la topología de red previamente presentada en este trabajo, integrando un router Cisco 7200, tres switches de acceso y una nube virtual (Cloud), como se ilustra en la Figura 87. La incorporación de estos equipos tiene como objetivo permitir que los dispositivos conectados a la red física, en la cual se encuentra anclado el servidor Proxmox VE y alojada la máquina virtual, puedan acceder a los servidores a través VXLAN. En este contexto, la nube virtual (Cloud1) se conecta a la interfaz del router G0/0 a través de la interfaz Eth0, que corresponde a la interfaz de la máquina virtual GNS3 VM. Es mediante esta conexión que todo el tráfico generado por los clientes externos ingresará a la topología de red configurada en este trabajo.

Figura 87

Topología de red modificada para el acceso desde la red física.



Comenzamos configurando las interfaces del router, donde las interfaces g1/0, g2/0 y g3/0 se asignan como el Gateway o puerta de entrada para servidores y clientes de cada VXLAN. Dado que el tráfico en la red se encuentra segmentado, es fundamental realizar esta adaptación de las tres interfaces de red. En la Figura 88, se detallan las direcciones IP asignadas a cada interfaz, las cuales servirán como puerta de enlace para servidores y clientes de las respectivas VXLAN.

Figura 88

Configuración de direcciones IP en interfaces de Router.

```
R1(config)#int g1/0
R1(config-if)#ip add
R1(config-if)#ip address 10.10.10.1 255.255.255.0
R1(config-if)#exit
R1(config)#int g2/0
R1(config-if)#ip add
R1(config-if)#ip address 10.10.20.1 255.255.255.0
R1(config-if)#no shu
R1(config-if)#no shutdown
R1(config)#int g3/0
R1(config-if)#ip ad
R1(config-if)#ip ad
R1(config-if)#ip add
R1(config-if)#ip address 10.10.30.1 255.255.255.0
R1(config-if)#no shu
```

En la Figura 89 se muestra la configuración de la interfaz del router que establece la conexión con la nube virtual (Cloud1). En esta configuración, se asigna una dirección IP que corresponde al segmento de red del servidor Proxmox VE, garantizando que los dispositivos dentro de este segmento puedan comunicarse correctamente. Además, se configura una ruta estática, en la cual se define el Gateway de la red del servidor Proxmox VE como la puerta de enlace. Esto permitirá que el tráfico se encamine de manera adecuada hacia redes externas.

Figura 89

Configuración de la interfaz WAN.

```
R1(config)#int g0/0
R1(config-if)#ip ad
R1(config-if)#ip add
R1(config-if)#ip address 10.0.0.249 255.255.255.0
R1(config-if)#no shu
R1(config-if)#no shutdown
R1(config-if)#
*Jun 30 19:27:18.731: %LINK-3-UPDOWN: Interface Giga
*Jun 30 19:27:19.731: %LINEPROTO-5-UPDOWN: Line prot
R1(config-if)#ip route 0.0.0.0 0.0.0.0 10.0.0.254
```

Ahora se procede a configura las listas de acceso (ACL) en esta configuración se utilizan para controlar y filtrar el tráfico de la red, permitiendo el paso de paquetes según criterios de dirección IP y Puertos. En este caso se realiza una lista de acceso extendida que se llama NAT_WEB para definir el tráfico debe ser dirigido a los servidores internos a través de la traducción de direcciones de red.

