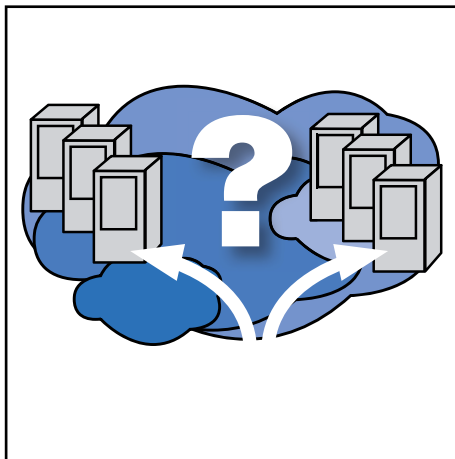


Routing im, vom und zum RZ

von Dr. Behrooz Moayeri

In den letzten Jahren wurde sehr häufig und sehr intensiv darüber diskutiert, wie eine skalierbare Layer-2-Netzstruktur für Rechenzentren aussehen soll. Diese Diskussion war angesichts der Herausforderungen, die durch die Servervirtualisierung verursacht werden, notwendig und wichtig.

Der Autor kann sich jedoch des Eindrucks nicht erwehren, dass wegen der Diskussion über Layer 2 ein anderer wichtiger Aspekt beim Design von RZ-Netzen in Vergessenheit geraten ist, nämlich das Routing im, vom und zum RZ. Oft kommt es dazu, dass sich Planer genau dazu wenig Gedanken gemacht haben, bis der Tag X kommt, an dem das Netz implementiert werden muss.



Häufig werden dann die Layer-3-Strukturen im und um das RZ irgendwie - vielleicht wie es gerade passt - aufgebaut. Das kann aber nicht im Sinne einer vorausschauenden Planung sein.

Warum Routing?

Fangen wir mit der grundsätzlichen Frage an, warum im Zusammenhang mit Rechenzentren Routing überhaupt noch erforderlich ist. Machen neue Layer-2-Verfahren das Routing im Zusammenhang mit RZ-Netzen nicht überflüssig?

weiter auf Seite 14

Zweitthema

Neue Protokolle im RZ: Funktionsumfang – Potential – Auswirkung auf die Infrastruktur

von Dipl.-Math. Cornelius Höchel-Winter

Reichen bestehende Netzwerktechnologien aus, um den Bedarf virtueller und verteilter Rechenzentren abzudecken? Die Idee, die in diesem Artikel diskutiert wird, ist die Verbindung von virtuellen Maschinen standortneutral in Form virtueller Layer-2-Netzwerke über beliebige dazwischen liegende Infrastrukturen hinweg.

Dahinter verbirgt sich die Frage, ob es überhaupt notwendig ist, dass virtuelle Maschinen zum Beispiel als Teil einer

skalierenden Web-Applikation unabhängig von ihrem Standort Layer-2-Verbindungen untereinander pflegen müssen. Hinzu kommt, dass in Zukunft im Rahmen einer fortschreitenden Automatisierung virtuellen Maschinen automatisch Ressourcen zugewiesen werden und damit auch automatisch initiierte Wanderungen von VMs verbunden sein können. Gleichzeitig werden Rechenzentren und Ressourcen immer mehr verteilt aufgesetzt, sei es nur über Brandschutzgrenzen oder sogar über Standortgrenzen hinweg.

Natürlich ist das Thema komplex, eben weil der Aufbau verteilter Infrastrukturen mit Zugang zu Datenbanken und Speichersystemen zwischen verteilten Rechenzentren alles andere als einfach und allgemeingültig umzusetzen ist. Trotzdem gibt es einen schnell zunehmenden Bedarf für diese Art von Architektur, weil eben Web-Anwendungen für große und stark schwankende Teilnehmerzahlen über die dynamische Generierung von virtuellen Maschinen zur Laufzeit skalieren.

weiter auf Seite 24

Geleit

Software-Defined Networking: wer braucht das?

ab Seite 2

Standpunkt

Virtualisierung ohne Grenzen auch bei hohem Schutzbedarf?

auf Seite 22

Neuaufgabe

Aktuelle Sonderveranstaltung

Professionelle Datenkommunikation

auf Seite 21

Wireless Networking

ab Seite 23

Zum Geleit

Software-Defined Networking: wer braucht das?

Seit einiger Zeit wird in der Netzbranche intensiv über Software-Defined Networking (SDN) und damit zusammenhängend auch über OpenFlow diskutiert. Das Open Networking Foundation (ONF) genannte Konsortium hat sich beides, SDN und OpenFlow, auf die Fahne geschrieben und behauptet auf der eigenen Webseite, mit SDN werde man über Netze nie wieder wie früher denken. Nun, ewig das Gleiche denken tun nur jene, die nichts dazu lernen. Insofern wird jeder, der etwas davon versteht, in ein paar Jahren über eine so dynamische Technologie wie Kommunikationsnetze anders denken als heute. Die Frage ist nur, welche Rolle SDN und OpenFlow dabei spielen werden.

