

VXLAN EVPN on Next-Generation Switches: Modern Data Center Network Architecture

Sasikumar Sadayan*

CISCO SYSTEMS, USA

* Corresponding Author Email: mr.sasikumarsadayan@gmail.com- ORCID: 0000-0002-5247-8850

Article Info:

DOI: 10.22399/ijcesn.4232

Received : 25 September 2025

Accepted : 03 November 2025

Keywords

VXLAN,
EVPN,
BGP,
Network Virtualization,
Data Center Fabric,
Multi-tenancy

Abstract:

Digital extensible LAN generation mixed with ethernet vpn to manipulate aircraft has emerged as a transformative solution for cutting-edge data center networking challenges, addressing the fundamental boundaries of traditional layer 2 switching in virtualized and multi-tenant environments. The exponential growth of cloud computing and the acceleration of multicloud adoption styles have created extraordinary needs for scalable community virtualization technology that may support vast community segmentation, even as it maintains gold standard performance characteristics. Vxlan EVPN leverages BGP as an advanced manipulation plane mechanism, enabling distributed knowledge of and marketing of endpoint information throughout fabric architectures without requiring centralized management factors or flood-and-learn mechanisms that plague conventional bridging protocols. The technology presents considerable expansion in community identifier capability through its improved addressing scheme, representing a sizeable development past the limitations of conventional VLAN implementations. Superior implementations make use of hardware-based encapsulation and decapsulation abilities in modern switching platforms, making sure that overlay networking functionality does not compromise forwarding performance, even as it allows obvious layer 2 connectivity throughout the infrastructure. The dispensed nature of EVPN manages plane operations gives inherent redundancy and scalability via general BGP mechanisms, with sophisticated commercial course techniques that optimize convergence conduct in large-scale deployments. Integration of routing and bridging capability through distributed anycast gateway configurations permits the most desirable forwarding paths for inter-subnet communication while keeping the operational advantages of dispersed routing architectures, which could scale horizontally throughout multiple physical locations.

1. Introduction

The rapid boom of cloud computing, virtualization, and multi-tenant environments has revolutionized statistics center networking desires, driven by the pace of multicloud times, while businesses are increasingly embracing hybrid cloud strategies that integrate public and private infrastructure [1]. Legacy Layer 2 switching solutions are severely challenged in terms of scalability, agility, and isolation features required by contemporary distributed applications, especially when taking into consideration that traditional VLANs are limited to at most 4,094 different identifiers as dictated by the IEEE 802.1Q standard. Virtual Extensible LAN (VXLAN) with Ethernet VPN (EVPN) control plane is a paradigm change in data center design,

with overlay network capabilities that overcome these limitations through network virtualization and sophisticated control plane innovation. The scalability needs of contemporary data centers reveal the weakness of conventional Ethernet switching when it is examined. Large virtualized environments often need tens of thousands of virtual machines spread over hundreds of physical servers, and each virtual machine could need its broadcast domain for isolation and security. IEEE 802.1Q's 12-bit VLAN ID field only offers 4,094 usable VLANs (without using reserved values 0 and 4,095), which is not enough in large multi-tenant environments where thousands of tenants are each likely to need multiple network segments. The multicloud adoption trends emerging through the latest industry research reflect that organizations

are putting more pressure on network infrastructures to be highly flexible network architectures that seamlessly integrate heterogeneous cloud platforms while preserving uniform security and performance attributes [1].

VXLAN EVPN allows network operators to build scalable multi-tenant network fabrics that can stretch across multiple physical sites while preserving logical isolation and optimal forwarding behavior. The 24-bit VXLAN Network Identifier (VNI) field offers more than 16.7 million distinct network segments, a 4,000-fold increase in available network identifiers over traditional VLANs. This technology stack uses BGP as the control plane protocol, allowing for distributed learning and advertisement of endpoint information across the fabric without the need for flood-and-learn mechanisms that afflict traditional bridging protocols. The deployment of BGP EVPN with all-active multi-homing functions further improves network resiliency by providing redundant connectivity paths with loop-free forwarding behavior using advanced route advertisement mechanisms [2]. The architectural benefits of VXLAN EVPN go beyond mere address space expansion. Decoupling the overlay logical network from the underlay physical infrastructure allows network operators to achieve independent scaling of compute and network resources. The underlay network is dedicated to high-performance IP forwarding with reduced routing protocols, while the overlay network caters to tenant-specific functions like segmentation, mobility, and insertion of services. This dichotomy allows the underlay to be tuned for the highest throughput and lowest latency, while the overlay offers the feature richness and rich flexibility necessary for multi-tenant services. The all-active multi-homing feature that is a part of newer BGP EVPN solutions guarantees that network outages do not cause any service disruptions because traffic can be redistributed dynamically along existing paths without the need for human intervention or involving complicated failover procedures [2].