Figura 90

Configuración de ACL.

```
R1(config)#ip access-list extended NAT_WEB
R1(config-ext-nacl)#static tcp 10.10.10.3 21 10.0.0.249 21 extendable
R1(config)#no ip nat inside source static tcp 10.10.10.3 21 10.0.0.249 21 exte$
R1(config)#de source static tcp 10.10.20.3 21 10.0.0.249 21 extendable
R1(config)#de source static tcp 10.10.20.3 3306 10.0.0.249 3306 extendable
R1(config)#no ip nat inside source static tcp 10.10.20.3 3306 10.0.0.249 3306 $
R1(config)#ce static tcp 10.10.30.3 3306 10.0.0.249 3306 extendable
%Translation not found
R1(config)#ip nat inside source static tcp 10.10.30.3 3306 10.0.0.249 3306 ext$
```


Finalmente, en las interfaces que conectan a las VXLAN, se activa la traducción de direcciones de red (NAT) indicando que esa interfaz es parte de la red interna que utilizará el NAT para poder acceder a las redes externas. Esta configuración como se muestra en la Figura 91, asegura que el tráfico proveniente de redes internas de esas interfaces sea correctamente gestionando por la NAT permitiendo la salida a la red externa.

Figura 91

Configuración de traducción de interfaces de redes internas.

```
R1(config)#int g2/0
R1(config-if)#ip nat ins
R1(config-if)#ip nat inside
R1(config-if)#exit
R1(config)#int g3/0
R1(config-if)#ip nat ins
R1(config-if)#ip nat inside
```

Como prueba del funcionamiento, en la Figura 92 se muestra una prueba realizada desde un terminal móvil conectado a la red física del Laboratorio de Fibra Óptica de la Universidad Técnica del Norte. En esta prueba el terminal móvil accede a la página web configurada en la VXLAN 100. Además se realiza la captura de paquetes con Wireshark como se indica en la Figura 93, demostrando que la configuración de la VXLAN y la conectividad de la red están funcionando correctamente, permitiendo que dispositivos de la red física puedan interactuar con los equipos que se encuentran en la red virtualizada.

Figura 92

Ingresa al servidor WEB desde un terminal móvil.

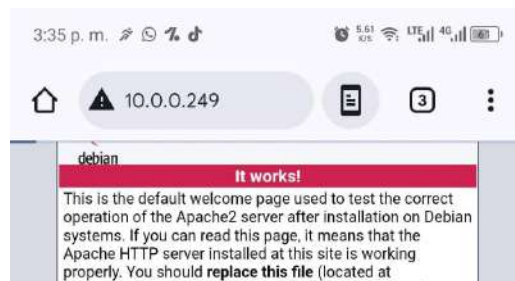


Figura 93

Captura del tráfico generado desde el terminal móvil al servidor Web.

89	176.779425	10.0.0.254	10.10.10.3	HTTP	500 GET /icons/openlogo-75.png HTT
90	176.951299	10.0.0.254	10.10.10.3	TCP	128 46228 → 80 [ACK] Seq=846 Ack=4
91	177.012434	10.0.0.254	10.10.10.3	TCP	116 46228 → 80 [ACK] Seq=846 Ack=6
92	177.012448	10.0.0.254	10.10.10.3	TCP	116 46228 → 80 [ACK] Seq=846 Ack=7
93	177.042435	10.0.0.254	10.10.10.3	TCP	116 46228 → 80 [ACK] Seq=846 Ack=9
94	177.042522	10.0.0.254	10.10.10.3	TCP	116 46228 → 80 [ACK] Seq=846 Ack=9
95	177.063404	10.0.0.254	10.10.10.3	HTTP	490 GET /favicon.ico HTTP/1.1
96	177.355593	10.0.0.254	10.10.10.3	TCP	116 46228 → 80 [ACK] Seq=1220 Ack=
100	182.498748	10.0.0.254	10.10.10.3	TCP	116 46228 → 80 [ACK] Seq=1220 Ack=
101	182.742985	10.0.0.254	10.10.10.3	TCP	116 46228 → 80 [FIN, ACK] Seq=1220


```

Frame 89: 500 bytes on wire (4000 bits), 500 bytes captured (4000 bits) on interface -, id 0
Ethernet II, Src: 0c:33:fe:04:00:01 (0c:33:fe:04:00:01), Dst: 0c:d0:04:b2:00:02 (0c:d0:04:b2:00:02)
Internet Protocol Version 4, Src: 192.168.0.6, Dst: 192.168.0.3
User Datagram Protocol, Src Port: 42579, Dst Port: 4789
Virtual eXtensible Local Area Network
Ethernet II, Src: ca:01:b6:bc:00:1c (ca:01:b6:bc:00:1c), Dst: 0c:af:54:3e:00:00 (0c:af:54:3e:00:00)
Internet Protocol Version 4, Src: 10.0.0.254, Dst: 10.10.10.3
Transmission Control Protocol, Src Port: 46228, Dst Port: 80, Seq: 462, Ack: 3365, Len: 384
Hypertext Transfer Protocol

