

Schwerpunktthema

Container: Hire and Fire im Rechenzentrum

von Markus Schaub

Container wie Docker führen zu einem Paradigmenwechseln in Rechenzentren. In der guten alten Zeit vergingen Stunden, wenn nicht Tage oder Wochen vom Antrag bis zur Inbetriebnahme eines neuen Server(-dienstes). Mit der Virtualisierung hat sich diese Zeit deutlich verkürzt. Verglichen mit Containern sind virtuelle Maschinen so träge wie Bohrseln: ja man kann sie von A nach B bringen, aber gerne tut man das nicht. Container ändern nämlich die Philosophie: man betreibt keinen Mehrwert-Server mehr, sondern man nutzt Wegwerf-Anwendungen. Diese Hire and Fire Mentalität hat Auswirkungen auf das Management und das Netzwerk, da



neue Dienste schnell entstehen, sich vervielfachen, nur um dann ebenso schnell wieder zu verschwinden.

Ein Blick auf Container lohnt also nicht nur für Entwickler und Anwendungsbetreiber, sondern auch für Rechenzentrumsbetreiber und Netzwerker.

weiter auf Seite 8

Zweitthema

Session Border Controller: Die Perimeter-Komponente für All-IP - Teil 4

von Dipl.-Inform. Petra Borowka-Gatzweiler

5.4 Produktbeispiel: OpenScape Session Border Controller von Atos/Unify

Der OpenScape SBC ist für den Einsatz als zentraler SBC in der OpenScape-Lösungsumgebung vorgesehen. Nach der

Atos/Unify Empfehlung hat er dieselbe "logische" Position wie der UC-Server und wird daher am selben Standort positioniert, zum Beispiel im Kunden-RZ. Soll ein verteiltes Konzept zum Einsatz kommen, sollte hierfür die OpenScape Branch mit Session Border Controller Funktionalität

implementiert werden. Laut Hersteller-Aussage gilt: OpenScape Branch hat "im Wesentlichen" eine identische Funktionalität wie der OpenScape SBC, unterstützt jedoch weniger Sessions.

weiter auf Seite 18

Geleit

iPad Pro: kommt eine neue Flut von mobilen Endgeräten?

auf Seite 2

Aktueller Kongress

ComConsult Netzwerk Forum 2016

ab Seite 5

Standpunkt

Manchmal sind es die einfachen Dinge...

auf Seite 15

Aktuelle Sonderveranstaltung

Voice und Video im WAN

auf Seite 16

Neuer Report, Seminar und kostenloses Video

Storage-Lösungen in Unternehmen

auf Seite 17

Zum Geleit

iPad Pro: kommt eine neue Flut von mobilen Endgeräten?

Naturgemäß spielt die erwartete Anzahl mobiler Teilnehmer eine erhebliche Rolle für die Infrastruktur-Planung. Dies betrifft auch die Frage, ob wir vor einem Wechsel von Kabel-gebundenen hin zu WLAN-basierten Endgeräten stehen und somit in Zukunft eine andere Form der Verkabelung in den Gebäuden benötigen (für die Access Points).

Nachdem die lange diskutierte Verdrängung der Desktop PCs durch Tablet Computer bisher nicht stattgefunden hat, stehen wir jetzt nach der Vorstellung des iPad Pro wieder mitten in dieser Diskussion. Und Apple ist nicht gerade zurückhaltend in der Klassifizierung des iPad Pro als PC-Killer. Gleichzeitig hat sich Apple mehr dem Unternehmens-Markt zugewandt, die Zusammenarbeit mit IBM ist ein Beispiel dafür (auch wenn Apple die Weiterentwicklung seiner Laptops und Desktop-Geräte sträflich vernachlässigt und offenbar hier nur noch einen eher unwesentlichen Zielmarkt sieht).

Die Frage ist also: ist die Erwartung einer neuen Welle von Pro-Anwendungen in Kombination mit dem iPad und speziell dem iPad Pro realistisch?

Als Basis für die Antwort soll die Produktstrategie von Adobe dienen. Die Zusammenarbeit zwischen Apple und IBM hat bisher nicht ausreichend überzeugende Ergebnisse gebracht. Gleichzeitig sieht Adobe in der Integration von Tablets einen wesentlichen Teil seiner Zukunfts-Strategie.

