

BWCAMPUSNETZ

Zukunftsfähige Konzepte für die Campusnetze
an Universitäten und Hochschulen

Architekturbausteine für moderne und virtualisierte Netze

Federführung bei der Erstellung dieses Dokuments:
Karlsruher Institut für Technologie, Universität Stuttgart, Universität Ulm
Kontakt: team@bwcampusnetz.de

Inhalt

1	Einleitung	5
2	Zielpublikum	6
3	Was macht Campusnetze speziell?	7
4	Motivation	8
5	Netzwerkarchitekturen	9
5.1	Traditionell	9
5.1.1	Collapsed Core/Backbone	9
5.1.2	Three-Tier Architektur	9
5.1.3	Technologien zur Netzwerksegmentierung	10
5.2	EVPN/VXLAN	10
5.2.1	Kernnetz auf Basis von EVPN/VXLAN	12
5.2.2	Distribution-Layer	14
6	Anwendungsfälle für EVPN/VXLAN	15
6.1	Arten der Anbindung	15
6.1.1	L2-Trunk	15
6.1.2	Routed Port	16
6.1.3	VTEP in Fabric	16
6.2	Anbindung Firewall	17
6.2.1	L3-Anbindung einer Firewall an einem Borderleaf	17
6.2.2	Layer 2 Anbindung	20
6.2.3	Firewall als Teil der Fabric	21
6.3	Anbindung WLAN und VPN	21
6.3.1	Local Breakout vs Centralized im WLAN	21
6.3.2	Anbindung WLAN-Controller (bei Centralized Breakout) und VPN-Server	24
6.3.3	L2 Trunk	25
6.3.4	L3 Routed Port	26

6.3.5	VTEP auf WLAN-/VPN-Controller	27
7	Fazit	29
8	Versionsverlauf	30
9	Glossar	31

Abbildungsverzeichnis

5.1	Beispiel einer typischen, symmetrischen Leaf/Spine-Architektur	12
5.2	Beispiel einer flexiblen Topologie auf Basis von EVPN/VXLAN: Leaf/Spine-Layer mit angebundenem Access-Layer, der keinen topologischen Einschränkungen unterliegt	13
6.1	L3-Anbindung einer Firewall an einem Borderleaf	17
6.2	Anbindung einer Firewall: Ingress (In) und Egress (Out) VRFs	19
6.3	Anbindung einer Firewall: Layer 2	21
6.4	WLAN: Central als auch den Local Switching Modus	23
6.5	VLAN-Override mittels AAA-Server	23
6.6	Local Breakout in einem virtualisierten Netz	24
6.7	Anbindung zentraler WLAN-Controller/VPN-Servern: Layer-2-Trunk	25
6.8	Anbindung zentraler WLAN-Controller/VPN-Servern: L3 Routed Port . . .	26
6.9	Anbindung zentraler WLAN-Controller/VPN-Servern: VTEP auf WLAN-/VPN-Controller	27

1 Einleitung

Mit dem zunehmenden Bedarf an Skalierbarkeit, Flexibilität und Sicherheit stoßen klassische, stark Layer-2-zentrierte Netzdesigns zunehmend an ihre Grenzen. Große Broadcast-Domänen, mangelnde Flexibilität bei der Netzsegmentierung und eingeschränkte Kontrolle über den Datenverkehr erschweren den Betrieb moderner Netzwerke - insbesondere in Campusnetzen, in denen oftmals Datacenter-Netze mit Anforderung aus der Virtualisierung oder dem High Performance Computing mit klassischen Access-Netzen kombiniert sind. Ethernet VPN (EVPN) hat sich als zukunftsweisende Technologie etabliert, um diesen Herausforderungen zu begegnen. EVPN ermöglicht eine konsistente Bereitstellung von Layer-2- und Layer-3-Diensten über ein gemeinsam genutztes IP Underlay und eröffnet neue Möglichkeiten für ein modulares und service-orientiertes Netzwerkdesign. In diesem Papier werden zentrale EVPN-Konzepte vorgestellt und aufgezeigt, wie sich mit ihrer Hilfe Firewalls, WLAN-Controller und andere Netzwerkdienste flexibel, skalierbar und sicher in moderne Architekturen integrieren lassen.

2 Zielpublikum

Dieses Dokument stellt ein Kompendium verschiedener Architekturbausteine dar, die gezielt dazu verwendet werden können, einzelne Aspekte der Netzarchitektur sowie die gesamte Netzstruktur einer Institution auf Grundlage aktueller Anforderungen zu modernisieren. Die vorgestellten Architekturbausteine sind als Module zu verstehen, die einerseits eine in sich schlüssige und moderne Architektur ergeben, andererseits aber auch isoliert eingesetzt und in ein bestehendes lokales Netzwerk integriert werden können. Die Bausteine beschäftigen sich mit den Bereichen Kernnetz, Datacenter-Netz, WLAN und Firewall.

Die Betrachtung der einzelnen Module erfolgt aus der Perspektive eines Architekten. Es werden grundlegende Gestaltungsmöglichkeiten aufgezeigt, jedoch werden ausdrücklich keine herstellerepezifischen Konfigurationsanleitungen erstellt, obwohl einzelne Abschnitte beispielhafte Konfigurationen enthalten können.

Ein besonderes Augenmerk wird auf die Interaktion der verschiedenen Komponenten und Module gelegt. Ziel ist es zu verdeutlichen, wie einzelne Module nahtlos in ein bestehendes Netzdesign integriert werden können.

Das übergeordnete Ziel dieses Dokuments besteht darin, einen umfassenden Überblick über die Bausteine eines modernen und resilienten Netzdesigns zu bieten, der den Leser zur eigenständigen Weiterbildung anregen soll.

3 Was macht Campusnetze speziell?

Campusnetzwerke an Universitäten und HAWen unterscheiden sich in vielerlei Hinsicht von klassischen Datacenter- oder Unternehmens-Netzwerken. Eine zentrale Besonderheit liegt darin, dass im akademischen Kontext häufig sowohl der klassische Campusbereich als auch eigene Rechenzentren innerhalb eines gemeinsamen Netzwerks betrieben werden. Dadurch müssen diese Netzwerke eine Vielzahl verschiedener Nutzungsszenarien abdecken – von alltäglichen Anwendungen in Forschung, Lehre und Verwaltung bis hin zu hochspezialisierten, datenintensiven Forschungsprojekten. Diese Vielfalt bringt besondere Anforderungen an die Netzarchitektur mit sich, denn das Netzwerk muss ein breites Spektrum an Anforderungen hinsichtlich Durchsatz, Flexibilität und Sicherheit erfüllen.

Im Vergleich zu typischen Unternehmensnetzwerken ist auch die räumliche Ausdehnung eines Campusnetzes oft deutlich größer: Universitätsgelände erstrecken sich nicht selten über weitläufige Flächen oder sogar mehrere Standorte oder Stadtteile, was zusätzliche Herausforderungen an die Netzplanung und -infrastruktur stellt. Hinzu kommt, dass nicht nur im Kernnetz, sondern auch im Access, beispielsweise in Forschungslaboren oder in Hörsälen mit moderner Medientechnik, hohe Bandbreiten benötigt werden. An vielen Standorten betreiben zudem einzelne Organisationseinheiten eigene lokale Rechenzentren für spezielle Anwendungen, sodass die Anforderungen an ein Datacenter-Netzwerk auch im Campusumfeld erfüllt werden müssen.

Ein weiterer Aspekt ist die Freiheit von Forschung und Lehre: Forschende und Lehrende benötigen ein offenes, innovatives und anpassungsfähiges Netzwerk, das ihnen erlaubt, neue Technologien und Anwendungen zu erproben und zu nutzen. Die hohe Dynamik der Forschungslandschaft führt zusätzlich dazu, dass sich die Anforderungen an das Netzwerk ständig ändern – neue Projekte entstehen, Arbeitsgruppen wechseln, etc. Gleichzeitig ist das Spektrum der Nutzenden sehr breit gefächert: Von Studierenden und Lehrenden ohne tiefere IT-Kenntnisse bis hin zu erfahrenen IT-Administratoren muss das Netzwerk für alle effektiv nutzbar sein.

Schließlich ist die Organisation der IT an Universitäten und Hochschulen oft dezentral aufgebaut: Verschiedene Fachbereiche oder Institute betreiben eigene IT-Infrastrukturen, die möglichst nahtlos in das Campusnetzwerk integriert werden müssen.

4 Motivation

Traditionelle Netzwerke basieren häufig auf bewährten Modellen wie Collapsed-Core- oder Three-Tier-Architekturen. Diese Architekturen bieten klare Strukturen mit definierten Schichten, bringen jedoch signifikante Herausforderungen mit sich. Insbesondere der Einsatz des Spanning Tree Protocols (STP) zur Vermeidung von Schleifen führt zu ineffizienter Nutzung der verfügbaren Bandbreite und langsamen Konvergenzzeiten bei Netzwerkänderungen. Darüber hinaus bieten klassische Mechanismen wie HSRP und VRRP nur eingeschränkte Möglichkeiten zur Lastverteilung und begrenzen damit die Effizienz der Infrastruktur. Ein weiteres wesentliches Defizit traditioneller Netzwerke liegt in der begrenzten Skalierbarkeit und Flexibilität bei der Integration moderner Technologien, insbesondere im Kontext von Cloud- und SDN-Umgebungen. Zudem stellt die logische Segmentierung von Netzwerken mittels Technologien wie VRF Lite oder MPLS/VPLS oft komplexe und nur eingeschränkt skalierbare Lösungen dar, welche die Flexibilität einer integrierten Layer-2- und Layer-3-Umgebung vermissen lassen.

