

IPv6: Fundamente richtig legen

von Dr. Behrooz Moayeri

Die Welt der Netze steht vor der nächsten großen Umstellung. Das Internet Protocol der Version 6 (IPv6) wird kommen. Daran zweifelt niemand mehr ernsthaft. IPv6 wird an den Grenzen keines Unternehmens halt machen. Jedes Unternehmen muss sich auf IPv6 vorbereiten. Wenn die Fundamente richtig gelegt werden, wird die Umstellung auf IPv6 weitgehend schmerzfrei verlaufen. Anderenfalls drohen instabile Zustände und ein großer Migrationsaufwand.

Aber welche sind die richtigen Fundamente? Die ersten Erfahrungen aus der Praxis liegen vor. Alle Unternehmen können von diesen Erfahrungen lernen. Und sie soll-

Die Zeit des Zurücklehns ist vorbei

ten das tun. Die Zeit dafür ist jetzt. Bevor äußere Zwänge zum überstürzten Handeln zwingen, müssen die Verantwortlichen für die IT-Infrastruktur die geplante und gut durchdachte Einführung von IPv6 im eigenen Unternehmensnetz vorbereiten. Der erste Schritt ist der Aufbau von Know-how über IPv6. Einiges bei IPv6 unterscheidet sich vom Pendant unter IPv4. Und da ist noch die lange Phase des Nebeneinanders der beiden Protokolle. Dieselben Systeme und Anwendungen können beide Protokolle nutzen. Dies sollte kontrolliert und gesteuert erfolgen. Alles dem Zufall zu überlassen kann heilloses Durcheinander, Ausfälle und viel Arbeit bedeuten.

weiter auf Seite 7

Zweitthema

SDN, NFV, OpenFlow und Virtualisierungs-Protokolle – Zusammenhänge, Perspektiven, Marktrelevanz

von Dipl.-Inform. Petra Borowka-Gatzweiler

SDN und NFV scheinen zur Zeit ganz oben im Hypecycle der Netzbetreiber zu stehen. Sind SDN und NFV dasselbe oder sind es unterschiedliche Technologien? Welche Märkte bedienen sie? Wohin gehören OpenFlow, Controller, Northbound, Southbound, East-Westbound API, Network Function Virtualisation Infrastructure sowie Virtualisierungs-Protokolle wie VXLAN, GENEVE und NSH? Welche Relevanz haben diese Technologien für

Enterprise Netzwerke? Der nachfolgende Beitrag beleuchtet diese Themen, ihre technologischen Hintergründe und ihre praktische Marktrelevanz.

1. SDN – Software Defined Networking

Heutige Netzwerke bestehen vielfach aus proprietärer Hardware, die mittels Hersteller-Software konfiguriert und administriert wird. Somit sind Hardware und Control Plane herstellerspezifisch. SDN ist

eine aufkommende Netzwerk-Technologie auf Basis einer standardisierten Netzwerk-Architektur, die den bislang proprietären Control Plane Konzepten etablierter Netzwerk-Komponenten ein Ende bereiten will: Die Control Plane wird aus den Netzkomponenten auf eine zentrale standardisierte Steuerung ausgelagert. Hierdurch wird sie für den Betreiber (mittels remote Steuerung standardisierter Clients / Agenten) zugreifbar und programmierbar.

weiter auf Seite 18

Geleit

Branded Whites Box-Switches: was soll das und was bedeutet das für die Branche?

auf Seite 2

Aktueller Kongress

ComConsult Netzwerk- und IT-Infrastruktur Forum 2015

ab Seite 4

Standpunkt

WLAN: Wird 5 GHz zum neuen 2,4-GHz-Band?

auf Seite 16

Neues Seminar

Strategien für Unternehmen zur richtigen Positionierung im Internet

auf Seite 17

Zum Geleit

Branded Whites Box-Switches: was soll das und was bedeutet das für die Branche?

Die Zeit proprietärer Netzwerk-Hardware könnte sich bald dem Ende nähern. Nicht in dem Sinne, dass normale Switches wie wir sie von den üblichen Anbietern kennen, komplett verschwinden sondern mehr in dem Sinne, dass eine neue Produktkategorie zunehmend an Bedeutung gewinnt. So beobachten wir in den großen Cloud-Rechenzentren seit Jahren eine zunehmende Ablehnung von Standard-Produkten, sowohl auf der Netzwerk als auch auf der Server-Seite. Dabei werden interessanterweise nicht nur die Anschaffungskosten kritisiert, sondern vor allem die Betriebskosten, die mit den bisherigen proprietären Lösungen einher gehen. Betrachtet man speziell das Open Compute Projekt von Facebook, dann wird durchaus deutlich, dass ein großer Markt, der bisher auf wenige starke Anbieter reduziert war, wichtige Verbesserungen und Veränderungen im Produktdesign total verschlafen hat. Es war ja auch nicht notwendig, die Kunden kauften ja auch so. In letzter Konsequenz sind so Lösungen entstanden, die den Kunden deutlich zu hohe Betriebskosten zumuten (stark abhängig von der Zahl der eingesetzten Switches und Server).