Um gleich mit der Tür ins Haus zu fallen, wollen wir mit den drei wesentlichen Herausforderungen starten, die Pro-Anwendungen auf einem iPad haben:

1. Das iPad ist nicht die ideale Basis für monolithische Mega-Anwendungen. Die Bedien-Schnittstelle muss übersichtlich und mit Touch benutzbar bleiben.
2. Weder die Hardware eines iPads noch das Betriebssystem IOS sind eine ideale Basis für Pro-Anwendungen.
3. Egal wieviel Pro das iPad liefern wird, eine Desktop-Integration wird unverzichtbar sein.

Da der zweite Punkt vielleicht etwas überraschend kommt, wollen wir mit ihm beginnen. Fangen wir bei der Hard-



ware an. Apple bewirbt die Leistung als Desktop-Leistung und kombiniert das mit einem recht guten Bildschirm. Von der Performance her stimmt das, aber was den Bildschirm angeht zeigt dieser im direkten Vergleich zu einem kalibrierten Desktop-Bildschirm zum Beispiel mit einer Spiegel-Anwendung wie Astropad deutliche Abweichungen. Aber das ist nicht das Kernproblem. Das Kernproblem ist, dass ein Tablet im verfügbaren Speicher doch sehr begrenzt ist. Das betrifft sowohl Hauptspeicher als auch Datei-Speicher. Daran ändern auch die in-

zwischen verfügbaren 256 GB nichts. Ich habe die Tests mit der Adobe-Produktfamilie mit einem meiner Fotos gemacht. Das Ausgangsfoto hat eine Dimension von 13343 x 5561 Pixel und ist 3,7 GB groß, ein Panorama erstellt mit einer 37 Megapixel Kamera. Wäre das nicht genug als Beispiel für die Grenzen von iPads, dann soll an dieser Stelle auch auf die nach wie vor fehlende Möglichkeit des Imports von Raw-Fotos hingewiesen werden.

Aber das gravierendste Problem liegt nicht in der Hardware, sondern in sehr kritischen Eigenschaften von IOS. Und der Kern allen Übels ist das Sandboxing. Sandboxing ist ein nettes Sicherheits-Feature für kleine und überschaubare Tablet-Anwendungen. Aber im Pro-Bereich haben wir definitiv die Anforderung, dass mehrere Anwendungen an denselben Dokumenten arbeiten können müssen. Hinzu kommen Anforderungen an Plug-Ins und eine direkte Zusammenarbeit zwischen Applikationen verschiedener Hersteller. Zum Beispiel sehen wir immer wieder die Anforderung nach Adobe Photoshop auf einem iPad. Diese Anforderung ist aber komplett sinnlos. Es gibt kaum einen Photoshop-Anwender, der Photoshop ohne Plug-Ins benutzt. Hinzu



Abbildung 1: Die neue Welt der Adobe Apps

iPad Pro: kommt eine neue Flut von mobilen Endgeräten?

kommt, dass Multi-Layer-Photoshop-Dokumente eine erhebliche Größe haben können und nicht wirklich gut zu einem iPad passen.

Adobe hat sich dieser Herausforderung angenommen und eine Lösung geschaffen, die wohl in der Qualität der Architektur und auch der Leistung im Moment von keinem anderen Hersteller auch nur annähernd erreicht wird. Die Frage ist jetzt, ob diese Top-Qualität der Lösung ausreicht um die Mängel des iPads auszugleichen.

Die Eckpfeiler der Adobe-Lösung sind:

1. Eine Integration des iPads mit dem Desktop über Lightroom in Kombination mit der Adobe Cloud. Der Anwender muss dabei zwingend ein Abo der Creative Suite und eine Adobe ID haben.
2. Lightroom Mobile übernimmt dabei die Aufgabe der lokalen Speicherung auf dem iPad, so dass andere Adobe Anwendungen/Apps darauf zugreifen können.
3. Da eine monolithische Anwendung auf einem iPad keinen Sinn macht, geht Adobe den Weg vieler kleiner und spezialisierter Anwendungen, die alle einfach mit Touch zu bedienen sind.
4. Um das Problem des Speicherplatzes zu lösen, arbeitet Adobe mit Proxies (Smart Previews in Form von reduzierten DNG-Versionen des ursprünglichen Dokuments), die die Basis der Integration zwischen Desktop und iPad sind (und das Format, das auf dem iPad zum Einsatz kommt).
5. Alle Apps auf dem iPad können die bearbeitete Version eines Dokuments über die Adobe Cloud zurück zum Desktop bewegen. Also kann ein Foto, das mit Adobe Fix bearbeitet wurde, direkt auf dem Desktop in Photoshop geöffnet werden (Photoshop wird dabei automatisch gestartet). Gleiches gilt für Layout-Lösungen mit InDesign.