Ein Blick auf die Initiatoren der ONF gibt darüber Aufschluss, wen SDN und OpenFlow am meisten interessieren: Deutsche Telekom, Facebook, Google, Microsoft, Verizon und Yahoo, also allesamt Unternehmen, die sehr große Rechenzentren betreiben müssen. Technisches Kernstück von SDN ist das OpenFlow, ein Protokoll, mit dem Netzkomponenten mit sogenannten Controllern kommunizieren. Der Controller übernimmt Aufgaben, die bei einem konventionellen Switch oder Router von der sogenannten Control Plane übernommen werden. Die Unterscheidung der Control Plane und der Data Plane hat sich in den letzten ca. zwei Jahrzehnten bei Netzkomponenten etabliert. Aufgabe der Data Plane ist die (möglichst schnelle) Vermittlung von Paketen anhand sogenannter Forwarding-Tabellen. Diese legen fest, wie ein Netzgerät ein Paket mit bestimmten Merkmalen (zum Beispiel MAC- und/oder IP-Adressen) behandeln, d. h. vor allem werfen oder weiterleiten und über welchen Port senden soll. Die Einträge Forwarding-Tabelle können jedoch von Mechanismen der Control Plane stammen, zum Beispiel einem Routing-Protokoll, das in der Regel von der Central Processing Unit (CPU) eines Routers oder Switches bearbeitet wird.

Welchen Sinn kann es für die Konzentration der Control Plane auf den Controller geben? Die Idee ist, dass in einem großen Netz die Konfiguration vieler intelligenter Komponenten entfällt und komplexe Aufgaben der Control Plane nur noch vom Controller übernommen werden. Die Namensgleichheit mit einem Wireless Local Area Network (WLAN) Controller ist nicht zufällig. Auch in einem Controller-basierenden WLAN ist die Intelligenz auf die



Controller konzentriert. Die WLAN Access Points werden von Controllern gesteuert. Der OpenFlow-Ansatz ist ähnlich. OpenFlow Switches sollen sich auf die Data Plane konzentrieren. Die Control Plane ist die Domäne des Controllers. Man stelle sich ein ganz großes Rechenzentrum vor, das aus hunderten virtuellen Servern besteht, die auf tausenden physikalischen Servern laufen. Der Betreiber eines solchen RZs hat ein großes Interesse daran, dass die ganze Infrastruktur aus gleich aufgebauten und standardisierten Modulen besteht. Die Netzkomponenten dieser Module müssen dann einheitlich konfiguriert sein und nach den gleichen Verfahren und Modellen funktionieren. Die Konzentration der Intelligenz auf den Controller erleichtert die Einrichtung, den Betrieb und die Erweiterung einer solchen Infrastruktur.

Daher ist das Interesse großer RZ-Betreiber an OpenFlow kein Zufall. Auch kein Zufall ist die zeitliche Koinzidenz der OpenFlow-Aktivitäten mit dem Trend Cloud Computing. Wenn Betreiber öffentlicher Clouds erfolgreich Kunden anwerben, müssen sie eines ihrer wichtigsten Versprechen, nämlich dass sie eine IT-Infrastruktur dank Synergieeffekten effizienter und wirtschaftlicher betreiben können als ihre Kunden, einhalten. Sie müssen ein RZ der o. g. Dimensionen insgesamt mit deutlich weniger Mitarbeitern betreiben als die Summe der bisher dafür eingesetzten Mitarbeiter ihrer Kunden.

Die Motivation für die Standardisierung von OpenFlow ähnelt der Motivation jeder anderen Standardisierung, vor allem: Vermeidung der Abhängigkeit von einzelnen Herstellern, unkomplizierter Austausch ein-

zelner Komponenten und Herstellung einer Marktmacht durch Bündelung der Potenziale mehrerer Hersteller. Ein Controller des Herstellers x soll eine Komponente des Herstellers y steuern können, solange beide die OpenFlow-Spezifikation einhalten.

Was hat das Ganze mit dem Begriff SDN zu tun? Der Controller basiert vor allem auf Software, die auf einem Standard-Server laufen soll. Dadurch soll die Möglichkeit eröffnet werden, dass ein Markt für Lösungsanbieter entsteht, die ihre Netzlösung auf Controller-Basis entwickeln. Die Konfiguration des Netzes erfolgt dann Software-gesteuert, womit wir bei der Namensgebung für SDN wären.

Soweit die Idee. Wie ist sie aber zu bewerten?