```

Además, se realiza una prueba para verificar que la red propuesta en esta investigación, basa en un diseño Spine-Leaf con VXLAN, pueda integrarse correctamente con la infraestructura de red física de la Universidad Técnica del Norte, utilizando la red Eduroam. Para ello, se implementa un proceso de traducción de direcciones (NAT) que permite el acceso a equipos dentro de la simulación.

En la Figura 94 se muestra la dirección IP asignada por DHCP en la red Eduroam, que es 172.20.147.12. A partir de esta dirección IP, nos conectamos al servidor web. Para verificar el proceso de conexión, se realiza una captura de tráfico utilizando Wireshark, en la cual se obtiene paquetes TCP por el puerto 80. En la Figura 95, se muestra la dirección IP de origen que corresponde a la computadora física, mientras que la dirección de destino es la del router cisco simulado, donde se lleva a cabo la traducción de direcciones (NAT).

Figura 94

Dirección IP de laptop asignada por DHCP de Eduroam.

```

Adaptador de LAN inalámbrica Wi-Fi:

Sufijo DNS específico para la conexión. . . : utn.edu.ec
Dirección IPv6 . . . . . : 2801:10:6800:128::1ac2
Vínculo: dirección IPv6 local. . . : fe80::f779:7aa4:5650:8c98%3
Dirección IPv4. . . . . : 172.20.147.12
Máscara de subred . . . . . : 255.255.224.0
Puerta de enlace predeterminada . . . . . : fe80::42f0:78ff:fe45:efda%3
                                           172.20.128.1

```

Figura 95

Captura de Wireshark desde computadora física hacia simulación por Wi-Fi.

4247	121.756108	172.20.147.12	10.0.0.249	TCP	60
4361	122.008486	172.20.147.12	10.0.0.249	TCP	60
5973	310.699026	172.20.147.12	10.0.0.249	TCP	60
5992	310.953343	172.20.147.12	10.0.0.249	TCP	60

Frame 4247: 66 bytes on wire (528 bits), 66 bytes captured (528 bits) on interface
Ethernet II, Src: AzureWaveTec 92:76:2f (b4:8c:9d:92:76:2f), Dst: Cisco_45:e
Destination: Cisco_45:ef:da (40:f0:78:45:ef:da)
Source: AzureWaveTec 92:76:2f (b4:8c:9d:92:76:2f)
Type: IPv4 (0x0800)
[Stream index: 1]
Internet Protocol Version 4, Src: 172.20.147.12, Dst: 10.0.0.249
Transmission Control Protocol, Src Port: 55000, Dst Port: 80, Seq: 0, Len: 0

Finalmente, se lleva a cabo una prueba con cable Ethernet, conectando una computadora a través de un cable UTP a uno de los puertos disponibles en el laboratorio de Fibra Óptica. En esta configuración, se asignó la dirección IP 192.168.34.143 a la computadora mediante DHCP.

Para comprobar la conectividad, se realiza pruebas con el servidor FTP utilizando FileZilla. En la Figura 96, se muestra una captura del tráfico con Wireshark, donde se observa que, al igual que en las pruebas anteriores, la conexión al servidor se realiza a través de la dirección IP NAT 10.0.0.249.

Figura 96

Prueba de ingreso a servidor FTP mediante conexión alámbrica de red física.