5 Netzwerkarchitekturen

5.1 Traditionell

5.1.1 Collapsed Core/Backbone

Das Collapsed-Core-Netzwerkdesign ist ein Layer-2-zentriertes Modell, das sich durch eine zentrale Konzentration wesentlicher Funktionen wie Routing im Core-Bereich auszeichnet. Dieser Kern ist bewusst klein und intelligent gehalten. Durch die zentrale Steuerung entstehen Vorteile in der Implementierung und Betriebskomplexität, da weniger Protokolle und Technologien erforderlich sind. Die Nutzung des Spanning Tree Protocols sorgt für Redundanz, erzeugt jedoch gleichzeitig ineffiziente Bandbreitennutzung, da redundante Pfade blockiert werden. Zudem treten lange Konvergenzzeiten bei Netzänderungen auf, was die Stabilität negativ beeinflusst. Fehlersuche gestaltet sich in dieser Architektur schwierig, da die Topologie statisch ist und Fehler oft mehrere Geräte betreffen können.

5.1.2 Three-Tier Architektur

Die Three-Tier-Architektur hingegen gliedert das Netzwerk in drei klar abgegrenzte Schichten: Core-, Distribution- und Access-Layer. Der Core-Layer stellt einen schnellen, zuverlässigen Backbone zur Verfügung, während der Distribution-Layer als zentrale Vermittlungsinstanz dient, die Routing- und Policy-Aufgaben übernimmt. Der Access-Layer verbindet direkt die Endgeräte und bildet typischerweise die Layer-2-Verbindung zum Distribution-Layer. Eine Alternative stellt der „Routed Edge“ dar, bei dem bereits im Access-Layer Layer-3-Konnektivität vorhanden ist. Ein wesentlicher Nachteil dieser Architektur besteht darin, dass durch die Layer-3-Trennung zwischen den Distribution-Layer Routern (oder falls es sich um ein „Routed Edge Design“ handelt auch zwischen Access- und Distribution-Layer) VLANs nicht über das gesamte Netzwerk hinweg ausgedehnt werden können. Subnetze sind auf bestimmte Bereiche beschränkt und lassen sich nicht flexibel auf jeden beliebigen Access-Switch ausrollen. Vorteile dieser Architektur sind klare Strukturen, bessere Skalierbarkeit und einfachere Segmentierungsmöglichkeiten. Trotzdem bleibt auch hier das grundlegende

Problem bestehen, dass redundante Pfade durch Spanning Tree oft blockiert werden, was zu ineffizienter Nutzung der Bandbreite führt. Ebenso sind Mechanismen zur Lastverteilung wie HSRP oder VRRP eingeschränkt, da sie primär auf einer Active/Backup-Logik basieren und somit keine echte Lastverteilung ermöglichen.

5.1.3 Technologien zur Netzwerksegmentierung

Darüber hinaus erfolgt in traditionellen Netzwerken oft eine logische Trennung von Routinginstanzen, üblicherweise Virtual Routing and Forwarding (VRF) genannt. VRF Lite stellt hier zwar eine einfache Möglichkeit zur Segmentierung dar, skaliert jedoch nur begrenzt, da auf jedem beteiligten Gerät mit Subinterfaces gearbeitet werden muss. Um diesem Skalierbarkeitsproblem zu begegnen, kommt häufig MPLS (Multiprotocol Label Switching) zum Einsatz, welches eine effiziente Layer-3-VPN-Segmentierung ermöglicht. Alternativ bietet Virtual Private LAN Service (VPLS) eine reine Layer-2-VPN-Konnektivität. Beide Technologien bieten zwar flexible Segmentierungsmöglichkeiten, sind jedoch jeweils auf Layer-2 oder Layer-3 beschränkt und integrieren diese Ebenen nicht effizient miteinander.

Alternative Ansätze wie LISP (Locator/ID Separation Protocol) oder Shortest Path Bridging (SPB) bieten theoretisch Lösungen für einige dieser Herausforderungen, konnten sich aber bislang aufgrund fehlender breiter Herstellerunterstützung und Interoperabilität nicht flächendeckend durchsetzen.

5.2 EVPN/VXLAN

VXLAN (Virtual Extensible LAN) in Kombination mit EVPN (Ethernet VPN) bildet eine moderne, standardisierte Lösung für Layer-2- (L2VNI) und Layer-3-VPNs (L3VNI), die sich besonders für den Einsatz in akademischen Campusnetzen eignet. Diese Netzwerkvirtualisierungstechnologie verfolgt das Ziel, flexible und hochverfügbare Netze auf Basis eines gerouteten IP-Underlay-Netzwerks bereitzustellen. Eine einzelne EVPN/VXLAN-Installation wird dabei oft als Fabric bezeichnet.

Das Overlay-Netzwerk, in dem produktive Netze und Dienste betrieben werden, basiert auf sogenannten VTEP (VXLAN Tunnel Endpoints). Diese Endpunkte übernehmen die Kapselung und Entkapselung der Datenpakete und sorgen dafür, dass die Kommunikation zwischen

den Geräten im Overlay über das Underlay-Netzwerk zu den jeweils relevanten VTEPs erfolgt. Die für den Betrieb eines EVPN-basierten Netzes erforderlichen Informationen werden dabei über das BGP-Protokoll verteilt, das nicht nur zur Verteilung von IP-Routen dient, sondern auch die Synchronisation von MAC-Adressen und weiteren Netzwerkparametern übernimmt.

Ein wesentlicher Vorteil von EVPN/VXLAN besteht darin, dass klassische, auf Layer 2 basierende Protokolle wie das Spanning Tree Protocol (STP) weitgehend obsolet werden. Während STP in traditionellen Netzen zwar Schleifen verhindert, jedoch die verfügbare Bandbreite limitiert und zu längeren Konvergenzzeiten führen kann, setzt EVPN auf moderne Routing-Mechanismen. Durch den Einsatz von Equal-Cost Multi-Path (ECMP) im Underlay können sämtliche Verbindungen parallel genutzt werden, was die Effizienz und Ausfallsicherheit des Netzwerks verbessert.

Darüber hinaus stellt EVPN mit Multi-Homing eine leistungsfähige Alternative zu klassischen, oft komplexen Lösungen wie MLAG (proprietär) bereit. EVPN ermöglicht echtes Aktiv/Aktiv-Multihoming über beliebige Kombinationen von Upstream-Geräten hinweg, wodurch sowohl Redundanz als auch Lastverteilung optimiert werden.

Ein weiteres zentrales Element der EVPN/VXLAN-Technologie ist das Anycast-Gateway. Damit kann jeder Router, der eine Broadcast-Domain bereitstellt, gleichzeitig als IP-Gateway für diese Domain fungieren. Der Datenverkehr wird somit möglichst früh geroutet, was die Auswirkungen eines Komponentenausfalls auf das übrige Netzwerk minimiert und das Routing nicht auf wenigen Komponenten zentralisiert.

EVPN/VXLAN zeichnet sich durch eine hohe Flexibilität aus und kann je nach Anforderungen und bestehender Infrastruktur in unterschiedlichen Architekturen umgesetzt werden. Für kleinere Campusnetze bietet sich beispielsweise ein zentrales Core-Design an, bei dem Routing und Bridging zentral gesteuert werden. In größeren Umgebungen, besonders im Datacenter, empfiehlt sich eine Leaf/Spine-Architektur, das eine hohe Skalierbarkeit ermöglicht. Auch hybride Ansätze sind realisierbar, bei denen bestehende Infrastrukturen schrittweise migriert und durch moderne EVPN/VXLAN-Komponenten ergänzt werden. Diese Architekturoptionen werden im folgenden Kapitel näher erläutert.

5.2.1 Kernnetz auf Basis von EVPN/VXLAN

5.2.1.1 Datacenter: Leaf/Spine-Architektur

In Rechenzentren hat sich die Leaf/Spine-Topologie (Abbildung 5.1) als Standard etabliert, die eine symmetrische Netzwerkstruktur mit vorhersagbaren Leistungswerten ermöglicht. Jeder Leaf-Switch ist mit allen Spine-Switches verbunden, wodurch eine gleichmäßige Lastverteilung und minimale Hop-Distanzen erreicht werden können. Die physische Verkabelung folgt hier einem klar definierten Muster, was im kontrollierten Datacenter-Umfeld deutlich einfacher umzusetzen ist als in WAN- oder Campusnetzen, wo die Verfügbarkeit von Glasfaserverbindungen der limitierende Faktor ist. Die Leaf-Switches fungieren als VTEPs, während die Spine-Ebene rein als IP-Transitlayer fungiert. Diese Trennung reduziert die MAC-Tabellenbelastung auf die Leaf-Ebene.