Grundsätzlich ist zu berücksichtigen, dass das Design und der Aufbau sehr großer Rechenzentren, die von Betreibern öffentlicher Clouds genutzt werden, anderen Gesetzmäßigkeiten und Prioritäten folgen als das Design und der Aufbau von wesentlich kleineren Umgebungen, die nur von einem Unternehmen genutzt werden. Ein privates RZ hat meistens nicht die dringenden Probleme, mit denen sich ein Public-Cloud-Anbieter auseinandersetzen muss. Andererseits kann sich ein Unternehmen für das eigene, privat genutzte RZ-Netz in der Regel keine Lösung leisten, für die kein einziger Lieferant verantwortlich ist. Stellen Sie sich die Situation vor, in der ein Controller- und ein Switch-Hersteller sich gegenseitig für eine Fehlfunktion im Netz verantwortlich machen. Bei der Komplexität von OpenFlow sind Fehlfunktionen und Interoperabilitätsprobleme trotz Standard nicht ausgeschlossen. Der Alptraum jedes Betreibers eines kleinen, privaten Rechenzentrums ist es, mit seinem Problem allein gelassen zu werden, weil ihm selbst die Expertise und der tiefe Einblick fehlen, um die Zuständigkeit für ein Problem eindeutig zuweisen zu können. Große Cloud-Anbieter haben völlig andere Möglichkeiten, Probleme selbst zu analysieren, den Verantwortlichen dafür zu finden und ihn zu einer Lösung zu motivieren. Eine Firma Microsoft oder Google ist für jeden Lieferanten so wichtig, dass die Bearbeitung jedes Problems auch für den Lieferanten oberste Priorität hat. Wie viele RZ-Betreiber gibt es, denen eine solche privilegierte Stellung zukommt?

Deshalb legen die meisten Betreiber privater Rechenzentren Wert darauf, in ihrer

Schwerpunkthema

Routing im, vom und zum RZ

Fortsetzung von Seite 1



Dr. Behrooz Moayeri ist bei der ComConsult Beratung und Planung GmbH als Mit-glied der Geschäftsleitung tätig. Er hat in den letzten Jahren unter anderem viele Unternehmen zu Themen der RZ-Vernetzung beraten.

Die Antwort ist ein klares Nein und wird im Folgenden begründet.

Man stelle sich ein RZ-Netz vor, an das verschiedene RZ-typische Endgeräte wie Server, Network Attached Storage (NAS) etc. angeschlossen sind (Abbildung 1).

Wenn das in der Abbildung 1 dargestellte RZ-Netz als eine Layer-2-Broadcast-Domäne aufgebaut wird, gilt diese Broadcast-Domäne zugleich als eine einzige Sicherheits- und Fehlerdomäne, wenn man vom Einsatz der bisherigen Layer-2-Netzverfahren (einschließlich neuer Verfahren wie Shortest Path Bridging SPB, Transparent Interconnection of Lots of Links TRILL etc.) ausgeht.

Das ist jedoch in fast keinem RZ erwünscht. Es darf nicht sein, dass die fehlerhafte IP-Konfiguration irgendeines Endgeräts im RZ potenziell das gesamte RZ-Netz oder jedes beliebigen Servers darin lahmlegen kann. Und es darf nicht sein, dass die Kompromittierung eines einzelnen Systems im RZ die Sicherheit des gesamten Rechenzentrums beeinträchtigen kann.

Also teilt man das RZ-Netz in mehrere Layer-2-Broadcast-Domänen auf, die aus Gründen der Verfügbarkeit über mindestens zwei Routing-Instanzen miteinander verbunden sind (Abbildung 2).

Die in der Abbildung 2 dargestellten Routing-Instanzen müssen nicht unbedingt klassische Router oder Layer-3-Switches sein. Infrage kommen auch Sicherheitskomponenten mit Routing-Funktion, zum Beispiel Firewalls.

Aus denselben Gründen wie für die interne RZ-Netzstruktur angeführt ist kein RZ-Netz zu empfehlen, das sich in derselben Broadcast-Domäne befindet wie alle Cli-

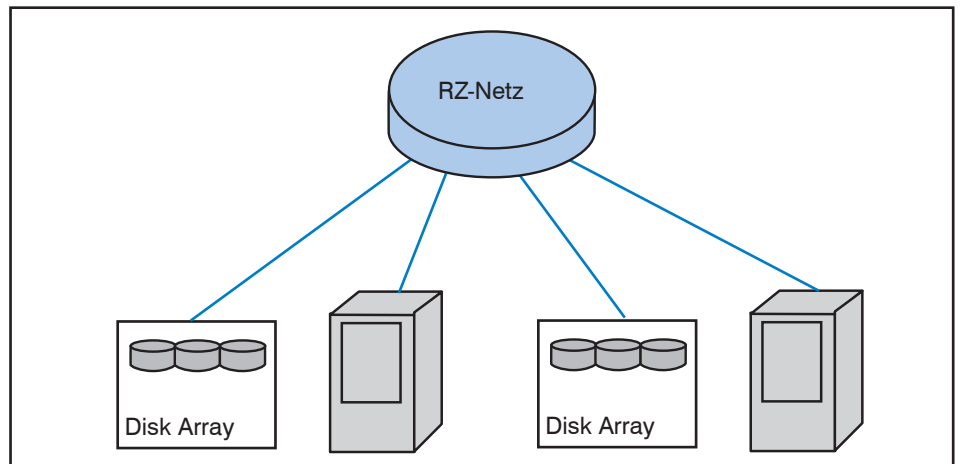


Abbildung 1: RZ-Netz mit Servern, NAS etc.

ents, die auf das RZ zugreifen. Das bedeutet, dass der Zugriff der Clients auf das RZ über Routing-Instanzen erfolgen muss. Soll dieser Zugriff von keinem Single Point of Failure abhängig sein, ist das RZ über mindestens zwei Layer-3-Instanzen mit den Clients verbunden sein (Abbildung 3).