2. Architectural Foundations and Control Plane Mechanics

2.1 VXLAN Overlay Architecture

VXLAN builds Layer 2 overlay networks on top of pre-existing Layer 3 infrastructure by encapsulating native Ethernet frames in UDP packets, creating virtual networks that are independent of the underlying physical network topology [3]. The 24-bit VXLAN Network Identifier (VNI) offers more

than 16.7 million distinct network segments, compared to the 4,094 VLAN limit of legacy IEEE 802.1Q implementations. VXLAN Tunnel Endpoints (VTEPs) are the points of encapsulation and decapsulation while preserving the overlay-to-underlay mapping between logical overlay networks and physical underlay infrastructure with advanced forwarding mechanisms correlating tenant networks with corresponding tunnel destinations.

Encapsulation is the process of enclosing the native Layer 2 frame in multiple protocol headers to facilitate transmission over IP networks, with VTEPs executing the important task of overlay-to-underlay addressing scheme translation [3]. This encapsulation process maintains the original frame format with the addition of required routing information to be traversed through the IP fabric. The underlying network, commonly realized as a spine-leaf structure, delivers IP connectivity between VTEPs via conventional routing protocols like OSPF or BGP to provide a solid base for overlay operation.

This division of overlay and underlay responsibilities allows each layer to scale and be optimized separately, with the underlay dedicated to high-performance IP forwarding and the overlay dealing with tenant-specific network requirements. Network operators can implement modifications to one layer without affecting service in the other using the modular architecture, which provides operational flexibility unavailable in traditional flat Layer 2 networks [3]. Current VXLAN solutions take advantage of switching platform hardware acceleration features to keep performance up to that of native switching, but with the additional benefits of network virtualization.

2.2 EVPN Control Plane Operations

EVPN builds on BGP by introducing new address families for Layer 2 VPN services, using advanced route distribution constructs that allow reachability information to be propagated efficiently throughout the fabric [4]. The protocol specifies various types of routes that facilitate effective allocation of MAC addresses, IP addresses, and multicast group information throughout the fabric, with each type of route playing specific roles in network state consistency maintenance. Type 2 routes support MAC/IP bindings advertisement, Type 3 routes support multicast traffic distribution, and Type 5 routes support inter-subnet routing within the overlay, forming a complete control plane solution. The distributed design of EVPN avoids the use of control plane elements, having built-in redundancy and scalability through tried-and-tested standard

BGP mechanisms in large-scale service provider networks [4]. Each VTEP holds local MAC address tables and advertises learned information to remote VTEPs via BGP, building a synchronized view of the overall overlay network without having to flood unknown unicast traffic. The control plane design takes advantage of the natural convergence characteristics of BGP and path selection policies to make optimal forwarding decisions throughout the fabric. BGP EVPN deployments in SP environments illustrate the protocol's ability to support massive-scale implementations with thousands of customer locations and millions of MAC addresses, setting the basis for carrier-grade Ethernet VPN services [4]. The interoperation of the protocol with current MPLS infrastructure facilitates the transparent extension of Layer 2 services over wide area networks, while interoperation with IP-based underlays provides support for contemporary data center fabric designs. This deployment model flexibility facilitates organizations in installing EVPN solutions that suit their respective infrastructure needs and operational styles.