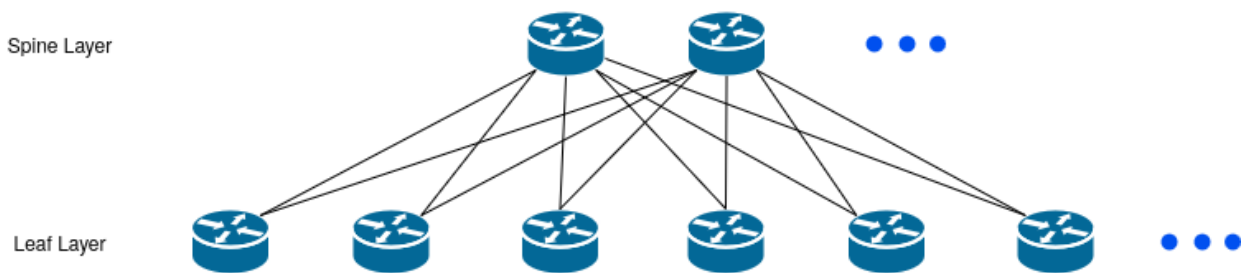


Abb. 5.1: Beispiel einer typischen, symmetrischen Leaf/Spine-Architektur

5.2.1.2 Campus-Netzwerke

In Campusumgebungen ermöglicht EVPN/VXLAN flexible Topologien jenseits starrer Baumstrukturen. Ringarchitekturen oder kleinere Spine-Layer erlauben redundante Pfade, bei denen der Ausfall einzelner Links oder Geräte nicht zu einem weitreichenden Ausfall führt. Im Gegensatz zu traditionellen Spanning-Tree-basierten Lösungen bleiben dabei alle physischen Verbindungen aktiv nutzbar, was die Gesamtbandbreite erhöht und Latenz reduziert. Der Einsatz von EVPN/VXLAN im Campus bietet sich besonders an, wenn sowohl auf Layer 2, als auch auf Layer 3 eine hohe Flexibilität benötigt wird. So können Broadcastdomänen (L2VPN) oder Routing-Kontexte (VRFs, L3VPN) flexibel verfügbar gemacht werden, und dabei die Vorteile von EVPN/VXLAN gegenüber klassischen Layer-2-zentrischen Netzen gewahrt werden. Die Access-Geräte fungieren dabei als VTEPs; analog zu Leaves in einer Leaf/Spine-Architektur. Abbildung 5.2 zeigt ein Beispiel einer flexiblen Netztopologie.

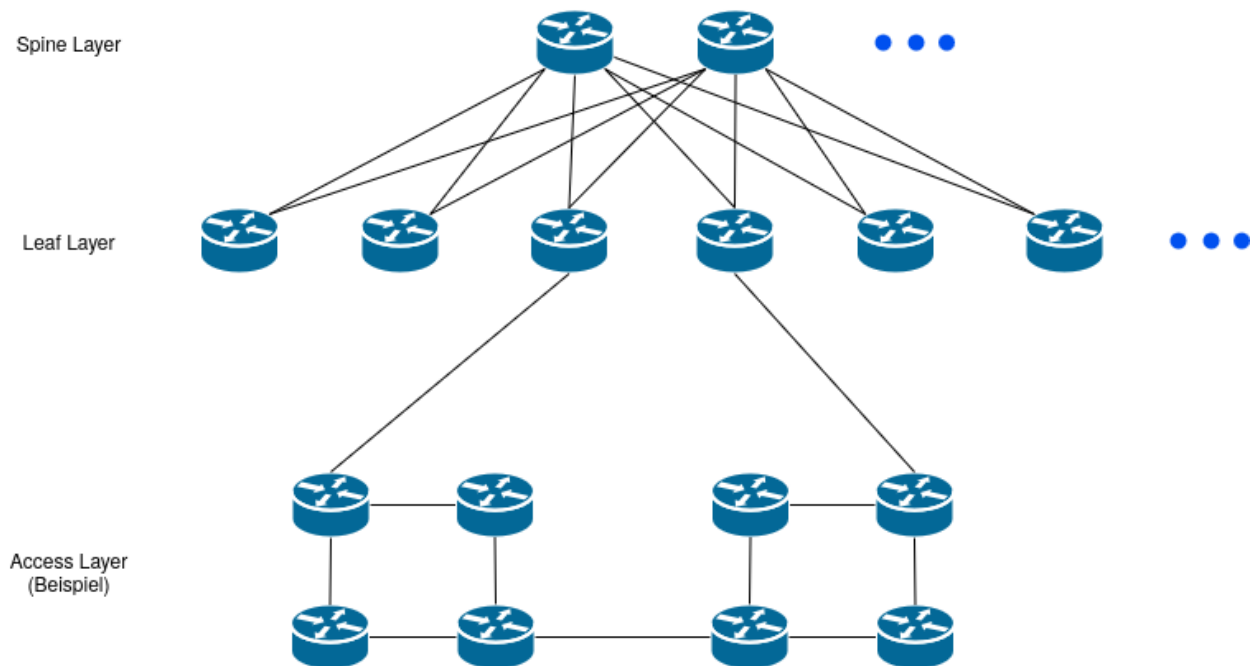


Abb. 5.2: Beispiel einer flexiblen Topologie auf Basis von EVPN/VXLAN: Leaf/Spine-Layer mit angebundenem Access-Layer, der keinen topologischen Einschränkungen unterliegt

5.2.1.3 Netzwerk-Virtualisierung oberhalb bestehender Produktivnetze

EVPN/VXLAN kann auch als Overlay auf beliebigen, gerouteten Underlay-Netzen betrieben werden. Dies ermöglicht es, L2- oder L3-Dienste flexibel an Stellen im Netzwerk bereitzustellen, an denen dies auf Grund der Netzwerk-Architektur im Normalfall nicht möglich ist (z.B. in rein gerouteten Netzen). Außerdem kann dieser Ansatz als Einstieg für eine schrittweise Modernisierung des Netzes hin zu EVPN/VXLAN verwendet werden. Dabei übernimmt das Underlay reinen IP-Transport, während das Overlay mandantenfähige Layer-2/Layer-3-Services bereitstellt.

Ein Vorteil dieser Lösung liegt in der Kompatibilität mit Legacy-Systemen: Traditionelle VLANs können über Hardware-VTEPs in das VXLAN-Overlay integriert werden. Gleichzeitig ermöglicht die Architektur die Einführung moderner Features wie Anycast-Gateways oder ARP/ND-Suppression, selbst wenn das Underlay-Netz nur grundlegende Routing-Funktionalitäten bietet.

5.2.2 Distribution-Layer

Je nach Größe des Netzes kann dieses über einen Distribution-Layer verfügen, also eine Verteilebene für die Anbindung der Geräte im Campus bereitstellen. Beim Einsatz von EVPN/VXLAN im Kernnetz kann der Distribution-Layer klassisch via Layer-2-Übergabe (LAGs mit VLAN-Trunks) angebunden werden. Alternativ kann der Distribution-Layer selbst am EVPN/VXLAN teilnehmen und als VTEP fungieren. Die Geräte, an welche der Distribution-Layer im Kernnetz physisch angebunden wird, können dabei selbst als VTEP fungieren oder ausschließlich am Underlay teilnehmen und Transport für das Overlay bereitstellen.

In Campus-Netzwerken kann bei größeren Installationen ein Distribution-Layer als Zwischenebene zwischen Access-Layer und Kernnetz implementiert werden. Die Struktur und Anbindung des Distribution-Layers variiert je nach Netzwerkgröße und -architektur.

Bei EVPN/VXLAN-basierten Kernnetzen bieten sich zwei Hauptansätze für den Distribution-Layer an:

5.2.2.1 Klassische Layer-2-Anbindung

Die Übergabe zwischen Kernnetz und Distribution-Layer kann via Link Aggregation Groups (LAGs) mit VLAN-Trunks hergestellt werden. Geräte im Distribution-Layer arbeiten rein auf Layer 2.

5.2.2.2 Aktive Teilnahme am Overlay

Der Distribution-Layer kann als VTEP (Virtual Tunnel Endpoint) fungieren und direkt am EVPN/VXLAN-Overlay partizipieren. In diesem Fall übernimmt er Paketkapselung/-entkapselung und kommuniziert mit anderen VTEPs über das IP-Underlay.

Die Upstream-Geräte im Kernnetz können dabei entweder selbst als VTEPs agieren und L2/L3-Dienste an angeschlossene Systeme bereitstellen oder rein transitiv im Underlay operieren und nur den Overlay-Datenverkehr transportieren.

6 Anwendungsfälle für EVPN/VXLAN

6.1 Arten der Anbindung

Nachgelagerte Systeme und netznahe Dienste können auf unterschiedliche Weise an eine EVPN/VXLAN-Fabric angebunden werden, wobei sich je nach Anforderungen und Systemumgebung verschiedene Ansätze etabliert haben.

6.1.1 L2-Trunk

Die einfachste und am weitesten verbreitete Methode ist die Terminierung an der Fabric mittels eines L2-Trunks. Hierbei werden auf einem physischen oder logischen Port, beispielsweise einem LAG, mehrere Broadcast-Domains über verschiedene VLAN-Tags gebündelt und an das nachgelagerte System übergeben. Diese Vorgehensweise bietet eine hohe Kompatibilität, da viele bestehende Systeme ausschließlich auf Layer-2-Basis arbeiten. Im Regelfall wird auf Seiten der Fabric für jedes übergebene VLAN ein Routing-Endpunkt, etwa in Form eines Switch Virtual Interfaces (SVI), bereitgestellt, um IP-Konnektivität herzustellen. Leistungsfähigere Systeme differenzieren dabei häufig zwischen Bridge-Domain und VLAN-Tag, was zusätzliche Flexibilität ermöglicht, allerdings vor allem im ISP-Umfeld verbreitet ist. Im Campus- und Datacenter-Umfeld wird hingegen meist direkt im VLAN terminiert. Die nachgelagerten Systeme erhalten Adressen in den jeweiligen Broadcast-Domains und sind so in den übergebenen Netzen erreichbar. Die einzelnen VLANs können auf dem vorgelagerten Router in unterschiedlichen Routing-Kontexten, sogenannten VRFs, enden, was eine Segmentierung auf Layer 3 ermöglicht. Vorteil dieses Ansatzes ist die hohe Kompatibilität mit nachgelagerten Systemen. Multihoming erfordert hier entweder EVPN Multi-Homing oder eine proprietäre MLAG-Lösung, deren Einsatz oft mit großer Komplexität auf Seiten der Router verbunden ist.