Abbildung 3 zeigt die minimale Routing-Struktur eines RZ-Netzes. Kein RZ, das diesen Namen verdient, kann ohne die dargestellte Routing-Funktion im, vom und zum RZ auskommen.

Ungünstiges Routing

Nun, da wir die grundsätzliche Frage geklärt haben, was die Notwendigkeit von Routing im Zusammenhang mit dem RZ-Netzdesign betrifft, betrachten wir die Probleme, um die es in diesem Beitrag hauptsächlich geht und die wir als „ungünstiges Routing“ zusammenfassen wollen.

Um die Problematik „ungünstiges Routing“ zu verstehen, müssen wir zunächst

definieren, was „optimales Routing“ bedeutet. Einigen wir uns auf die Definition, dass optimales Routing die Auswahl der Pfade mit den kürzesten Signallaufzeiten und den niedrigsten Paketverlustraten ist. Letzteres setzt häufig voraus, dass die Übertragungskapazitäten im Netz optimal genutzt werden, zum Beispiel durch Verteilung der Last auf verschiedene Wege. Das heutige Internet weist diesbezüglich einerseits noch Optimierungspotenziale auf und ist andererseits das Ergebnis von über vierzig Jahren Erfahrungen bei der Optimierung von Routing.

Die Abbildung 4 zeigt zwei verschiedene Wege zwischen einem Server und einem Speichersystem in unserem „Minimal-RZ“.

In unserem Minimaldesign sind die beiden in der Abbildung 4 dargestellten Pfade hinsichtlich der Anzahl der Layer 3 Hops gleichwertig, d. h. bei Anwendung optimaler Routing-Verfahren kommt es zur Nutzung beider Pfade.