3. Data Plane Forwarding and Tunnel Encapsulation

VXLAN data plane operations include complex forwarding decisions based on destination MAC addresses and their corresponding VNI contexts that demand expert-level processing abilities, which have been proven in large-scale data center installations like Google's Jupiter network fabric [5]. If a vtep receives a body, this is meant for a faraway MAC deal with. It does a lookup in its nearby forwarding desk to pick out the proper far-off VTEP and VNI, the use of forwarding architectures that support the large-scale needs of today's statistics facilities. It is then encapsulated with physical Ethernet headers, UDP, IP, and VXLAN before it is sent over the underlay network, forming a multi-layered packet construction that maintains tenant isolation but allows for travel over IP infrastructure. Jupiter network design illustrates how Clos topologies can efficiently accommodate overlay networking demands at record scale, with each data center fabric supporting more than 100,000 servers connected via several stages of switching equipment [5]. Encapsulation maintains the original frame format, along with introducing essential overlay routing data via systematic header construction mechanisms that have to function under the limitations of centralized control systems managing tens of thousands of network devices. The header of the outer IP includes source and destination VTEP addresses taken from the

underlay routing table, and the VXLAN header includes the 24-bit VNI and additional control flags supporting transparent Layer 2 connectivity over IP networks with tenant isolation through VNI separation. Specialized forwarding engines utilized in contemporary data plane implementations are optimized for addressing the computational intensity of overlay networking actions without affecting performance attributes necessary in hyperscale networks [5]. The forwarding operation consists of several table lookups, including a MAC address table lookup, VNI-to-tunnel mapping resolution, and underlay routing table lookup, all of which have to be executed within microsecond timescales to provide line-rate performance over fabric bisection bandwidths of multiple petabits per second. Sophisticated implementations enable hardware-based encapsulation and decapsulation based on special silicon architectures that have developed substantially over ten years of large-scale deployment heritage. The VXLAN packet processing pipeline in contemporary implementations utilizes programmatic processing methods that allow for adaptive flexibility in accommodating changing protocol demands and deployment environments [6]. Network processors specialize in programmable packet processing features that can be tailored via programming languages at a higher level and optimized for network data plane processes. These programmable architectures conduct concurrent tasks such as header parsing, table lookups, encapsulation processing, and quality-of-service marking, allowing end-to-end overlay networking capability with no loss of throughput or addition of jitter that would affect application performance. Dedicated VXLAN processing is made possible by today's switching platforms through programmable packet processing engines that can be programmed using unified high-level programming methods [6]. These next-generation architectures feature flexible processing pipelines that can be reorganized to support protocol-independent packet processing demands, yet retain the performance profiles required for production data center use cases. Programmability within these systems allows network operators to specialize forwarding behavior and optimization approaches according to the particular deployment needs, while the shared programming model streamlines the coding and deployment of sophisticated networking capabilities across heterogeneous hardware platforms.

4. Design and Scalability Factors

4.1 Planning Fabric Size and Performance

VXLAN EVPN fabric design involves special attention to scale factors such as the number of VTEPs, tenants, and endpoints per tenant, where current hyperscale data center architecture has proven capable of serving BGP deployments in thousands of network devices within a single facility [7]. BGP control plane scaling is based on the number of BGP peers and on the amount of route advertisement, making strategic deployment of route reflectors in large-scale deployments inevitable to cope with the computational complexity that arises when deploying BGP at unprecedented scale in data center contexts. Route reflector hierarchies can be used to distribute control plane load and enhance convergence times, with deployment patterns that have been tested and proven in production environments managing enormous volumes of traffic and intricate routing needs. The characteristics of scaling with BGP in large-scale data center deployments demand advanced methods of managing route distribution and convergence behavior, especially while supporting the diverse connectivity needs of contemporary cloud infrastructure [7]. Bandwidth planning should consider both overlay encapsulation overhead and tenant traffic patterns, with VXLAN imposing an overhead of around 50 bytes per packet with the use of protocol headers that are necessary for overlay networking functionality. This encapsulation overhead will have implications for network utilization calculations and necessitate adjustments to the interface MTU settings across the fabric to avoid fragmentation and optimize performance across a range of application workloads typical in modern data center environments. New fabric topologies make use of optimized BGP methods that have been formulated and honed by large-scale deployment in hyperscale settings, where conventional routing protocols are confronted with scale and performance demands unlike anything previously seen [7]. Hierarchical route reflector topologies are used in implementation to facilitate effective routing information distribution and minimize processing loads on specific network components while taking redundancy and failover measures into account for uninterrupted performance during maintenance tasks or unplanned failure situations. Performance planning needs to also consider the intricate interplay between overlay and underlay traffic patterns, especially in situations where tenant workloads produce substantial traffic flows that need thorough engineering to ensure optimal network performance.