So ist in den letzten Jahren die White-Box-Bewegung entstanden, die sich vor allem an die sehr großen Rechenzentren dieser Welt wendet. Dabei werden Hardware und Software entkoppelt. Ein Anbieter konzentriert sich auf Hardware-Massenware, ein anderer auf die Software. Hardware-Massenware ist dabei durchaus naheliegend und liegt im Trend der letzten Jahre. Mehr und mehr Hersteller von Netzwerk-Komponenten haben die eigene ASIC-Entwicklung entweder eingestellt oder reduziert. Viele der heute angebotenen Switches sind bereits verkapselte OEM-Boxen, die nur durch die Software des Anbieters und die Art des Gehäuses eine gewisse Eigenständigkeit bekommen. Das mag nicht unbedingt für die Top 10% der Hochleistungs-Switches gelten, aber für die Massenware im RZ oder im Access-Bereich ist dies definitiv der Fall. Access-Switches sind bereits weitestgehend austauschbar und die Hersteller gehen lange Wege um diese Austauschbarkeit durch Software-Eigenschaften zu kaschieren (manchmal durchaus mit Sinn, zum Beispiel wenn es um die Rechts- und Zugangsverwaltung für Endgeräte geht oder um die Vereinheitlichung eines Wi-



red und Wireless Access). Für den RZ-Betreiber bedeutet der neue Trend vor allem, dass er sich betriebsseitig für eine Software entscheiden kann und nach Bedarf die Hardware zukaufen kann, die am besten zu seinem Betrieb passt. Der Betrieb wird einfacher, und natürlich sind die Einkaufspreise deutlich günstiger.

Der hier angesprochene Markt ist so groß, dass ihn die traditionellen Anbieter nicht einfach ignorieren können. Trotzdem überrascht Hewlett Packard den Markt mit seiner Zusammenarbeit mit Accton und Cumulus. HP wird in den nächsten Wochen zwei Typen von Switch-Systemen zur Vermarktung durch Accton bereit stellen. Damit entsteht eine branded White-

Box, die vor allem die Eigenschaft haben wird, mehrere unterschiedliche Betriebssysteme zu betreiben. Um dem Kunden die Entscheidung zu erleichtern, hat HP gleichzeitig mit Cumulus einen Vertrag geschlossen und stellt die Cumulus-Software als Betriebssystem für die Accton-Switches zur Verfügung (ein Linux-basiertes Netzwerk-Betriebssystem). Was auf den ersten Blick vielleicht etwas paradox wirkt, kann sich bei näherer Betrachtung zu einem gewinnbringenden Geschäftsmodell für HP entwickeln. Die bisherigen White-Box-Lösungen haben ihren typischen Nachteil im Bereich des Services. Die Kunden müssen die Integration und mögliche Probleme selber bearbeiten. HP steigt mit seinen neuen Kooperationen in diesen Markt ein und will sich als das führende Service-Unternehmen für branded White-Ware positionieren. Es bleibt abzuwarten, ob das gelingt, aber der Schritt ist sicher interessant.

Die im Moment von HP und Accton geplanten Produkte sind 10/40-Gigabit Massenprodukte, wie wir sie in RZs in großen Stückzahlen sehen. Sie werden später im Jahr ergänzt durch die neuen 25/50/100G Switches.

HP ist mit dieser Vorgehensweise nicht allein, reduziert seine Kooperationen aber auf genau definierte und somit eingeschränkte Produkte. Vermutlich will man auch abwarten, ob das angestrebte Ser-

Seminar

Netzwerk-Design für Enterprise Netzwerke 18.05. - 19.05.15 in Köln

LAN-Technik wird im Moment neu erfunden. Neue Anforderungen erfordern neue Lösungen. Neue Fabric-Konzepte, ein Umdenken bei VLAN-Technik, eine Neupositionierung von QoS und neue Nutzungsformen im Rahmen von Audio-/Video-Bridging sind herausragende Beispiele. Das Seminar zum Thema Netzwerk-Design für Enterprise Netzwerke erklärt, was im Moment passiert und wie Sie sich auf die Zukunft vorbereiten. Es geht auf RZ- und Campus Design-Alternativen im Zeitalter neuer Layer-2 Technologien wie Fabrics, Multichassis-Link Aggregation, Shortest Path Bridging und Hochgeschwindigkeits-Datenraten von 10/40/100 Gbit ein.