6.1.2 Routed Port

Eine alternative Möglichkeit besteht darin, die Anbindung über einen gerouteten Port am Overlay zu realisieren. Hierbei wird ein eigenes Transfernetz direkt auf dem physischen oder logischen Interface konfiguriert und einem bestimmten Routing-Kontext zugeordnet. Das angebundene System kann in diesem Transfernetz dynamische Routing-Protokolle wie BGP nutzen, um die gewünschte IP-Konnektivität herzustellen. Diese Methode bietet insbesondere Vorteile beim Multihoming, da die angeschlossenen Systeme nicht auf komplexe Layer-2-Multihoming-Technologien wie MLAG angewiesen sind. Stattdessen können sie einfach mehrere Layer-3-Verbindungen zu verschiedenen Routern aufbauen, sofern der entsprechende Routing-Kontext auf diesen Routern zur Verfügung steht. Dies erhöht die Flexibilität und vereinfacht das Design, insbesondere in größeren Umgebungen mit hohen Anforderungen an Redundanz und Skalierbarkeit.

6.1.3 VTEP in Fabric

Darüber hinaus besteht die Möglichkeit, netznahe Dienste direkt als VTEP in die Fabric zu integrieren. In diesem Fall wird das System, vorzugsweise über einen gerouteten Port, mit dem Underlay der EVPN/VXLAN-Fabric verbunden und nimmt aktiv mittels BGP an der EVPN Control Plane teil. Das System terminiert dabei VXLAN-Tunnel und kann sowohl Layer-2- als auch Layer-3-Konnektivität in Form von Broadcast-Domains oder Routing-Kontexten (VRFs) bereitstellen. Dieser Ansatz ermöglicht es, eine große Anzahl von Netzsegmenten effizient zu terminieren, ohne die Konfiguration der vorgelagerten Router anpassen zu müssen. Besonders vorteilhaft ist dies in Szenarien, in denen Dienste mit vielen verschiedenen Netzsegmenten interagieren müssen, wie es beispielsweise bei VPN- oder Firewall-Lösungen der Fall ist. Eine Herausforderung dieses Ansatzes stellen die hohen Anforderungen hinsichtlich der Integration in eine bestehende EVPN/VXLAN-Fabric dar: Hier muss sichergestellt werden, dass sich verschiedenste Parameter im Kontext des Betriebs der EVPN/VXLAN-Fabric zwischen virtuellen VTEPS und den Hardware-VTEPS decken.

6.2 Anbindung Firewall

Eine Firewall zwischen Netzsegmenten kann in einer EVPN/VXLAN Fabric auf verschiedene Arten angebunden werden. Prinzipiell zu unterscheiden ist eine Anbindung, bei der Netze tatsächlich auf einer Firewall terminieren und eine Anbindung, bei der eine Firewall lediglich als L3 Hop fungiert. Auch der topologische Ort der Anbindung ist ein wichtiges Unterscheidungsmerkmal: Ist die Firewall Teil der Fabric oder nicht? Die folgenden Anbindungen sollen im folgenden betrachtet werden:

- L3 an Borderleaf
- Anbindung via klassischem Trunk
- Firewall als Teil der Fabric

6.2.1 L3-Anbindung einer Firewall an einem Borderleaf

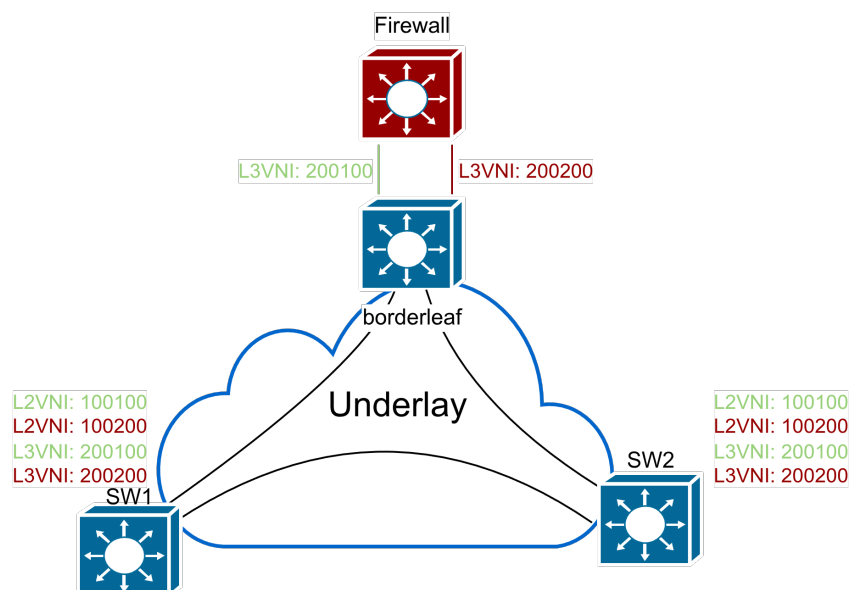


Abb. 6.1: L3-Anbindung einer Firewall an einem Borderleaf

In diesem Szenario wird eine Firewall zwischen zwei Netzsegmenten an einem Border-Leaf angebunden und selbst nicht Teil der Fabric. In einem klassischen Netzdesign werden verschiedene Segmente oftmals in verschiedenen VRF-Instanzen (Virtual Routing and Forwarding, kurz VRFs) separiert. In jeder VRF existiert eine eigene Routingtabelle. Damit ein

Host von einer VRF in eine andere VRF kommunizieren kann, muss eine Verbindung der VRFs geschaffen werden. Im einfachsten Fall, können Routen der einen in die andere VRF geleakt werden. Das ermöglicht ein Routing auf dem selben Gerät. Möchte man nun jedoch, dass Traffic zwischen einzelnen VRFs über eine Firewall geleitet wird, so bietet es sich an, Routinginformationen über diese Firewall auszutauschen, sodass diese jeweils der Next-Hop zwischen den VRFs ist. So kann beispielsweise die Firewall mit jeweils einem Port an jede der beiden VRFs mit einem Linknetz angebunden werden. Routen können dann über ein dynamisches Routingprotokoll ausgetauscht werden. In einer EVPN/VXLAN Fabric kann der gleiche Ansatz gewählt werden. Innerhalb der Fabric werden VRFs durch sogenannte L3VNIs abgebildet. An einem Borderleaf kann nun die Firewall so angebunden werden, dass sie mit jeweils einem Linknetz in einer L3VNI verbunden ist. Danach ist der Ansatz exakt identisch zum klassischen. In Abbildung 6.1 ist eine Fabric zu sehen, an die eine Firewall an ein Borderleaf angebunden wurde. Dabei gehört jeweils ein Link zu einer VRF. Beiden VRFs sind in der Fabric jeweils eine L3VNI zugeordnet. Die jeweiligen L2VNIs die der VRF zugehörig sind, können nun durch die Router der Fabric geroutet werden (bspw. als Anycast-Gateway) und sind dann auf jedem beteiligten Router verfügbar. Zu beachten ist, dass die Linknetze zur Firewall nicht notwendigerweise physikalische L3-Interfaces sein müssen. Hier kann ein Trunk benutzt werden, um mehrere VRF Linknetze auf einem physikalischen Interface zu transportieren. Im obigen Beispiel gilt, dass die klassischen VRFs „red“ und „green“ durch die folgenden VNIs repräsentiert werden:

VRF	VNI	L2VNI/L3VNI
Green	200100	L3VNI
Red	200200	L3VNI

Die zugehörigen L2VNIs sind jeweils farblich markiert. Der Vorteil einer Anbindung auf diese Art ist, dass die Regelsätze allein basierend auf einem L3/L4-Tupel formuliert werden können. Darüber hinaus terminiert kein Subnetz direkt auf der Firewall. Man ist also nicht darauf angewiesen, dass die Firewall bzw. deren Betriebssystem eine gute Implementierung typischer Edge-Funktionalitäten aufweist. Das ermöglicht mehr Flexibilität bei der Auswahl der Komponenten. Dieses Konzept entspricht am ehesten der klassischen Perimeter-Firewall zwischen zwei Netzsegmenten. Sollte man eine große Zahl Netzsegmente voneinander isolieren wollen - bspw. für Mikrosegmentierung - stößt man schnell an Grenzen. So ist die Anbindung als L3 Hop kann bei einer großen Anzahl an VRFs sehr aufwändig werden. So müssen Linknetze zu jeder VRF an die Firewall angebunden werden. Unabhängig davon ob dies physikalisch oder mittels Trunk geschieht, erfordert dieses Vorgehen einen hohen Grad

an Automatisierung. Die folgenden Ansätze können benutzt werden, um die Skalierbarkeit zu verbessern.