4.2 Multi-Tenancy and Isolation

Successful multi-tenancy depends on effective VNI assignment strategies and routing policies that support total isolation between tenant networks while preserving the operational simplicity and scalability features that provide EVPN with appeal for large-scale installations [8]. VNI provisioning must be in sync with the needs of tenants and administrative boundaries, cognizant of growth and future service development that might necessitate adjusting network segmentation plans as business needs change over time. Route targets and route distinguishers should be allocated systematically to support efficient route import and export behavior among tenant networks, with careful consideration of the hierarchical relationship between these identifiers and the impact on overall network performance. The use of strong multi-tenancy in VXLAN EVPN networks demands advanced policy systems that can impose tenant separation while supporting the flexibility and scalability advantages that make Ethernet VPN technology appropriate for heterogeneous deployment scenarios [8]. Contemporary deployments take advantage of automated provisioning platforms that can dynamically assign network resources to meet tenant needs, mitigating the operational complexity of having multiple autonomous network segments within common physical infrastructure. These automated systems need to have robust validation mechanisms in place to avoid configuration errors that may lead to compromised tenant isolation or introduce operational issues that affect network performance and stability. Sophisticated multi-tenancy deployments use the built-in strengths of Ethernet VPN technology to deliver scalable network virtualization solutions that can host hundreds of tenants with thousands of network endpoints spread across several geographic locations [8]. These solutions' scalability hinges on effective network resource utilization and prudent optimization of control plane functions so that multi-tenancy does not compromise the overall network performance or add operational complexity that has the potential to affect service delivery. Latest EVPN architectures enable the deployment of use cases with large tenant bases and multicasting requirements involving complexity, facilitated by advanced resource management methods that maximize network resource optimization across competing tenant needs.

5. Implementation Strategies and Configuration Patterns

Effective deployment of VXLAN EVPN adheres to standard configuration models that guarantee uniform behavior throughout the fabric, with

deployment strategies that take advantage of the doctrines of Network Functions Virtualization to attain enhanced operational efficiency and service agility [9]. VTEP configuration involves the definition of NVE interfaces, VNI-to-VLAN mappings, and BGP EVPN address family activation, which necessitates systematic methods of configuration management that align with NFV architectures intended to simplify operations and enhance service delivery capacities. Every VTEP needs router IDs that are distinct and correctly set up BGP neighbor relations with route reflectors or full mesh peering, with close attention to the scalability ramifications of various peering methods as fabric deployments scale to support more intricate service needs. The level of complexity in VXLAN EVPN deployments requires advanced automation and orchestration features that reflect the operational philosophy of Network Functions Virtualization, in which software-defined designs substitute conventional hardware-oriented configuration management [9]. Contemporary implementations take advantage of centralized configuration management systems that can dynamically create and deploy identical configurations to numerous network devices with the minimum amount of operations overhead while facilitating new tenant networks and services rapid deployment. These automated systems will need to have thorough validation procedures in place to guarantee that configuration updates preserve the operational advantages that NFV architectures are intended to deliver, such as enhanced service agility, lower time-to-market for new services, and better operational efficiency. Integrated Routing and Bridging (IRB) capability provides inter-subnet communication within tenant networks using distributed anycast gateway configurations that deliver optimal paths of forwarding while retaining the scalability advantage of distributed routing designs [10]. This solution avoids the necessity for centralized gateway appliances that would act as

bottlenecks within high-traffic environments, dispersing routing capabilities between multiple VTEPs to provide horizontal scaling of the gateway capacity. SVI interfaces assigned with the same IP addresses and virtual MAC addresses on multiple VTEPs provide transparent gateway capability, with advanced load balancing and failover functions that maintain ongoing service availability in the event of maintenance procedures or unplanned device failures. Distributed anycast gateway functionality requires meticulous planning of different IRB design models that have been created to solve particular deployment scenarios and operational needs [10]. High-end implementations use advanced algorithms to distribute gateway duties among accessible VTEPs, with dynamic load balancing features that can keep up with evolving traffic patterns and device availability. These design models need to take into consideration the rich interplay between routing and bridging functions such that forwarding behavior is consistent and predictable across failure scenarios and levels of traffic load. Quality of Service (QoS) policies need to be translated for overlay environments, taking into consideration both overlay and underlay marking strategies for traffic such that tenant service level agreements are maintained and efficient use of resources is provided across shared physical infrastructure [10]. DSCP marking strategies must retain tenant QoS requirements without allowing appropriate underlay network treatment of encapsulated traffic, necessitating advanced policy frameworks to translate tenant-specific QoS requirements into proper underlay forwarding behaviors. Hierarchical QoS policies in current IRB implementations can enforce tenant-specific bandwidth allocations and priority classes and guarantee important network control traffic appropriate treatment for fabric stability and convergence performance in various deployment scenarios.