Die Bewertung der Qualität der Lösung hängt stark von der Anwendung ab. Was Adobe hier für InDesign als App geschaffen hat, ist absolut genial und speziell für Layout-Entwürfe eine perfekte Basis. Die Arbeit an Fotos ist ok solange man sie nicht wieder zurück zum Desktop und dort in Photoshop bringen will. Das ist zwar möglich, auch inklusive einer Multi-Layer-Bearbeitung. Aber natürlich enden die Proxies dann auf dem Desktop und Photoshop hat keine Möglichkeit das Ursprungs-Dokument wie-

der zu integrieren (zumindest habe ich keine gefunden).

Ich will das hier im Rahmen eines Geleits nicht vertiefen. Wer hier tiefer einsteigen will, der sei auf ein Interview zwischen Scott Kelly und dem Entwicklungsleiter für mobile Apps von Adobe verwiesen (siehe Kelby-One).

Kommen wir damit zurück zu der Frage ob wir vor einer Welle von professionellen Anwendungen auf iPads stehen und dies Auswirkungen auf unser Infrastruktur-Design haben wird.

Dazu folgende Antworten:

1. Ohne Frage wird es professionelle Anwendungen zum Beispiel im medizinischen Bereich geben, die eine neue Form von Nutzung schaffen werden. Aber diese werden individuell sein und müssen entsprechend auch einer sehr individuellen Infrastruktur-Bewertung unterworfen werden. Speziell für Krankenhäuser könnte das eine erhebliche Herausforderung generieren, da hier ggf. ein Bedarf für sehr hohe Bandbreiten besteht und die vorhandene WLAN-Infrastruktur inklusive der Verkabelung der Access Points dem nicht gewachsen sein können.
2. Was Adobe hier geschaffen hat, ist schlicht herausragend (und man mün-

kelt über weitere Apps, die in der Pipeline sind). Aber trotz der Qualität der Lösung kann Adobe die inhärenten Nachteile eines iPads nicht aufheben. Aber es entstehen neue Arbeitsabläufe, die Sinn machen. Dabei entsteht auch durchaus eine neue Form von Kreativität. Aber diese ist immer als mobile Ergänzung zum Desktop zu sehen. In dieser Kombination bzw. Integration aus beiden Geräte-Klassen entsteht ein ernstzunehmender Mehrwert.

3. Die Schlussfolgerung muss entsprechend sein, dass ein iPad auch wenn man Pro in den Namen aufnimmt, kein Ersatz für einen professionellen Desktop-Arbeitsplatz ist. Es ist ein möglicher Ersatz für einfach ausgestattete Office-Arbeitsplätze, aber wo sollte die Motivation und die wirtschaftliche Begründung für einen solchen Schritt liegen?

Also können wir für die Infrastruktur-Planung aufatmen? Das wird in großen Zügen von Microsoft abhängen. Die Lösungen, die Microsoft mit Windows 10 und den diversen Hardware-Ausprägungen liefert, sind architektonisch weitergehend als das was Apple macht. Allerdings ist Microsoft in der Umsetzung schlicht schlampig. Was man sich hier mit dem Surface Book an Qualitäts-Problemen geleistet hat und was man auf der Seite der Integration des Touch-Interfaces leis-