6.2.1.1 Policy Based Routing

Hierbei gibt es zwei VRFs: Eine public VRF und eine firewalled VRF. Die Firewall dient als L3 Hop zwischen beiden VRFs. Dabei kann sie als Kontrollpunkt für Traffic zwischen den zwei VRFs dienen, allerdings per se keinen Traffic innerhalb der VRF kontrollieren. Das liegt daran, dass innerhalb einer VRF direkt anliegende Routen lokal immer präferiert werden. Um dieses Problem zu lösen, kann mit Policy-based Routing (PBR) gearbeitet werden. Dabei wird mittels einer Regel der next-hop innerhalb der VRF von „local“ auf eine Adresse der Firewall umgeschrieben. So wird intra-VRF Traffic auch immer zur Firewall geschickt, die diesen dann behandeln kann. Dieser Ansatz erfordert die Implementierung von PBR-Regelsätzen auf allen L3-Interfaces, die dem Client zugewandt sind.

6.2.1.2 Ingress und Egress VRFs

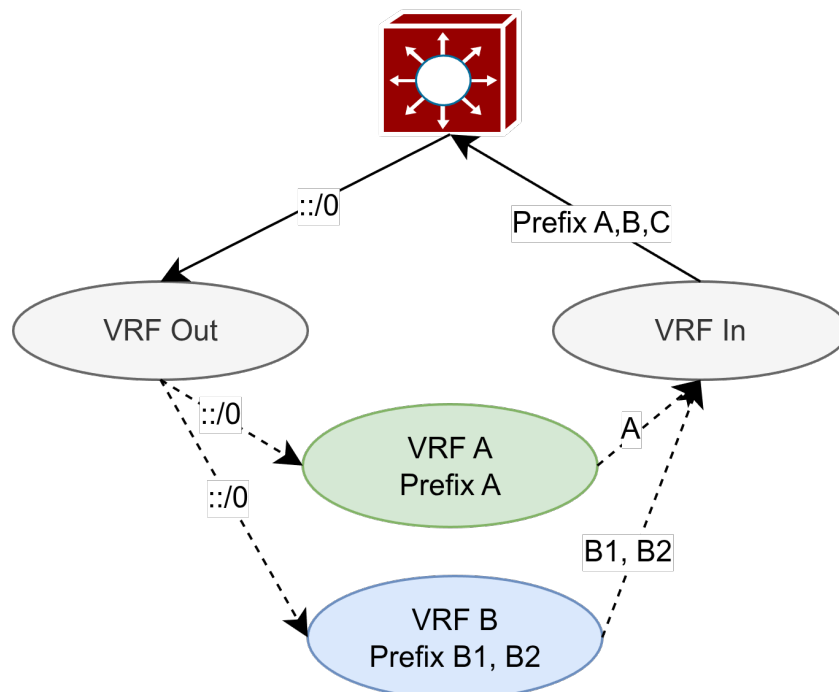


Abb. 6.2: Anbindung einer Firewall: Ingress (In) und Egress (Out) VRFs

Ein weiterer Ansatz ist, jeden Tenant in seiner eigenen VRF zu isolieren und dabei durch Route-Leaking zu verhindern, dass zu jeder VRF dedizierte Linknetze benötigt werden. Dabei werden zwei Transit-VRFs angelegt. Die Firewall annouciert eine default-route in die Outgoing VRF. Aus dieser VRF wird nun in alle Tenant-VRFs ein route-leak für die default-route konfiguriert. Die Routen für alle Prefixes in den jeweiligen Tenant-VRFs werden in die Incoming VRFs geleakt. Von dort werden die Routen per BGP an die Firewall annouciert. Will VRF A nun mit VRF B kommunizieren, so geht der Traffic zunächst in die Outgoing VRF, da die default-route dort hin zeigt. Von dort wird der Traffic zur Firewall geroutet. Der Rückweg erfolgt über die Ingress VRF. Das Routing ist also asymmetrisch, wodurch eine direkte Kommunikation zweier VRFs ohne Firewall verhindert wird.

6.2.1.3 Implikationen für den Firewall-Regelsatz

Die meisten Firewall Implementierungen verwendet First-Match Semantik. Sobald eine Regel ein Paket behandelt hat, werden keine weiteren Regeln geprüft. In einem Ansatz, in dem mehrere Tenants durch eine Firewall geschützt werden, muss daher Wert auf die korrekte semantische Isolation von Regelsätzen gelegt werden. Ansonsten könnte ein Outgoing Regelsatz der VRF A, der eine Klausel „Allow All“ enthält, fälschlicherweise Traffic zu VRF B erlauben.

6.2.2 Layer 2 Anbindung

Die EVPN/VXLAN Fabric kann auch ausschließlich zum Transport von Layer-2 Frames benutzt werden. Dann entspricht die Funktionalität in etwa der eines klassischen, VLAN-basierten Ansatzes. Hierzu werden L2VNIs am Port, an dem die Firewall angebunden werden soll, zu regulären VLAN-tagged Paketen umgesetzt. Beispielsweise:

VLAN	VNI	L2VNI/L3VNI
100	100100	L2VNI
200	100200	L2VNI

Nun kann die Firewall selbst als Abschluss der Netze fungieren. Dabei hat die Fabric keinerlei Kenntnis mehr über Routinginformationen für die jeweiligen Netze.

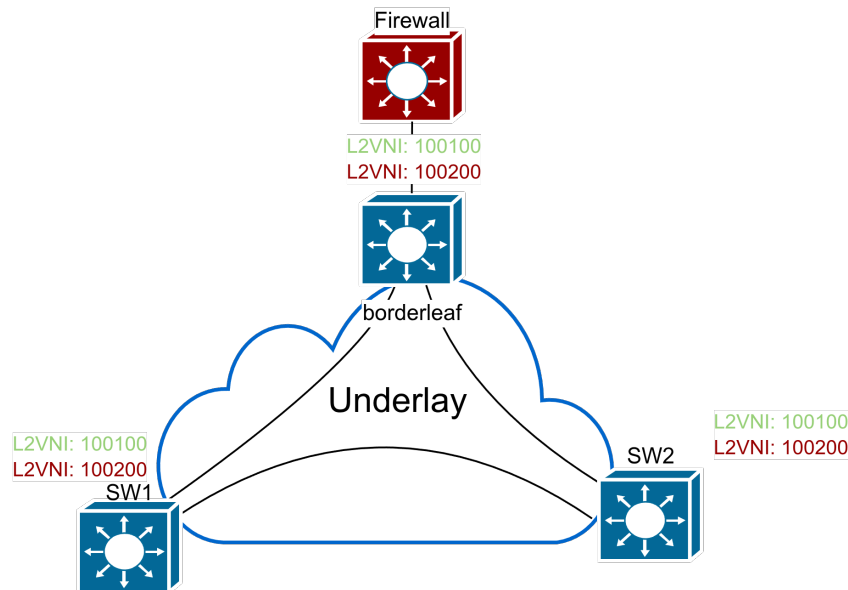


Abb. 6.3: Anbindung einer Firewall: Layer 2

6.2.3 Firewall als Teil der Fabric

Bei diesem Ansatz wird die Firewall selbst Teil der Fabric, also bspw. zu einem VTEP. Das heißt, dass keine externe Anbindung mehr wie in Fall 1 erfolgen muss, sondern das Routing direkt lokal auf der Firewall bzw. dezentral in der Fabric erfolgen kann. Dieser Ansatz ist der flexibelste, da das Routing zwischen L3VNIs nun direkt innerhalb der Fabric erfolgen kann. Jedoch ist bisher kein Hersteller bekannt, der diese Funktionalität bietet.

6.3 Anbindung WLAN und VPN

6.3.1 Local Breakout vs Centralized im WLAN

In drahtlosen Netzwerken gibt es zwei grundlegende Ansätze für die Verarbeitung des Datenverkehrs: Centralized Switching (tunnelbasiert) und Local Switching (lokaler Breakout). Je nach Architektur und Anwendungsfall kann entweder der gesamte Datenverkehr über einen zentralen Controller geleitet oder direkt am Access Point (AP) ins lokale Netzwerk ausgeleitet werden.

Centralized Switching ist die klassische Methode, bei der der gesamte Datenverkehr der Access Points über einen Tunnel (bspw. per CAPWAP oder GRE) an den Wireless LAN Controller (WLC) gesendet wird. Dort erfolgt die Verarbeitung, einschließlich Sicherheitsrichtlinien, Quality of Service (QoS) und Routing.

Local Switching (Local Breakout) hingegen erlaubt es den APs, den Datenverkehr direkt ins lokale Netzwerk zu leiten, ohne den Umweg über einen zentralen Controller. Eine Technologie, die diese Funktionalität auf Cisco-Geräten ermöglicht, ist FlexConnect (ehemals H-REAP). FlexConnect erlaubt es APs, sowohl im zentralisierten Modus als auch im Standalone-Modus zu arbeiten. Ist die Verbindung zum WLC aktiv, kann der Datenverkehr entweder weiter über den Controller geleitet oder lokal geschwitcht werden. Falls die Verbindung zum WLC unterbrochen wird, können die APs autonom weiterarbeiten und den lokalen Breakout weiterhin ermöglichen.

Abbildung 6.4 veranschaulicht dabei sowohl den Central als auch den Local Switching Modus. Im Folgenden wird dabei beispielhaft jeder SSID ein bestimmtes Default VLAN (bzw. Subnetz) zugeordnet, welchem Nutzer nach dem Verbinden mit der SSID zugeordnet werden. Dabei wird die SSID „SSID1“, zentral am Controller geschwitched, wohingegen die SSIDs „SSID2“ und „SSID3“ lokal am Access Point geschwitched werden.