Referentin: Dipl.-Inform. Petra Borowka-Gatzweiler

Preis: € 1.590,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Schwerpunktthema

IPv6: Fundamente richtig legen

Fortsetzung von Seite 1



Dr.-Ing. Behrooz Moayeri hat viele Großprojekte mit dem Schwerpunkt standortübergreifende Kommunikation geleitet. Er gehört der Geschäftsleitung der ComConsult Beratung und Planung GmbH an und betätigt sich als Berater, Autor und Seminarleiter.

Deshalb ist es wichtig, dass die Experten im Unternehmen die Unterschiede zwischen IPv4 und IPv6 kennen. Ferner ist es von Bedeutung zu wissen, welche Probleme und Unwägbarkeiten es in einer dualen Umgebung mit beiden Protokollen geben kann.

Die letzten Zweifler überzeugen

Noch sind nicht alle, auf deren Stimme und Einfluss es in der IT ankommt, davon überzeugt, dass man sich mit IPv6 befassen muss. Immer seltener, aber häufig genug hört man das Argument, dass ein Protokoll, von dem schon seit fast 20 Jahren gesprochen wird, noch lange auf sich warten lassen wird. Vor allem weil viele Anwendungen und Systeme IPv6 noch nicht unterstützen, brauche man sich mit diesem Protokoll gar nicht zu befassen. So das Argument der Zweifler. Zu diesen gehören leider auch viele, die über IT-Budgets entscheiden und die Prioritäten in Unternehmen setzen. Diese Zweifler sind der Meinung, die Unternehmens-IT habe dringendere Aufgaben zu bewältigen. Die dringenden Probleme seien wichtiger als ein Protokoll einzuführen, das zurzeit von den meisten Applikationen und Geräten noch gar nicht unterstützt werde.

Wahrscheinlich lassen sich diese Zweifler von der rhetorischen Frage kaum beeindrucken, die lautet: Warum muss eine Aufgabe erst zu einer dringenden werden, bevor sie angepackt wird? Es gilt nun man: Dringenderes zuerst.

Was vielleicht überzeugender auf die Zweifler wirkt, ist der Hinweis, dass die IPv6-Einführung insgesamt nicht als „Big Bang“ durchgeführt werden muss. Das Stichwort heißt „Lifecycle“. Die IPv6-Umstellung betrifft die Gesamtheit der IT-Systeme: Switches, Router, Sicher-

heitskomponenten, Server, Endgeräte, Betriebssysteme, Middleware und Anwendungssoftware. Es ist unvorstellbar, dass die Umstellung all dieser Komponenten auf IPv6 nach einem Stichtagsprinzip gelingen kann. Deshalb ist mit einer langen Umstellungsphase zu rechnen. Vor diesem Hintergrund ist es geschickter, die IPv6-Umstellung nicht losgelöst, sondern im Zuge vom normalen Lebenszyklus der Komponenten vorzunehmen. Konkret bedeutet dies: Ist die Ablösung einer Anwendung, eines Systems, einer Middleware zu einem bestimmten Termin ohnehin vorgesehen, ist dieser Termin der geeignetste Zeitpunkt der Umstellung auf den dualen Betrieb mit IPv4 und IPv6 bzw. auf reinen Betrieb mit IPv6.

Auf diese Weise wird die Hemmschwelle gesenkt, die viele Unternehmen bisher von der Beschäftigung mit IPv6 abgehalten hat. IPv6-Einführung als Begleiterscheinung des normalen Lifecycles der IT-Komponenten lässt sich einfacher durchsetzen als ein Umstellungsprojekt mit eigenem hohem Budget.

Aber die Dringlichkeits- und Budgetüberlegungen sind nicht die einzigen Argumente gegen IPv6. Auch technische Argumente werden manchmal bemüht. So wird behauptet, die Adressknappheit als Hauptmotivation für die Einführung von IPv6 sei ein Phantom. Das Argument lau-

tet konkret, global gültige IPv4-Adressen seien schon seit den 1990er Jahren knapp, und man habe gelernt, damit umzugehen. Diese Position weist auf die Abhilfe durch Network Address Translation (NAT) hin, die angeblich das Problem der Adressknappheit eliminiere. (siehe Abbildung 1)

In ihrer Not haben sogar die Service Provider für die eigenen Netze NAT eingeführt. Wer die IPv4-Adresse seines mobilen Gerätes im öffentlichen Netz überprüft, stellt oft fest, dass die Adresse einem privaten Adressraum wie 10.0.0.0/8 oder 100.64.0.0/10 entstammt. Im Umkehrschluss bedeutet dies, dass die heutigen Applikationen auf eine global eindeutige Adresse nicht angewiesen sind.