Nun stelle man sich vor, einer der Wege

Routing im, vom und zum RZ

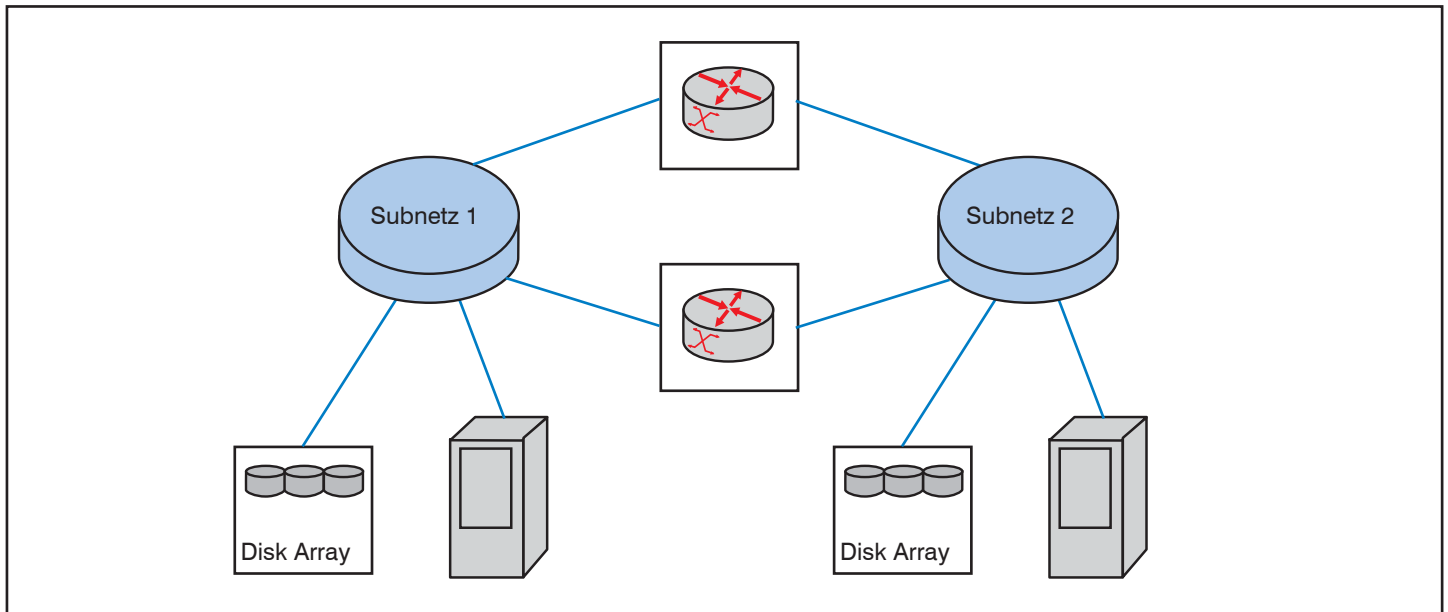


Abbildung 2: Aufteilung des RZ-Netzes in Broadcast-Domänen

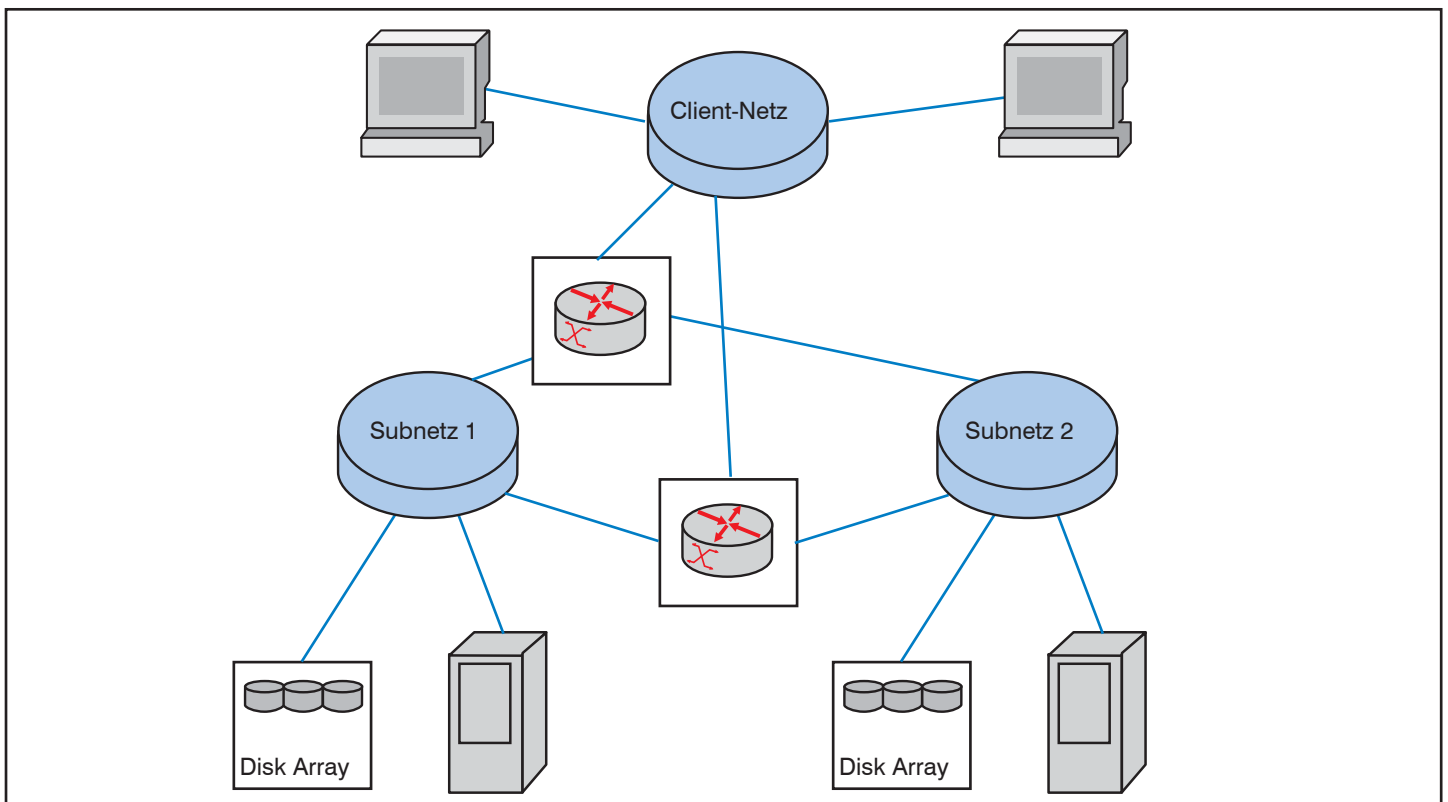


Abbildung 3: Routing im RZ sowie zwischen RZ und Clients

sei hinsichtlich der tatsächlichen Signallatenz weniger optimal als der andere. Das gilt zum Beispiel, wenn der Server, das NAS-System und einer der Router am selben Standort, aber der andere Router an einem anderen Standort aufgestellt sind. Dazu kann es leicht kommen, wenn die beiden RZ-Standorte, wie häufig gefordert, über Layer 2 transparent verbunden sind.

Wegewahloptimierung im RZ

Um sicherzustellen, dass eine optimale Layer-3-Route auch mit optimalen Layer-2-Pfaden verbunden ist, dürfen die Wegewahlentscheidungen auf Layer 3 und Layer 2 nicht entkoppelt werden. Damit kommen wir zu dem in der Abbildung 5 dargestellten Design.