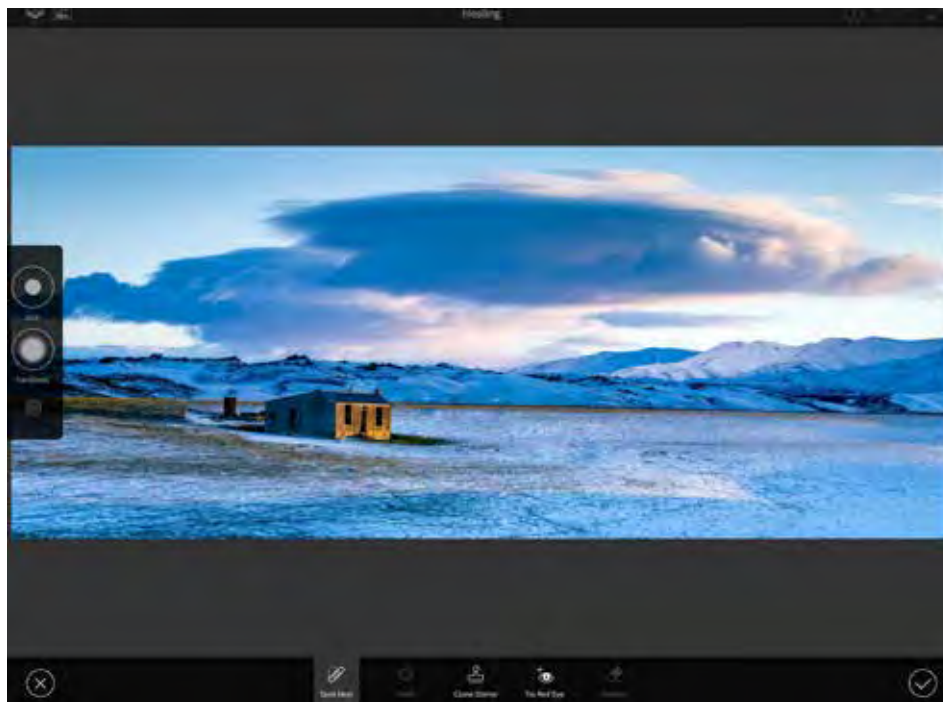


Abbildung 2: Photoshop Fix auf dem iPad mit einem Foto, das über Lightroom Mobile von einem Desktop-Rechner übernommen wurde, hier mit einer Healing-Brush im Einsatz. Aus dem 13343 x 5561 Pixel großen Foto wird dabei ein Proxy mit 2048 x 854 Pixel.

iPad Pro: kommt eine neue Flut von mobilen Endgeräten?

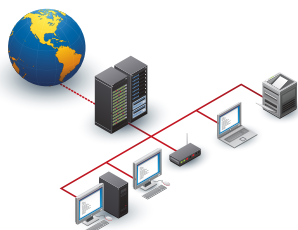
tet, ist eher als bescheiden einzustufen. Aber da Microsoft ja historisch häufig eine Weile gebraucht hat, um eine ausgereifte Lösung zu bringen, ist hier vielleicht noch Hoffnung gegeben (wer kann sich noch ein PowerPoint 1.0 erinnern? und was ist daraus geworden?). Dies gilt auch für die Hardware-Hersteller, die inzwischen im Design der Geräte sicher auf Apple-Niveau aufgeholt haben. Trotzdem sind aber viele der angebotenen Geräte nicht rund in ihrer Leistung. Dann noch eine Anmerkung zum iPad Pro zum Schluss. Das Gerät wird zwar

speziell durch die Eigenschaften von IOS zu sehr eingeengt, um für bestimmte professionelle Anwendungen gut geeignet zu sein. Das ändert aber nichts daran, dass es ein gutes Tablet ist (vielleicht das beste im Moment kaufbare, das hängt ein wenig davon ab wie man das Surface Tablet sieht). Und die Kombination mit der Astropad App als Konkurrenz zu einem Wacom Cintiq ist schon gelungen (und wer damit etwas anfangen kann, für den relativiert sich auch auf einen Schlag der Preis völlig). Auch das Lesen von komplexen Dokumenten

mit vielen Grafiken ist jetzt endlich gut möglich. Das Gerät ist also in sich durchaus gut. Aber es ist keine Basis für eine Flutwelle von Pro-Anwendungen, die den Desktop überflüssig machen werden. Wenn Apple den Desktop-PC verdrängen will, muss es auf der IOS-Seite deutlich mehr leisten und vor allem die Anforderungen von Pro-Anwendungen wirklich ernst nehmen.