Diese Argumentation lässt außer Acht, dass die Service Provider mit Skalierbarkeits- und Betriebsproblemen im Zusammenhang mit NAT kämpfen. Der Verkehr im Netz explodiert. Dies gilt auch für die Anzahl paralleler Verbindungen. NAT-Komponenten, die die rapide steigende Zahl der Verbindungen in Tabellen verwalten, werden zu Nadelöhren im Netz. Immer mehr solche Komponenten müssen betrieben werden. Dies bedeutet mehr Aufwand für die Service Provider und oft auch Ärgernis für ihre Kunden.

Ferner legt IPv4 der Cloud Ketten an.

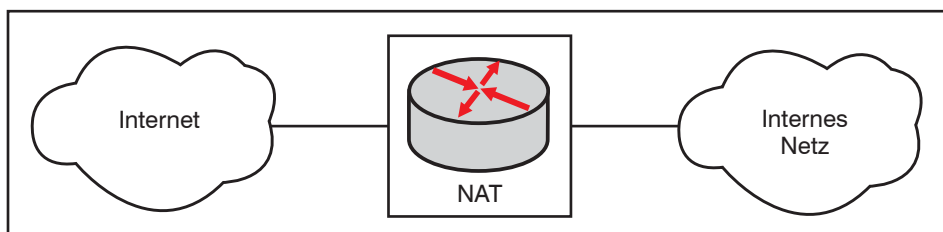


Abbildung 1: Network Address Translation

IPv6: Fundamente richtig legen

Eine Cloud lebt von der Skalierbarkeit, auch was die Anzahl der in der Cloud realisierbaren Server betrifft. Jede, auch jede virtuelle, Maschine in der Cloud braucht mindestens eine IP-Adresse, über die sie erreichbar ist. Diese Adresse sollte eine öffentliche sein. Hinter NAT verborgen kann ein Server nur ca. 60.000 gleichzeitige Verbindungen unterstützen, weil die zur Indizierung genutzte Portnummer nur 2^{16} Werte ermöglicht. Nicht von ungefähr zahlte Microsoft 7,5 Millionen Dollar für ca. 666.000 IPv4-Adressen aus der Konkursmasse von Nortel. Jeder Cloud-Betreiber braucht eine schnell wachsende Zahl von öffentlichen Adressen.

Übrigens gibt es die NAT-Restriktion auch beim ausgehenden Zugriff aus einem Unternehmensnetz auf einen öffentlichen Server. Eine öffentliche IP-Adresse kann per NAT nur ca. 60.000 gleichzeitige Verbindungen mit demselben Server erlauben.

Es gibt auch Sicherheitsbedenken gegen IPv6. Hier lautet das Argument, die Nutzung privater IPv6-Adressen sei ein Sicherheitsvorteil. Eine private Adresse sei von außen nicht erreichbar, und das könne nur gut sein. Wer so argumentiert, verkennet, dass die direkte IP Connectivity für die meisten der heutigen Angriffsmuster nicht erforderlich ist. Oder anders formuliert: Es ist um die Sicherheit des Unternehmens nicht gut bestellt, wenn sich das Unternehmen nur auf private Adressen im internen Netz setzt. Dem Autor sind Unternehmen bekannt, die seit über zwei Jahrzehnten keine anderen als global gültige IPv4-Adressen einsetzen. Die Netze dieser Unternehmen gehören zu den sichersten, die der Autor kennt. Den IP-Durchgriff ins Unternehmensnetz hinein zu verhindern gehört zu den leichtesten Sicherheitsaufgaben und ist auch ohne NAT zu lösen. Heutige Angriffe sind viel komplexer und nutzen erlaubte Kanäle, um selbst bei Anwendung von NAT in das Unternehmensnetz zu gelangen. (siehe Abbildung 2)

Eine weitere Argumentationslinie gegen die Einführung von IPv6 im Unternehmensnetz besteht darin, IPv6 als eine reine Internet-Angelegenheit zu bezeichnen. Wer so argumentiert, sieht die Notwendigkeit des neuen Protokolls im öffentlichen Netz ein, um im gleichen Atemzug diese Notwendigkeit in Unternehmensnetzen abzustreiten. Man könne doch im internen Netz ein anderes Protokoll nutzen als im Internet, wo sei das Problem? Diese Argumentation baut darauf, dass die Anbieter von Hardware und Software IPv4 unbegrenzt unterstützen