Zusätzlich zum Default VLAN können Clients per VLAN-Override vom AAA-Server einem anderen VLAN zugeordnet werden (siehe Abbildung 6.5).

Im Weiteren ist Layer-2-Roaming im Falle von Local Switching nur möglich, wenn sich die Access Points in derselben Broadcast-Domäne befinden. Im vorliegenden Beispiel muss der Client eine neue IP-Adresse beziehen, da das Roaming auf Layer 3 stattfindet, was zu einer Unterbrechung der Verbindung führt.

Werden wie im Beispiel oben unterschiedliche VLAN-IDs verwendet, muss die Zuordnung zum entsprechenden VLAN ortsabhängig, anhand des Access Points oder des Distribution Routers, erfolgen.

Betreibt man allerdings entweder ein traditionelles Netz entsprechend eines Collapsed Core oder ein modernes Netz, wie in Kapitel 5.2.1 vorgestellt, hat man die Möglichkeit campusübergreifend Layer-2 Konnektivität bereitzustellen, wie in Abbildung 6.6 gezeigt. Damit muss der gesamte WLAN-Verkehr nicht zwingend über einen zentralen Controller getunnelt werden muss, wodurch ebenfalls potenzielle Engpässe bei der Bandbreite vermieden werden.

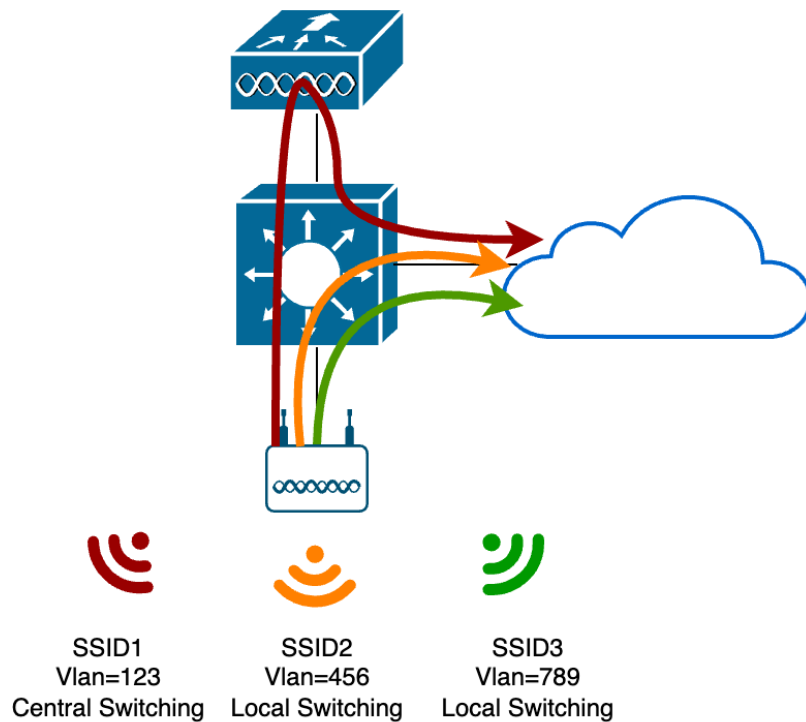


Abb. 6.4: WLAN: Central als auch den Local Switching Modus

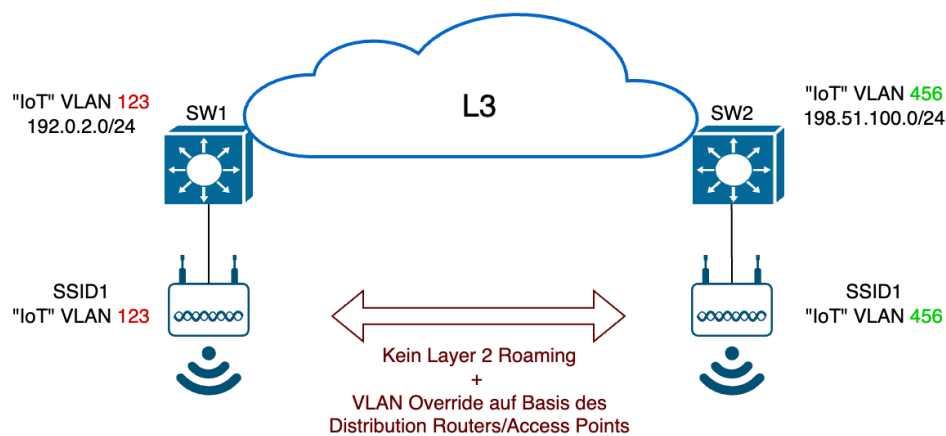


Abb. 6.5: VLAN-Override mittels AAA-Server

Beim Einsatz einer EVPN/VXLAN-Fabric sollte beachtet werden, dass durch MAC-Mobility, also dem Wandern von Clients zwischen verschiedenen Access Points bei Nutzung von Local Breakout, eine erhöhte Last innerhalb der Fabric entstehen kann. Dies liegt daran, dass sich die MAC-Adressen der Clients bei einem Wechsel des Access Points zwischen den beteiligten VTEPs (VXLAN Tunnel Endpoints) verschieben. Jeder dieser sogenannten MAC-Moves löst in der Regel ein Update im BGP-Protokoll aus, das für die Verteilung der MAC-Adressen

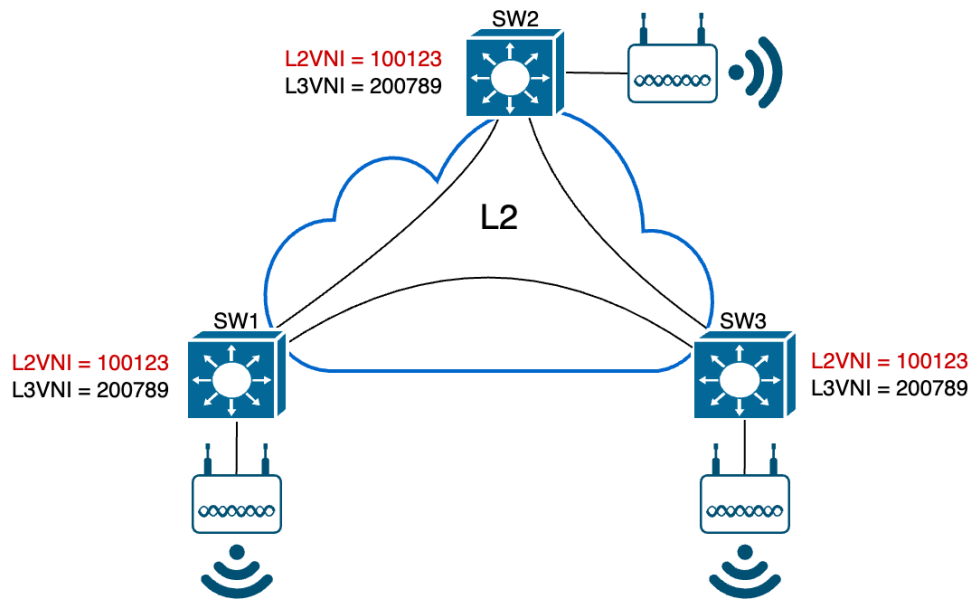


Abb. 6.6: Local Breakout in einem virtualisierten Netz

in der Fabric verantwortlich ist. Insbesondere in Umgebungen mit vielen mobilen Clients, wie etwa in drahtlosen Netzwerken, kann dies zu einer hohen Frequenz an BGP-Updates führen. Die Folge ist eine erhöhte Belastung der Control Plane, da jeder MAC-Move von allen beteiligten Switches verarbeitet werden muss. Zusätzlich ist zu berücksichtigen, dass es eine gewisse Latenz gibt, bis die Information über den neuen Aufenthaltsort einer MAC-Adresse in der gesamten Fabric verteilt ist. Während dieser Zeit kann es passieren, dass Datenpakete zunächst noch an den alten VTEP gesendet werden, was zu kurzzeitigen Paketverlusten führen kann.

Im Folgenden werden die Möglichkeiten der Anbindung eines WLCs beschrieben, für den Fall des Centralized Switchings, bei dem sämtlicher Datenverkehr zum Controller getunnelt und dort verarbeitet wird.

6.3.2 Anbindung WLAN-Controller (bei Centralized Breakout) und VPN-Server

Es gibt insgesamt drei Möglichkeiten, den WLC sowie VPN-Controller an die EVPN-Fabric anzubinden. Einerseits die Anbindung über einen Layer-2-Trunk oder Layer-3 Routed Port

zwischen einem der Leafs sowie über einen VTEP auf dem WLC- und/oder VPN-Controller selbst, womit dieser effektiv Teil der Fabric wird.

6.3.3 L2 Trunk

Die wohl gängigste Art ist es, Controller über einen Layer-2-Trunk direkt mit einem Leaf-Switch oder Leaf-Switch Paar zu verbinden. Die SVIs (Switched Virtual Interfaces), also die Default Gateways der jeweiligen Broadcast-Domänen, werden dabei auf den Leaf-Switches konfiguriert.