Ihr
Dr. Jürgen Suppan

Kongress



ComConsult Netzwerk Forum 2016 18.04. - 21.04.16 in Königswinter

Das ComConsult Netzwerk Forum 2016 stellt die momentan dominantesten Netzwerk-Themen in den Mittelpunkt der Veranstaltung: Neue Technologien und IT-Architekturen, Netzwerk-Design, WLAN-Design und Sicherheit in Netzwerken. Drei Vortragstage und ein optionaler Vertiefungstag bilden den perfekten Rahmen um effizient und kompakt auf den neuesten Stand zu kommen. Das ComConsult Netzwerk Forum 2016 ist die herausragende Veranstaltung im Jahr 2016. Wie immer ein Treffpunkt der Branche und schlicht der beste Ort um in kürzester Zeit auf den neuesten Stand zu kommen.

Aus der Sicht von ComConsult Research werden die folgenden Fragen die Entwicklung unserer Netzwerke in den nächsten Jahren dominieren:

1. Brauchen wir mehr Intelligenz und Service-Orientierung in unseren Netzwerken als bisher? Wann ja, wo kommt sie her?
2. Müssen wir dynamisch wachsende Rechenzentren mit mehr dynamisch skalierenden Netzwerken begleiten? Wenn ja, wie können Netzwerke wirtschaftlich dynamisch skalieren?
3. Welche Bandbreiten brauchen wir wo, warum und wann? Ist die Zeit von Scale up vorbei und brauchen wir intelligentere Konzepte um ein Optimum aus Preis, Verfügbarkeit und Leistung zu erreichen? Welche Auswirkungen haben die neuen WLAN-Standards im Access Bereich?

Diese Fragen sind ebenso dominant wie komplex. Ihre Beantwortung erfordert die Auseinandersetzung mit Technologien wie Fabrics, SDN, ACI, NSX, IPv6, WLAN, 25, 50, 100, 400 Gbit/s Ethernet. Die Bewertung des Bedarfs erfordert eine genaue Analyse der neuesten Anwendungs-, Speicher- und Server-Architekturen. Und das ganze muss zudem immer sicherer werden. Sicherheit wird zum Schlüsselkriterium, kann aber nicht losgelöst von der Netzwerk-Architektur gesehen werden. So erfordern dynamisch skalierende Netzwerke eine dazu passende dynamisch skalierende Sicherheits-Lösung.

Die Anforderungen an Netzwerk-Infrastrukturen waren noch nie so komplex und gleichzeitig mit großen Fragezeichen der Wirtschaftlichkeit versehen.

Hier setzt das ComConsult Netzwerk Forum 2016 an:

- wir analysieren wo der Bedarf herkommt und wer davon betroffen ist
- wir bewerten die dominanten Lösungstechnologien
- wir geben Empfehlungen zu den anstehenden Investitionen und deren Wirtschaftlichkeit

Moderatoren: Dipl.-Inform. Petra Borowka-Gatzweiler, Dipl.-Math. Cornelius Höchel-Winter, Dr.-Ing. Behrooz Moayeri

Preis: € 2.590,- netto 4 Tage
€ 2.390,- netto 3 Tage
€ 990,- netto 1 Tag



Buchen Sie über unsere Web-Seite
www.comconsult-akademie.de

Die Top 3 Netzwerk-Themen der nächsten 5 Jahre

ComConsult Netzwerk Forum 2016

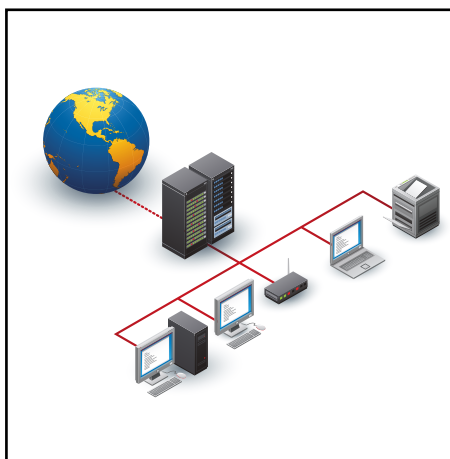
18.04. - 21.04.16 in Königswinter

Die ComConsult Akademie veranstaltet vom 18.04. bis 21.04.16 ihr "ComConsult Netzwerk Forum 2016" in Königswinter.