Wie aus der Abbildung 5 hervorgeht, müssen die beiden Netzkomponenten nicht nur Layer-3-, sondern Layer-2-Intelligenz besitzen, um für eine aus Gesamt-sicht optimale Wegewahlentscheidung treffen zu können. Eine solche Entscheidung kann zur Auswahl des in der Abbildung 5 dargestellten Pfades vom Server zum NAS-System führen. Der Pfad geht über die erste Layer-2/3-Komponente in

Zweitthema

Neue Protokolle im RZ: Funktions- umfang – Potential – Auswirkung auf die Infrastruktur



Dipl.-Math. Cornelius Höchel-Winter ist Leiter des Testlabors der ComConsult Research GmbH. In dem Labor werden regelmäßige Messungen und Evaluierungstests neuester Hard- und Softwareprodukte durchgeführt und ausgewertet. Herr Höchel-Winter besitzt langjährige Erfahrung in der Konzeptionierung, im Aufbau und Betrieb von Windows- und Unix-netzen; so hat er als verantwortlicher Projektmanager die Rechenzentren und Netzwerke auf dem Gelände der EXPO2000 in Hannover aufgebaut und während der Weltausstellung betrieben.

Fortsetzung von Seite 1

Dahinter steckt das Phänomen, dass einzelne Nutzer einer Web-Applikation kaum Last erzeugen und die Last nur als Funktion der Anzahl der Teilnehmer entsteht. Es entsteht also der Bedarf nach einer dynamisch wachsenden Infrastruktur in Abhängigkeit von der Teilnehmerzahl. Wenn also ein Unternehmen eine neue Web-Anwendung für Kunden oder Zulieferer im Internet anbieten will, dann entsteht automatische die Frage nach der Skalierung.

Der andere Megatrend ist die automatische Provisionierung von Anwendungsserver in einer virtuellen Infrastruktur. Heute werden Anwendungsserver in der Regel von Hand aufgesetzt und physischen Servern zugewiesen. Auch mögliche Wanderungen dieser virtuellen Server werden manuell konfiguriert. Damit werden über den Tag hinweg die Ressourcen in einer virtuellen Infrastruktur nicht optimal ausgenutzt. Ideal wäre eine vollautomatische Zuweisung und Pflege dieser Zuweisungen. Damit verbunden sind automatische Wanderungen von virtuellen Maschinen auch zwischen Standorten oder über Brandschutzbereiche hinweg. Heutige Lösungen dieser Art sind häufig zu komplex und auch zu teuer, sie werden häufig unter dem Schlagwort Private Cloud angeboten und umgesetzt, da die automatische Provisionierung ein zentrales Merkmal von Cloud-Lösungen ist. Aufgrund dieses hohen Preises kommen

diese Lösungen bisher eher selten zum Einsatz. Aber es ist erkennbar, dass Hersteller wie VMware oder auch HP intensiv in dieser Richtung arbeiten. Und damit entsteht ebenfalls die Frage, in welchem Umfang virtuelle Layer-2-Verbindungen über Netzwerkengrenzen hinweg erforderlich sein werden, um dieses Konzept der automatischen Provisionierung wirklich rund zu machen.

Brandaktuell wird das Thema durch die Aktivitäten von VMware mit VXLAN und der großen Unterstützung, die dieses Thema seitens der Netzwerkhersteller erhält. Auch der Gegenentwurf NVGRE schafft vergleichbare Lösungen. Im Folgenden soll daher speziell untersucht werden, welches Potenzial diese Protokolle haben.

An dieser Stelle soll auch erwähnt werden, dass es weitere technische Ansätze zur Lösung dieses Problems gibt, die zum Teil auch deutlich weiter gehen. Das sind zum einen Software Defined Networks (SDN) und zum anderen Service-Architekturen unter Nutzung des Service Tags von Shortest Path Bridging. Wir werden in weiteren Artikeln auf diese Optionen eingehen und Lösungsstrategien natürlich auch auf dem Rechenzentrum Redesign-Forum der ComConsult Akademie Anfang November in Köln diskutieren.

Bevor wir uns an dieser Stelle aber Ein-

satzszenarien und Zukunftsperspektiven widmen, lassen Sie uns zunächst die technischen Grundlagen der beiden Protokolle VXLAN und NVGRE klären.

VXLAN – Die Grundlagen

VXLAN ist als sogenannter „Internet-Draft“ bei der IETF veröffentlicht. Das aktuelle Dokument (draft-mahalingam-dutt-dcops-vxlan-02 vom 22.8.2012) hat den Status „Experimental“, mit einer Verabschiedung als „Internet-Standard“ ist vorerst nicht zu rechnen, was aber erfahrungsgemäß keinen Einfluss darauf hat, in wie weit das Protokoll in neuen Produkten zum Einsatz kommt.

Autor ist Mallik Mahalingam, führender Software-Architekt bei VMware, weitere Autoren sind von Arista, Broadcom, Cisco, Citrix und Rad Hat.