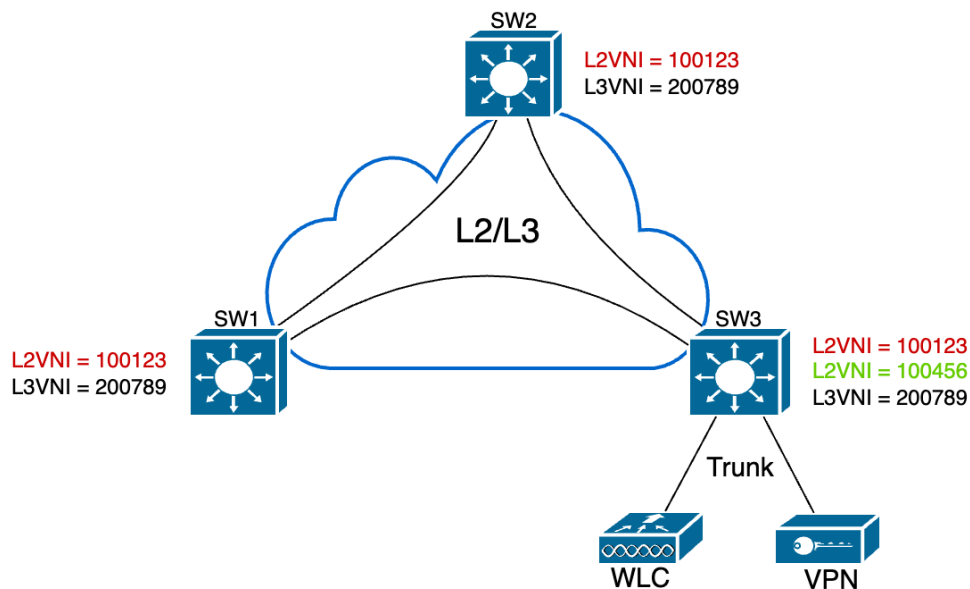


Abb. 6.7: Anbindung zentraler WLAN-Controller/VPN-Servern: Layer-2-Trunk

VLAN	VNI	L2VNI/L3VNI
123	100123	L2VNI
456	100456	L2VNI
789	100789	L3VNI

Wenn der Controller per Trunk an den EVPN-Leaf angebunden wird, ermöglicht dies einerseits das Layer-2-Stretching innerhalb der Fabric. Dies wird im Folgenden anhand des L2VNIs 100123 veranschaulicht, wobei als L3VNI 200789 verwendet wird. Damit besteht die Möglichkeit, die Broadcast-Domäne, welcher das L2VNI 100123 zugeordnet wurde, über die Switches SW1, SW2 und SW3 zu stretchen, wobei der WLC/VPN-Controller an SW3 angebunden ist. Dadurch kann Kunden campusübergreifend eine gemeinsame Broadcast-Domäne

im LAN sowie WLAN/VPN bereitgestellt werden. In diesem Szenario werden innerhalb der EVPN-Fabric sowohl Type-2-Routen (MAC/IP Routen) als auch Type-5-Routen geadvertised, sodass die Endgeräte nahtlos in der gleichen Layer-2-Domäne kommunizieren können.

Alternativ dazu kann der Ansatz eines sogenannten Schwesternetzes umgesetzt werden. Dabei werden Broadcast-Domänen nicht über den Campus zum Controller gestretched (LAN <-> WLAN/VPN), sondern stattdessen lediglich Layer-3-Konnektivität zwischen LAN und WLAN/VPN bereitgestellt. Dies geschieht, indem der Broadcast-Domäne bzw. dem VLAN im WLAN/VPN ein anderes L2VNI – in diesem Fall grafisch visualisiert mittels L2VNI 100456 – zugeordnet wird. In dieser Variante werden keine Type-2-Routen an entsprechende Leafs innerhalb der Fabric verteilt. Die Kommunikation zwischen den verschiedenen L2VNIs erfolgt über das zugeordnete L3VNI, wobei die EVPN-Fabric in diesem Fall Type-5-Routen (IP-Präfixe) advertised, um die Layer-3-Konnektivität zwischen den Segmenten sicherzustellen.

6.3.4 L3 Routed Port

Als Alternative zur Anbindung per Trunk kann der WLC bzw. VPN-Controller, ähnlich zur Firewall in Kapitel 6.2.1, per Layer-3-Routed-Port angebunden werden. Dabei erfolgt das Routing zu den WLAN/VPN-Netzen entweder statisch, über ein Interior Gateway Protocol (z.B. OSPF) oder mittels BGP. Die entsprechenden Routen werden in die EVPN-Fabric als Type-5-Routen importiert.

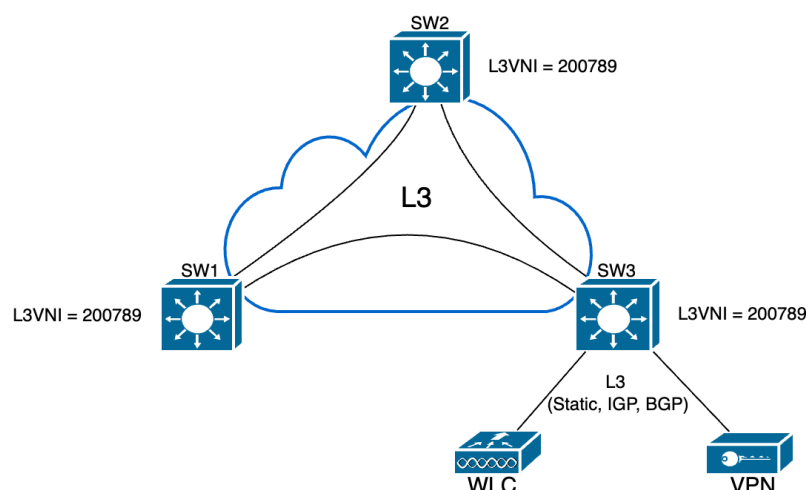


Abb. 6.8: Anbindung zentraler WLAN-Controller/VPN-Servern: L3 Routed Port

VLAN	VNI	L2VNI/L3VNI
789	100789	L3VNI

6.3.5 VTEP auf WLAN-/VPN-Controller

Der Controller könnte ebenso als Teil der EVPN-Fabric konfiguriert werden, was den direkten Import von Type-2-Routen ermöglicht. Im Folgenden Schaubild wird der Controller als erweiterter Teil der Fabric dargestellt. Dabei wird beispielhaft das L2VNI 100123 sowie die dazugehörige Broadcastdomäne über den gesamten Campus gestretched.

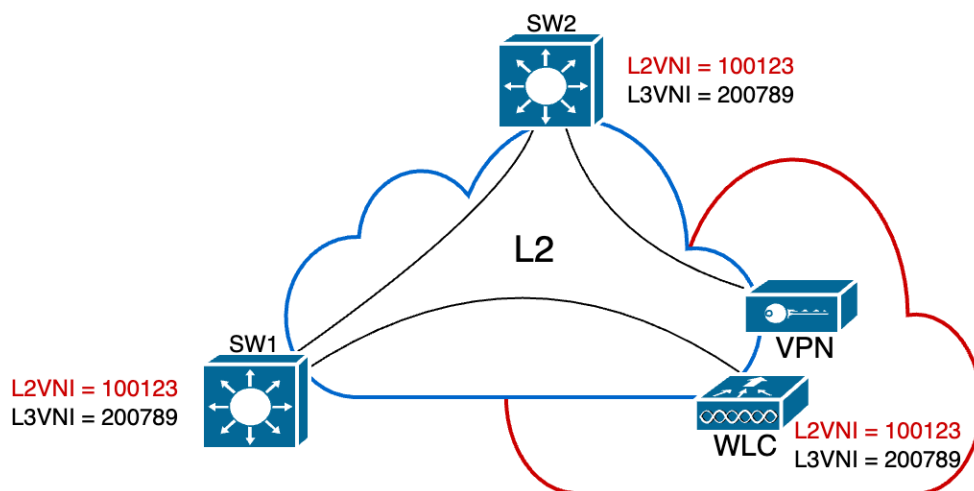


Abb. 6.9: Anbindung zentraler WLAN-Controller/VPN-Servern: VTEP auf WLAN-/VPN-Controller

VLAN	VNI	L2VNI/L3VNI
123	100123	L2VNI
789	100789	L3VNI

Solch eine Konfiguration bietet wesentliche Vorteile, insbesondere die Optimierung des Ost-West-Traffics stellt einen wesentlichen Vorteil dar, indem VXLAN es ermöglicht, virtuelle Maschinen oder Container, die sich zwar auf dem gleichen physischen Host, aber in unterschiedlichen VLANs oder Subnetzen befinden, virtuell in einer gemeinsamen Layer-2-Domäne zu gruppieren. Ohne VXLAN wäre diese direkte Kommunikation nur innerhalb desselben VLANs möglich; jeglicher Traffic zwischen unterschiedlichen VLANs müsste über externe Netzwerkgeräte, wie etwa Leaf-Switches oder Router, geleitet werden. Durch die VXLAN-Enkapsulierung verbleiben die Pakete auf dem Host, was die Kommunikation beschleunigt,

die Latenz verringert und die Bandbreitennutzung erheblich verbessert. Zudem erfolgt die ARP- und Neighbor-Discovery-Suppression direkt am Host, was Broadcasts und Multicasts erheblich reduziert und unnötigen Overhead im Netzwerk vermeidet. Ein weiterer Vorteil liegt in der Möglichkeit, standardisierte Routingprotokolle für das Underlay-Netzwerk zu verwenden, wodurch ein effektives ECMP-Loadbalancing zur optimalen Lastverteilung ermöglicht wird.