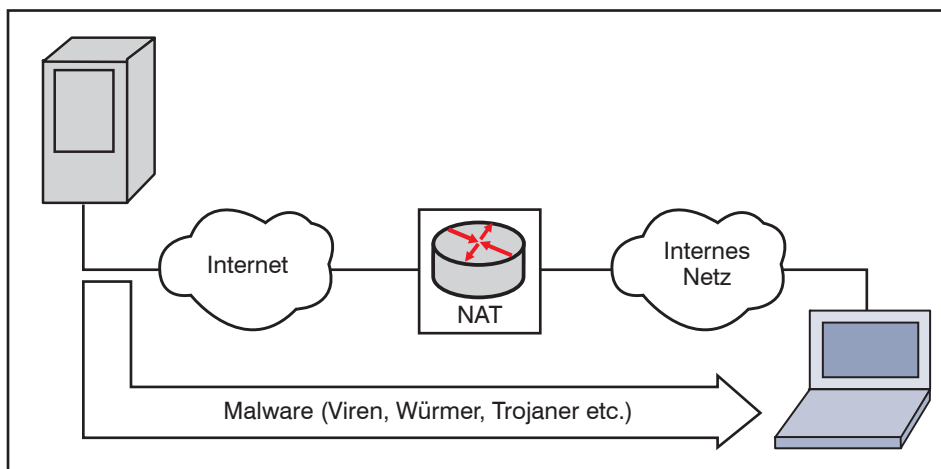


Abbildung 2: Bedrohungen, für die NAT keine ernste Barriere darstellt

werden. Die Erfahrung der letzten Jahrzehnte spricht aber dagegen. Sobald die seit einigen Jahren schnell wachsende IPv6-Nutzung im Internet IPv4 überflügelt, wird die duale Welt infrage gestellt. Viele Entwicklungen in der IT greifen vom Verbraucher- auf den Unternehmensmarkt über. So war das bei PCs, Smartphones, Tablets, und so wird es bei IPv6 sein. Das Leben auf der IPv4-Insel der Glückseligen, während ringsherum Milliarden von Geräten über IPv6 kommunizieren und das Internet of Things alles erfasst, was mit Strom oder Batterie arbeitet, wird irgendwann nicht mehr möglich sein. Es werden Systeme und Applikationen kommen, die nur IPv6 und kein IPv4 mehr unterstützen. Erste Anzeichen dafür gibt es schon. Microsoft schränkt den Support für Windows-Systeme ein, auf denen IPv6 abgeschaltet ist.

IPv6-Adresskonzept ausarbeiten

Ist es gelungen, die letzten Zweifler im eigenen Unternehmen zu überzeugen, muss man zur Vorbereitung der IPv6-Einführung zunächst konzeptionelle Arbeit leisten. Dazu gehört die Ausarbeitung eines Plans für die im Unternehmen zu nutzenden IPv6-Adressen.

IPv6 wird an einigen Prinzipien des Netzdesigns nichts ändern. Zu diesen Prinzipien gehört die Einteilung der Netze in Subnetze. Auch IPv6-Subnetze müssen über Router miteinander kommunizieren. Deshalb werden automatisch generierte Link-Local Addresses (LLAs) nicht ausreichen. LLAs erlauben nur die Kommunikation in einem Subnetz.