Component	Function	Layer	Key Characteristics
Spine Switches	IP Forwarding Underlay Routing	Layer 3	High Bandwidth ECMP Support
Leaf Switches (VTEPs)	Encapsulation Decapsulation	Layer 2 + Layer 3	Tunnel Endpoints VNI Mapping
BGP EVPN Control Plane	Route Distribution MAC/IP Learning	Control Plane	Distributed Learning No Flooding Required
VXLAN Overlay	Virtual Networks Tenant Isolation	Overlay	16.7M Network Segments Multi-tenant Support
IP Underlay	Physical Connectivity VTEP Reachability	Underlay	OSPF or BGP Routing High Performance
VNI (24-bit Identifier)	Network Segmentation Tenant Identification	Data Plane	Unique per Tenant Scalable Addressing
Route Reflector	BGP Scalability Route Propagation	Control Plane	Reduces Peering Hierarchical Design
IRB (Integrated Routing)	Inter-subnet Routing Anycast Gateway	Layer 2 + Layer 3	Distributed Gateway Optimal Traffic Path

Figure 1. VXLAN EVPN Architecture Components [3, 4].

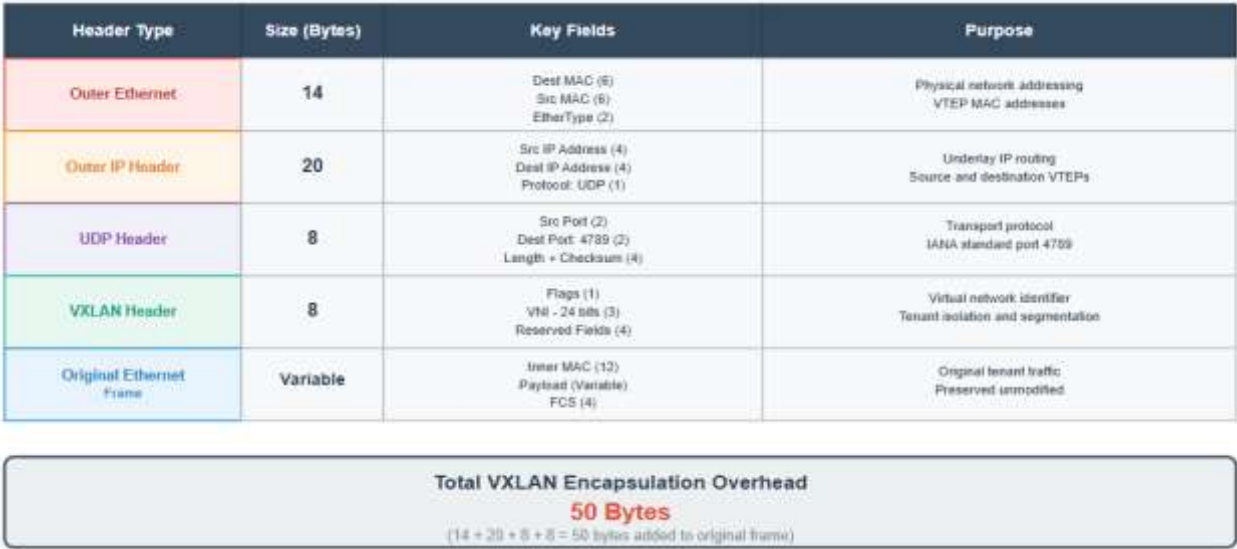


Figure 2. VXLAN Packet Encapsulation Headers [5, 6].