Auf der anderen Seite bestehen bei dieser Konfiguration auch einige Herausforderungen. Ein zentraler Nachteil ist die mangelnde Unterstützung durch viele Hersteller, da VXLAN/EVPN bisher noch selten Bestandteil existierender WLC- oder VPN-Portfolios ist. Zudem dominieren proprietäre Lösungen den WLAN-Markt, und es mangelt an geeigneten Open-Source-Lösungen, die die Anforderungen eines Unternehmens-WLANs erfüllen könnten. Darüber hinaus steigen die Anforderungen an die Hardware erheblich, da VXLAN-Encapsulierung und Routing direkt auf dem Host durchgeführt werden müssen. Besonders ohne Hardware-Offloading können hierdurch erhebliche Performance-Einbußen auftreten. Nicht zuletzt steigt der Betriebsaufwand durch die komplexere Konfiguration und Verwaltung erheblich an, da neben der eigentlichen VPN-Funktionalität auch der Routing-Daemon auf dem Host konfiguriert und betrieben werden muss. Dies erfordert leistungsfähige Automatisierungs- und Orchestrierungslösungen, um die Administration effizient zu gestalten.

7 Fazit

Die Anforderungen moderner Netzwerke, insbesondere hinsichtlich Skalierbarkeit und Flexibilität, lassen sich mit traditionellen, Layer-2-zentrierten Architekturen zunehmend schwer erfüllen. Klassische Ansätze wie STP, HSRP oder VRF Lite stoßen an ihre Grenzen, wenn es darum geht, performante, dynamisch anpassbare und serviceorientierte Infrastrukturen bereitzustellen. Insbesondere in Campusumgebungen, die eine Vielzahl unterschiedlicher Anforderungen vereinen, zeigen sich die Schwächen konventioneller Netzdesigns besonders deutlich.

Wie in diesem Konzeptpapier gezeigt wurde, bietet Ethernet VPN hier einen zukunftsweisen- den Lösungsansatz. Durch die Trennung von Control- und Data-Plane sowie die einheitliche Unterstützung von Layer-2- und Layer-3-Services über ein IP-Backbone schafft EVPN die Voraussetzungen für ein modernes, modulares Netzwerkdesign. Es ermöglicht eine effiziente Nutzung von Ressourcen, verbessert die Lastverteilung und erlaubt eine flexible Integration von Netzwerkdiensten wie Firewalls oder WLAN-Controllern.

EVPN erweist sich somit nicht nur als technologisches Upgrade, sondern als notwendiger Entwicklungsschritt für Netzwerke, die den Anforderungen heutiger und zukünftiger Anwendungen gerecht werden wollen.

8 Versionsverlauf

Version	Datum	Änderungen
1.0	6.10.2025	Initiale Veröffentlichung

9 Glossar

EVPN Ethernet VPN. Moderne Netzwerkvirtualisierungstechnologie.

IP Internet Protocol.

WLAN Wireless Local Area Network.

STP Spanning Tree Protocol. Layer-2-Technologie zur Herstellung von Schleifenfreiheit und Redundanten in Layer-2-Netzen.

HSRP Hot Standby Router Protocol. Von Cisco entwickeltes, offenes Protokoll für Gateway-Redundanz. Ähnliche Funktionalität via VRRP.

VRRP Virtual Router Redundancy Protocol. Offenes Protokoll für Gateway-Redundanz. Ähnliche Funktionalität via HSRP.

SDN Software-Defined Networking. Sammelbegriff für verschiedene Konzepte moderner Netze, insbesondere Trennung von Control-Plane und Data-Plane.

VRF Lite Virtual Routing and Forwarding Lite:- Funktionalität, welche auf Routern mehrere lokale Routing-Kontexte bereitstellt (eigene Routing-Tabellen).

MPLS Multiprotocol Label Switching. Einfache Forwarding-Plane auf Basis lokal gültiger Label an Paketen, oft einfacher und günstiger als eine IP-Dataplane.

VPLS Virtual Private LAN Service. Protokoll zur Herstellung virtueller Layer-2-Sgmente auf Basis von MPLS.

VRF Virtual Routing and Forwarding. Ermöglicht die Implementierung von mehreren virtualisierten Routing-Kontexten über eine zugrundeliegende Fabric (MPLS L3VPN / EVPN L3VNI).

LISP Locator/ID Separation Protocol. Protokoll zur Trennung von Identität (ID) und Lokalisierung (Locator) im IP-Routing, um Mobilität und Skalierbarkeit zu verbessern.

SPB Shortest Path Bridging:- Offene Netzwerkvirtualisierungstechnologie basierend auf Ethernet und IS-IS.

IS-IS Intermediate System to Intermediate System. Alternatives IGP vergleichbar zu OSPF.

VXLAN Virtual Extensible LAN. Tunnel-Protokoll mit 24-Bit ID-Nummernraum oberhalb von UDP/IP.

L2VNI Layer 2 VXLAN Network Identifier:- Im Kontext von EVPN ein virtualisiertes Layer-2-Segment (Broadcastdomain) mit eindeutigem Identifier.

L2VPN Layer 2 Virtual Private Network:- Oberbegriff für Netzwerkvirtualisierungstechnologien, welche das Aufspannen von Layer-2-Segmenten (Broadcastdomains) über mehrere Geräte hinweg ermöglichen.

L3VNI Layer 3 VXLAN Network Identifier:- Im Kontext von EVPN ein virtuelles Layer-3-Segment (VRF) mit eindeutigem Identifier, wessen Routingtabellen über mehrere Geräte hinweg synchronisiert werden.

L3VPN Layer 3 Virtual Private Network:- Oberbegriff für Netzwerkvirtualisierungstechnologien, welche die Implementierung von VRFs ermöglichen.

IGP Interior Gateway Protocol. Oberbegriff für Routing-Protokolle, die innerhalb eines Netzwerks verwendet werden.

VTEP VXLAN Tunnel Endpoint. Geräte oder virtuelle Funktionen, die VXLAN-Tunnel-Endpunkte darstellen; sie kapseln und entkapseln VXLAN-Pakete.

ECMP Equal-Cost Multi-Path. Routingtechnologie, welche eine Lastverteilung über mehrere gleichwertige Pfade erlaubt.

MLAG Multi-Chassis Link Aggregation Group. Oberbegriff und Name verschiedener Technologien, welche die Bildung von LAGs über mehrere separate Geräte hinweg zwecks Redundanz ermöglicht. Von den meisten Herstellern proprietär implementiert.

Leaf/Spine-Architektur Netzwerktopologie, welches Netzwerkgeräte, an die Endgeräte angeschlossen werden (Leaves) mittels einem symmetrischen Spine-Layer verbinden. Dabei gibt es keine Querverbindungen zwischen Leaves oder Spines.

WAN Wide Area Network. Weitverkehrsnetzwerk.

ARP/ND-Suppression Technologie zur Reduktion von Multicast/Broadcast-Datenverkehr in Broadcastdomains, indem ARP/ND-Anfragen durch das nächstgelegene Gateway beantwortet werden, statt durch das Netzwerk geflutet zu werden.

LAG Link Aggregation Group. Technologie zur Bündelung mehrerer Verbindungen auf Layer 2 zur Herstellung von Redundanz und Erhöhung der Bandbreite.

VLAN-Trunk Verbindungen, die mehrere virtuelle LANs (VLANs) über einen einzelnen Link transportieren, gewöhnlich mittels 802.1Q-Tagging.

SVI Switch Virtual Interfaces. Virtualles Interface auf Routern, welches Konnektivität in virtuellen Netzwerksegmenten (VLANs) bereitstellt.

ISP Internet Service Provider. Internetanbieter.

VLAN Virtual Local Area Network. Logisches Teilnetz identifiziert durch eine VLAN-ID zwischen 1 und 4095.

BGP Border Gateway Protocol. Flexibles, policy-basiertes Routing-Protokoll für verschiedene Anwendungsfälle wie Internetrouting, strukturiertes Routing innerhalb eines Netzwerks sowie zum Transport von Daten für Overlay-Netzwerke.

PBR Policy Based Routing. Regelbasiertes-Routing auf Basis weiterer Merkmale als die Ziel-Adresse von Paketen.

VPN Virtual Private Network. Begriff beschreibt im Rahmen von Netzwerkvirtualisierungstechnologien die Bildung virtueller Layer-2- und Layer-3-Segmente, im weiteren Sinne auch Zugangstechnologien, welche den Zugang zum Netzwerk über das Internet ermöglichen.

AP Access Point. Netzwerkgerät, das kabellosen Zugang zum Netzwerk via WLAN bereitstellt.

CAPWAP Control And Provisioning of Wireless Access Points. Protokoll für die Steuerung und Anbindung von Access Points durch zentrale Controller mittels Tunneln.

GRE Generic Routing Encapsulation. Tunneling-Protokoll oberhalb von IP.

WLC Wireless LAN Controller. Meist proprietäre Komponente des Netzwerks, welches Access Points via Tunnel-Technologien zentral anbindet.

QoS Quality of Service. Technologie zur Priorisierung von Netzwerkverkehr.

H-REAP Hybrid Remote Edge Access Point. Access Point, der trotz zentralem Controller ein lokales Breakout bereitstellt.

OSPF Open Shortest Path First. Dynamisches IGP zur Etablierung effizienter Routing-Pfade innerhalb eines Netzwerks.