Geroutete IPv6-Adressen sind entweder Global Addresses (GAs) oder Unique Lo-

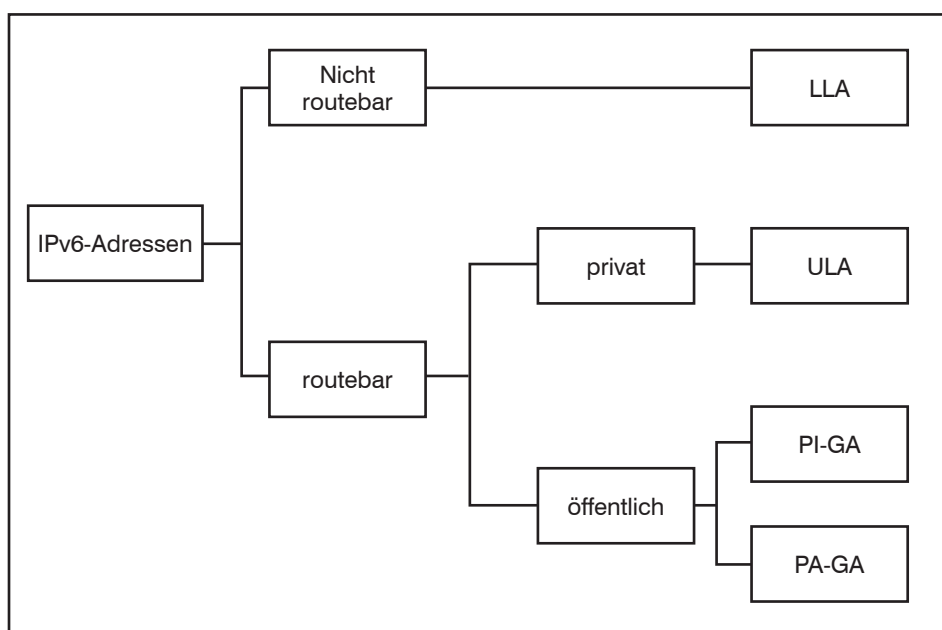


Abbildung 3: Kategorisierung von IPv6-Adressen

Zweitthema

SDN, NFV, Open-Flow und Virtualisierungs-Protokolle

Zusammenhänge, Perspektiven, Marktrelevanz

Fortsetzung von Seite 1

SDN soll somit herstellerübergreifend in Multivendor-Umgebungen einsetzbar sein. Um dies für alle Produkte zu erreichen, wird die Control Plane offengelegt, es werden offene, standardisierte Protokolle entwickelt, die Control Plane setzt auf einem Standard-Server und Standard-Betriebssystem auf.

SDN will die Netzwerk-Kontrolle (Lernen, Weiterleitungs-Entscheidungen und Policies) von der Netzwerk-Hardware, -Topologie und Paketweiterleitung (physikalische Verbindungen, Interfaces und deren Beziehung zueinander) entkoppeln.

Da die Control Plane ausgelagert ist, kann sie verschiedenste Arten von Netzkomponenten steuern: vSwitches, Hardware-Switches, Router, WLAN Access Points, Load Balancer, Traffic Shaper, Firewalls etc. Somit kann eine Gesamtsteuerung der kompletten verkabelten und kabellosen Netzwerk-Infrastruktur aus einer gemeinsamen Topologie-Sicht heraus implementiert werden. In diesem Sinn zielt SDN auf eine übergeordnete Orchestrierung des Netzwerks und der Netzwerk-Dienste hin, bei der die Netzwerk-Dienste von den physikalischen Netzwerk-Interfaces und physikalischen Netzwerk-Infrastrukturen losgelöst sind.

Globale Service-Definitionen müssen nicht mehr auf die physikalischen Interfaces gemappt werden, Service-Einheiten (zum Beispiel eine Applikation, eine VM) können über verschiedene Interfaces hinweg migrieren, ohne ihre Identität zu ändern oder getroffene Spezifikationen zu verletzen. In dem Maß, wie globale Service-Definitionen nicht mehr auf alle Interfaces und alle Interface-Lokationen gemappt werden müs-

sen, soll eine Vereinfachung des Netzwerk-Betriebs erreicht werden.

Die SDN Control Plane nutzt zur Kommunikation zwischen Controller und Netzkomponente typischerweise ein spezielles Kontroll-Protokoll, vergleichbar der Signalisierung in Kommunikations-Lösungen.

Die Vision ist nun: Netzwerk-Administratoren können anhand der logischen Netzwerk-Abstraktion das Verhalten des Gesamt-Netzwerks programmieren, anstatt die Konfiguration vieler einzelner Netz-

Dipl.-Inform. Petra Borowka-Gatzweiler leitet das Planungsbüro UBN und gehört zu den führenden deutschen Beratern für Kommunikationstechnik. Sie verfügt über langjährige erfolgreiche Praxiserfahrung bei der Planung und Realisierung von Netzwerk-Lösungen und ist seit vielen Jahren Referentin der ComConsult Akademie. Ihre Kenntnisse, internationale Veröffentlichungen, Arbeiten und Praxisorientierung sowie herstellerunabhängige Position sind international anerkannt.



werk-Komponenten anfassen zu müssen. Über den SDN Controller mit seiner zentralen Intelligenz könnte das Netzwerk-Verhalten in Echtzeit geändert werden und könnten neue Netzwerk-Dienste / -Applikationen innerhalb von Stunden oder Tagen anstelle von Wochen oder Monaten ausgerollt werden. Die logische Gesamtsicht soll zudem die effiziente Ausnutzung aller möglichen denkbaren Netzwerk-Ressourcen ermöglichen. (siehe Abbildung 1)

Eine weitergehende Sichtweise setzt auf die Applikationsschicht der ONF-Architektur.