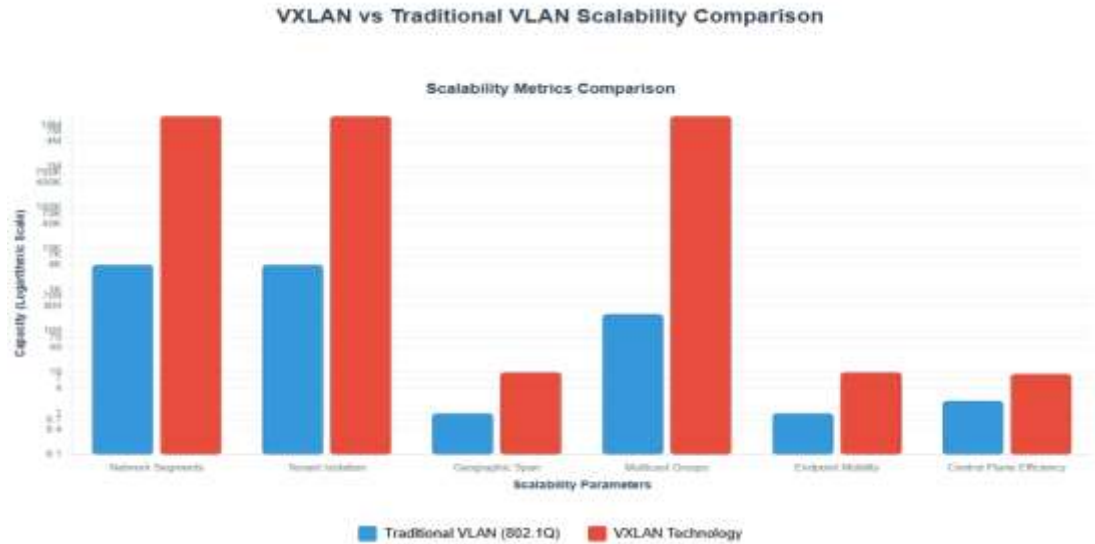


Figure 3. VXLAN vs Traditional VLAN Scalability Comparison [7, 8].

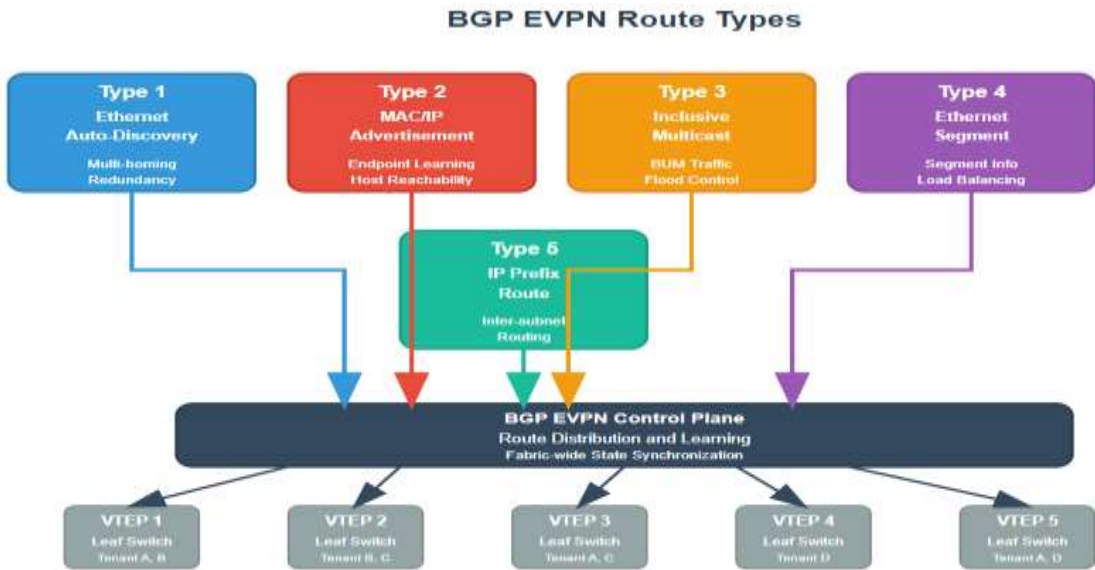


Figure 4. BGP EVPN Route Types and Functions [9, 10].