Das Protokoll selbst kann als MAC-in-IP oder genauer als MAC-in-UDP-Verfahren bezeichnet werden, das Frame-Format ist in Abbildung 1 dargestellt.

Der Original-Frame (ohne dessen ursprünglicher FCS) wird also nacheinander um vier Header erweitert:

- Der VXLAN-Header enthält im Wesentlichen einen 24 Bit großen Identifier (VNI = VXLAN Network Identifier). Damit

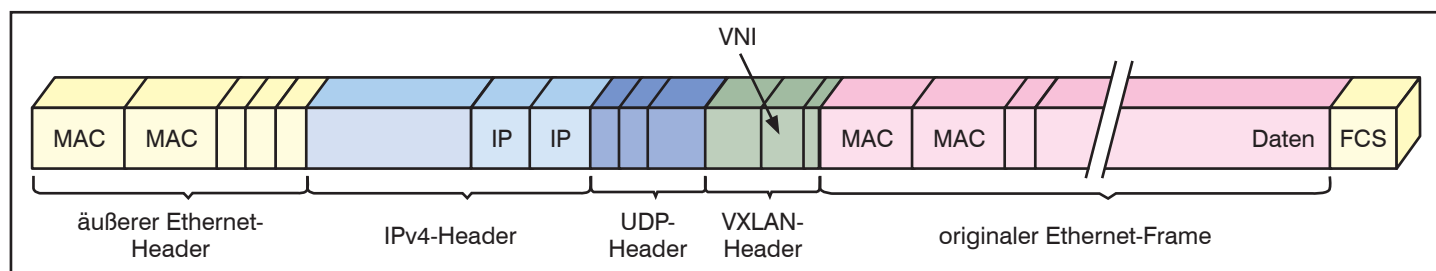


Abbildung 1: VXLAN Frame-Format

Neue Protokolle im RZ: Funktionsumfang – Potential – Auswirkung auf die Infrastruktur

können mehr als 16 Millionen VXLAN-Segmente unterschieden und somit deutlich mehr als mit der nur 12 Bit großen 802.1Q-VLAN-ID. Wie bei 802.1Q können virtuelle Maschinen in verschiedenen VXLAN-Segmenten nicht miteinander kommunizieren.

- Die UDP- und IP-Header sorgen dafür, dass der eingekapselte Frame über Layer 3 transportiert werden kann. Die Zieladresse im IP-Header ist hierbei die IP-Adresse des Tunnelendpunktes, über den das eigentliche Zielsystem (MAC-Adresse im originalen Frame) erreicht werden kann.
- Und abschließend folgt ein passender äußerer MAC-Transport-Header, der sich vermutlich beim Durchgang durch das IP-Netz regelmäßig ändert.

Das Verfahren definiert also pro VNI je ein Layer-2-Overlay-Netz über ein Layer-3-Transportnetz und verbindet so alle Layer-2-Segmente, denen dieselbe VXLAN-ID zugewiesen wurde. Oder anders ausgedrückt: Layer-2-Netze werden so über Layer 3 hinweg aufgespannt.

Segmente mit unterschiedlicher VNI sind völlig voneinander getrennt und können nicht miteinander kommunizieren. Daher spielt es auch keine Rolle, wenn in solchen Segmenten gleiche VLANs oder gleiche MAC-Adressen verwendet werden.

Der Draft schränkt den Einsatzbereich von VXLAN explizit auf RZ-Netze mit vir-

tuellen Servern ein („VXLAN addresses ... data center network infrastructure in the presence of VMs in a multitenant environment“). Die Tunnelendpunkte (VTEP = VXLAN Tunnel End Point) werden dementsprechend auch innerhalb des Hypervisors der Virtualisierungs-Hosts positioniert. Gleichwohl merken die Autoren aber auch an, dass VTEPs auch in physischen Komponenten realisiert werden können („Note that it is possible that VTEPs could also be on a physical switch or physical server and could be implemented in software and hardware.“).

Grundlage des Protokolls sind im Wesentlichen Zuordnungstabellen, in denen jeder VTEP einerseits die MAC-Adressen der lokalen Endsysteme (also im Wesentlichen die „seiner“ virtuellen Server) einer VNI und andererseits die MAC-Adressen der entfernten Zielsysteme der IP-Adresse des zugehörigen VTEPs zuordnet:

- Durch die Zuordnung zu einer VNI wird die jeweilige MAC-Schnittstelle einem bestimmten VXLAN-Segment, sprich Layer-2-Segment zugeordnet. Diese Zuordnung geschieht initial via Management auf dem System, das die betreffende VM hostet. Diese Zuordnung ist vergleichbar mit der im vSwitch üblichen Zuordnung eines Endsystems zu einer VLAN-ID.
- Die Zuordnung zur IP-Adresse des zugehörigen VTEPs wird benötigt, um bei der Einkapsulierung eines Datenpakets die Ziel-IP-Adresse im äußeren IP-Header setzen zu können. Wie diese

Zuordnung erfolgt lässt der Draft explizit offen, als ein mögliches Verfahren wird lediglich das sogenannte „data plane learning“ beschrieben. Hierbei lernt, wie bei normalen Switches auch, die Zielkomponente die Zuordnung des Absenders anhand der Adressen im empfangenen Paket.