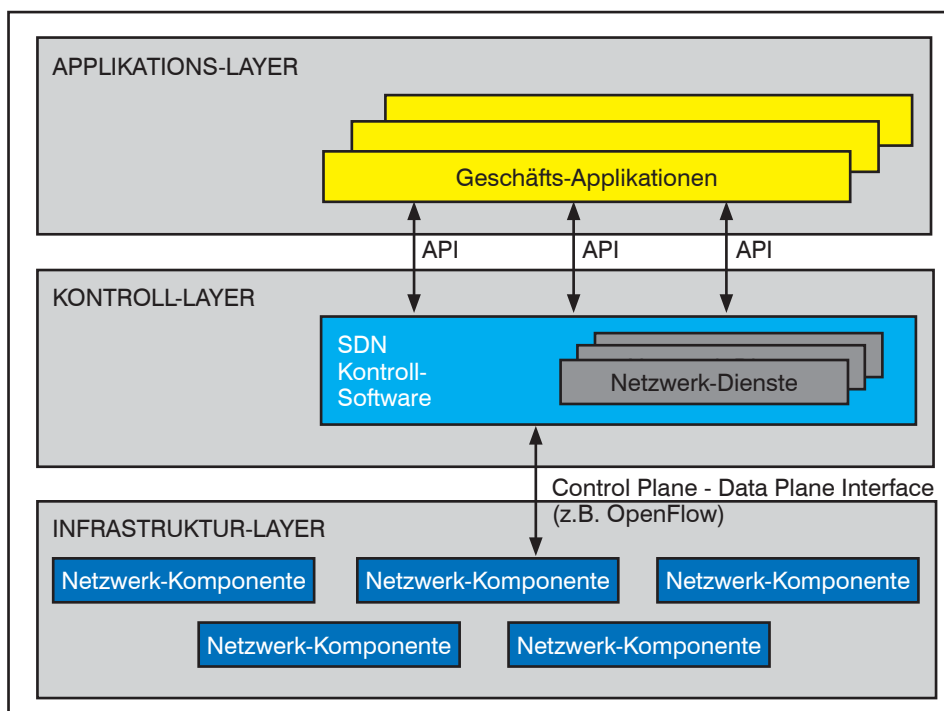


Abbildung 1: Logische Übersicht der SDN-Architektur gemäß ONF

SDN, NFV, OpenFlow und Virtualisierungs-Protokolle – Zusammenhänge, Perspektiven, Marktrelevanz

tektur noch eine Orchestrierungs-Ebene auf, die der übergeordneten Administration und Automatisierung der verschiedenen SDN-Applikationen gilt. (siehe Abbildung 2)

SDN und ONF

Anfang 2011 gründeten die Deutsche Telekom, Facebook, Google, Microsoft, Verizon und Yahoo die Open Networking Foundation als non-Profit Konsortium mit dem Ziel, Software Defined Networking (SDN) nach vorne zu treiben. Später erfolgte der Beitritt von Broadcom, Brocade, Ciena, Cisco, Citrix, Dell, Ericsson, Force10, HP, IBM, Juniper, Marvell, Microsoft, NEC, Netgear, NTT, Oracle, Riverbed Technology und VMware. Aktuell hat die ONF viele Mitglieder aus verschiedenen marktrelevanten Technologiebereichen: Provider, Netzwerk-Komponenten-Hersteller, Chip-Hersteller, Software- und System-Hersteller. Die ONF führt auch Studien zur Entwicklung offener APIs für entsprechende übergreifende Management Tools durch.

Einsatzbereiche von SDN

SDN Openflow kann ein Standard zur Handhabung aller virtualisierten Netzwerk-Lösungen werden, insbesondere auch für OpenVirtual Switch (OVS) und OpenStack. Der gesamte Komplex "komponentenübergreifendes Management" ist ein potenzieller SDN-Kandidat, da weder OSI-Management noch SNMP noch http hier umfassende Lösungen gebracht haben (ketzerische Frage: Trauen wir SDN dies nunmehr zu?). Insbesondere Automation von Management und Provisionierung stehen ganz oben auf der Liste von Providern, Data Center und Netzwerkbetreibern.

Der Service-Bereich AaaS, IaaS, Konsolidierung von Überkapazität sowie Cloudlösungen würde von übergreifenden SDN-Lösungen profitieren. Bei Einsatz von Overlay-Verfahren könnte SDN/OpenFlow Ende-zu-Ende konsistent die erforderlichen Frame- und Paket-Änderungen steuern (NAT, Tunnel-Header, Tag Rewriting etc.).

Die Entwicklung neuer Routing Protokolle als "Open Source Routing" wird zwar von den SDN-Vätern immer wieder gerne herangezogen, steht aber aktuell wohl nicht wirklich im Fokus der Problemlösung. OSPF und MPLS gut funktionierende Verfahren.