4. Conclusions

The implementation of VXLAN EVPN generation represents an essential shift in the statistics center networking structure, offering community operators a cutting-edge system for building scalable, flexible, and green network fabric that could adapt to the evolving demands of cutting-edge allotted computing environments. The combination of VXLAN overlay abilities with EVPN manipulates aircraft intelligence, creates an effective framework for implementing multi-tenant architectures that transcend the constraints of conventional layer 2 technology, at the same time, while keeping the operational simplicity and performance traits required for manufacturing environments. The allotted nature of BGP EVPN manipulates plane operations, mixed with hardware-extended information aircraft forwarding in present-day switching systems, permits deployment of large-scale fabrics with predictable performance characteristics and deterministic behavior across various traffic patterns and failure situations. The technology stack supports sophisticated multi-tenancy necessities through systematic VNI allocation techniques and complete routing policy frameworks that ensure complete isolation among tenant networks, even as allowing managed inter-tenant communication when required by way of application architectures. Superior implementations leverage programmable information aircraft skills and automated orchestration structures to reduce operational complexity whilst retaining the flexibility and scalability advantages that make VXLAN EVPN attractive for hyperscale deployments. The integration of routing and bridging capability through allotting anycast gateway configurations offers top-of-the-line forwarding paths for inter-subnet site visitors whilst disposing of single points of failure that could affect service availability. Excellent carrier guidelines adapted for overlay environments ensure that tenant provider-level agreements are preserved, even as they allow efficient utilization of shared bodily infrastructure resources across numerous deployment situations.

Author Statements:

- **Ethical approval:** The conducted research is not related to either human or animal use.
- **Conflict of interest:** The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper

- **Acknowledgement:** The authors declare that they have nobody or no-company to acknowledge.
- **Author contributions:** The authors declare that they have equal right on this paper.
- **Funding information:** The authors declare that there is no funding to be acknowledged.
- **Data availability statement:** The data that support the findings of this study are available on request from the corresponding author. The data are not publicly available due to privacy or ethical restrictions.

References

- [1] Kip Compton, "Cisco's Global Cloud Index Study: Acceleration of the Multicloud Era," Cisco, 2018. [Online]. Available: <https://blogs.cisco.com/news/acceleration-of-multicloud-era>
- [2] Derek Cheung, "BGP EVPN and All-Active Multi-Homing Ethernet Segment on GNS3," Medium, 2020. [Online]. Available: <https://derekcheung.medium.com/bgp-evpn-and-all-active-multi-homing-ethernet-segment-d797172cabf1>
- [3] Jason Reeves, "Understanding VXLAN: Virtual Extensible LAN Explained," FiberMall, 2024. [Online]. Available: <https://www.fibermall.com/blog/vxlan.htm?srltid=AfmBOoom3bCBBSQzQNzLn8tWLB2bcGjzcLUu9qEkiAyAvS6cJSE0qWpn>
- [4] Hassib Fowad Rezai, "MPLS VPN In a Service Provider Network," Bachelor of Engineering Information and Communication Technology, 2021. [Online]. Available: https://www.theseus.fi/bitstream/handle/10024/356060/Rezai_Hassib.pdf?sequence=2
- [5] Arjun Singh et al., "Jupiter Rising: A Decade of Clos Topologies and Centralized Control in Google's Datacenter Network," SIGCOMM, 2015. [Online]. Available: <https://conferences.sigcomm.org/sigcomm/2015/pdf/papers/p183.pdf>
- [6] Zijun Hang et al., "Programming Protocol-Independent Packet Processors High-Level Programming (P4HLP): Towards Unified High-Level Programming for a Commodity Programmable Switch," MDPI, 2019. [Online]. Available: <https://www.mdpi.com/2079-9292/8/9/958>
- [7] Alexey Andreyev et al., "Running Border Gateway Protocol in large-scale data centers," Engineering at Meta, 2021. [Online]. Available: <https://engineering.fb.com/2021/05/13/data-center-engineering/bgp/>
- [8] Piotr Jakacki, "Ethernet VPN explained: benefits, use cases, and key concepts," Codilime, 2023. [Online]. Available: <https://codilime.com/blog/ethernet-vpn-benefits-use-cases-key-concepts/>

- [9] Alepo, "Network Functions Virtualization: Basics to Benefits," 2020. [Online]. Available: <https://www.alepo.com/network-functions-virtualization-nfv-basics-to-benefits/>
- [10] ipSpace, "Integrated Routing and Bridging (IRB) Design Models," 2022. [Online]. Available: <https://blog.ipspace.net/2022/11/irb-design-models/>