Eine Besonderheit dieses „data plane learning“ ist, dass Broadcasts, Multicasts und Unicasts an unbekannte Zieladressen mittels IP-Multicasts weitergeleitet werden. Damit auch diese Kommunikation innerhalb des jeweiligen VXLAN-Segments bleibt, erhält jedes VXLAN-Segment eine eigene IP-Multicast-Adresse. Diese Zuordnung von IP-Multicast-Adresse zu VNI erfolgt laut Draft-Dokument auf Management-Ebene und wird über einen nicht näher spezifizierten „management channel“ an die verschiedenen VXLAN-Endpunkte verteilt werden („This mapping is done at the management layer and provided to the individual VTEPs through a management channel.“). Die Endpunkte selbst propagieren dann ihre jeweilige Zugehörigkeit zu bestimmten Multicast-Gruppen (also welche VXLAN-Segmente von ihnen bedient werden) dann via IGMP an das Transportnetzwerk.

NVGRE – Die Grundlagen

NVGRE („Network Virtualization using GRE“) ist ein alternativer Vorschlag von Microsoft in Zusammenarbeit mit Arista, Intel, Dell, HP, Broadcom und Emulex. Das Dokument (aktueller Stand: draft-sridharan-virtualization-nvgre-01 vom

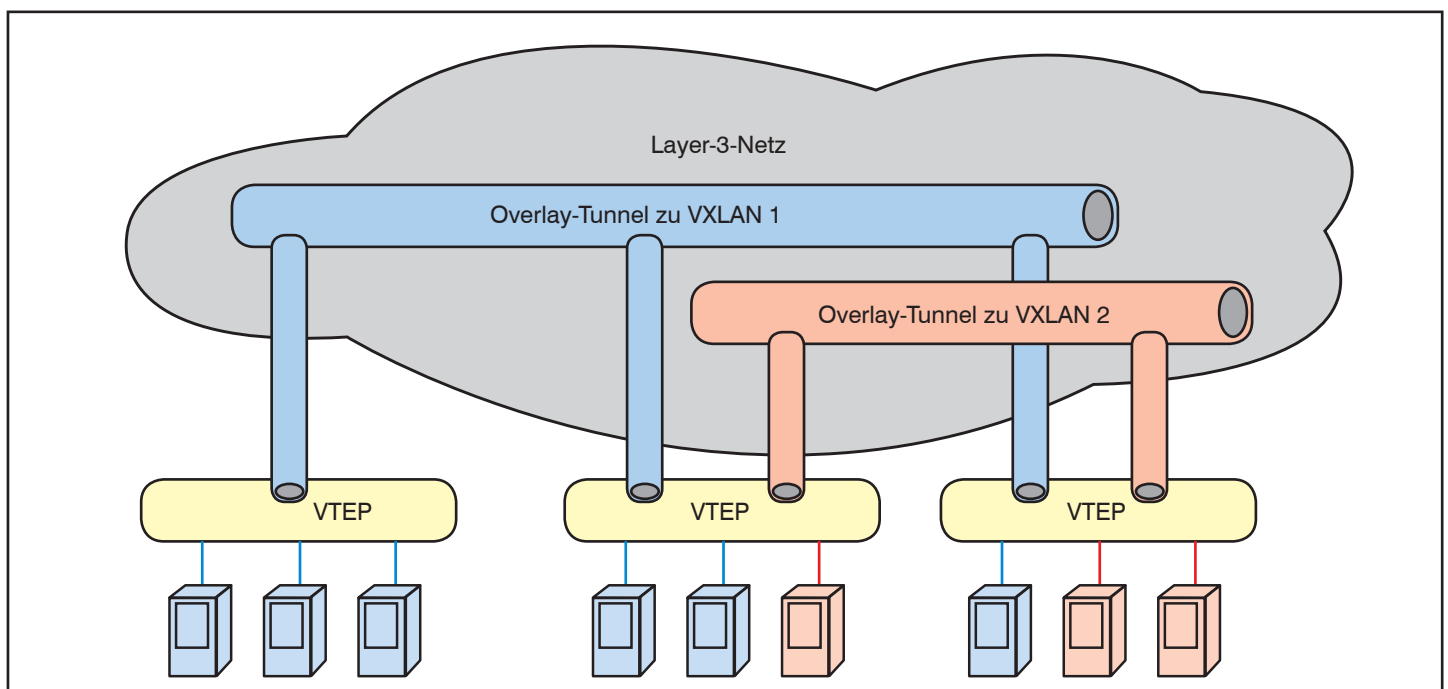


Abbildung 2: Overlay-Netze durch VXLAN