Aber selbst wenn die etablierten Layer-2 und Layer-3 Standards nicht ersetzt werden sollen, gibt es auch im Enterprise

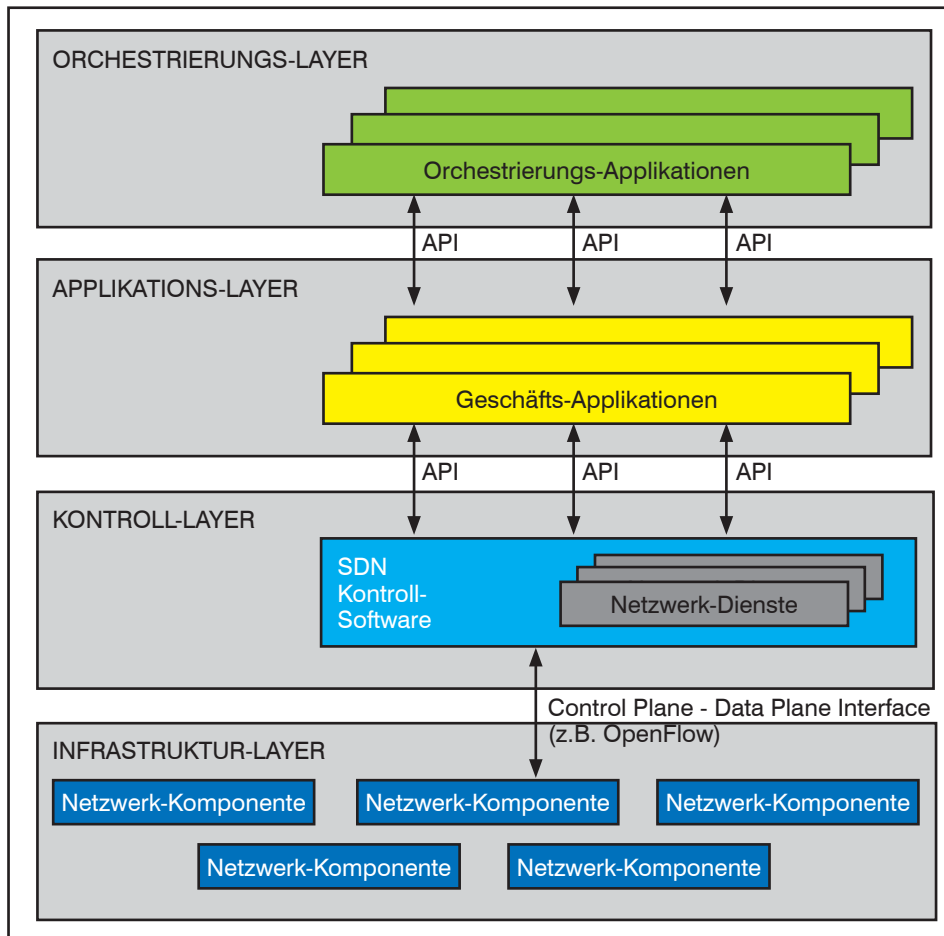


Abbildung 2: Erweiterte SDN-Architektur

Netzwerkbereich Mehrwert-Möglichkeiten mit SDN/OpenFlow: Übergreifende Lastverteilung / Traffic Engineering / netzweite QoS-Regeln als Erweiterung der L2/L3-Standardnetzwerksteuerung.

Ebenso sind Mandantentrennung und Zugangssteuerung mit netzweiten Sicherheits-Regeln und NAC, dynamische Umleitung von Flows für jede Art von IPS / IDS-Überprüfung und Erkennung von Sicherheits-Incidents sehr spannende SDN-Anwendungsbereiche. Gleiches gilt auch für netzweites Monitoring und Statistiken als Funktionserweiterungen der Netzwerksteuerung, die mit SDN/OpenFlow unterstützt werden könnten.

Abschließend ist eine Zusammenfassung von Einsatzszenarios aufgeführt, teilweise sind dies auch schon Applikationsmodule auf verfügbaren Controllern:

- Management
 - Management von OpenVirtual Switch (OVS), OpenStack
 - zentrales komponentenübergreifendes Management
 - Gemeinsame Management-Domäne für physische und virtuelle Switch-In-

frastruktur

- Programmierbarkeit von Netzkomponenten
- Automation von Management und Provisionierung
- Monitoring, Reporting, insbesondere
 - dynamische und selektive Flows bei Last-Peaks und außergewöhnlichen Lasten
 - dynamische und selektive Flows zur Fehlerdiagnose
- Mandanten
- IaaS (Infrastruktur)
- AaaS (Applikationen)
- Konsolidierung von Überkapazität (Cloud)
- Zugangs-Dienste (Carrier, dynamisch)
- Sicherheit
 - Enterprise NAC
 - Enterprise Security-Regeln
 - Triggern von Security-Incidents
 - dynamische und selektive Flow-Umleitungen zur Überprüfung, an ein IDS / IPS
- Traffic Engineering
 - Kontrolle des Überbuchungsverhältnisses
 - Bandwidth on Demand
 - Flow Klassifizierung
 - Forward zentraler Netzwerksteue-