Das ComConsult Netzwerk Forum 2016 stellt die vier momentan dominantesten Netzwerk-Themen in den Mittelpunkt der Veranstaltung:

- Neue Technologien und IT-Architekturen: wie müssen sich Netzwerke ändern?
- Netzwerk-Design: skalierbare Kapazitäten, Service-orientiert und sicher
- WLAN-Design mit 802.11ac Wave 2 und die saubere Integration ins LAN-Design
- Sicherheit in Netzwerken: Zertifikate und NAC

Am ersten Tag analysieren wir u.a. ob wir zentral gesteuerte Netzwerk-Lösungen brauchen. An einer Reihe ausgewählter Anwendungsbeispiele wird untersucht, ob eine ausgelagerte Data-Plane den Betrieb, die schnelle Bereitstellung und die Gestaltung unterstützt oder ob das Ganze zu komplex wird. Hintergrund dazu ist die Frage, wie man Overlay-Netzwerke am besten konfigurieren und betreiben kann (verbunden natürlich mit der Frage, ob man sie überhaupt braucht).



Diese ergänzen wir um die praktische Frage nach der Zukunft des WAN: können sich WANs gegenüber dem Internet durchsetzen?

Am zweiten Tage steht Netzwerk-Design mit allen neuen Technologien im Vordergrund:

- Network Function Virtualization
- SDN
- 25/50/100: neue Bandbreiten für wen?
- Trill kontra Fabricpath kontra SPB kontra VXlan

- Layer 3 Design mit modernsten Technologien: wo stehen wir?

Der dritte Tag stellt zwei Sonderthemen in den Vordergrund, die in allen aktuellen Projekten eine tragende Rolle spielen und auch speziell das Jahr 2016 bestimmen werden:

- WLAN-Design nach 802.11ac Wave 2 und seine Integration in LAN-Design
- Sicherheit mit Zertifikaten und NAC

Das ComConsult Netzwerk Forum 2016 ist das richtige Forum zur richtigen Zeit.

Wir analysieren exklusiv für Sie:

- welche neuen Technologien und Produkte stehen für bessere und wirtschaftlichere Netzwerke zur Verfügung?
- wie verändern sich Anforderungen an Netzwerke?
- wie verändert sich Netzwerk-Design und wie können Sie die Vorteile zu Ihren Gunsten nutzen ohne das gesamte Netzwerk ablösen zu müssen?

Unser Vertiefungstag in diesem Jahr dreht sich komplett um IPv6 und die aktuellen Projekterfahrungen in diesem Bereich.

Fax-Anmeldung an ComConsult 02408/955-399

ComConsult Netzwerk Forum 2016

Ich buche den Kongress
ComConsult Netzwerk Forum 2016
18.04. - 21.04.16 in Königswinter

☐ 18.04. - 21.04.16 in Königswinter
zum Preis von € 2.590,- netto - 4 Tage

☐ 18.04. - 20.04.16 in Königswinter
zum Preis von € 2.390,- netto - 3 Tage

☐ Bitte buchen Sie mir ein Hotelzimmer

 Buchen Sie über unsere Web-Seite
www.comconsult-akademie.de

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift

Programmübersicht - ComConsult Netzwerk Forum 2016

Montag, den 18.04.2016 - IT-Architekturen und neue Technologien

9:30 bis 10:30 Uhr

Die Top-Themen 2016

- Warum sich SDN in Unternehmensnetzen nicht durchsetzt (Unterschiede zwischen Unternehmens- und Hyperscaler-Netzen)
- WANs unter zunehmendem Einfluss der Entwicklungen im Internet
- Risiken des Aufschiebs der IPv6-Einführung

*Dr. Behrooz Moayeri,
ComConsult Beratung und Planung GmbH*

10:30 bis 11:30 Uhr

Cloud Computing: Einsatz im Unternehmen

- Anspruch vs. Marketing: Was ist Cloud Computing eigentlich?
- Cloud-Produkte im Unternehmenseinsatz:
 - Nutzbarkeit: Cloud-Produkte sind „anders“
 - Anforderungen an die Netzwerke und die Infrastruktur
 - Wo liegen Nutzungsgrenzen und typische Probleme
- Erfahrungen aus konkreten Projekten

*Dipl.-Math. Cornelius Höchel-Winter,
ComConsult Research GmbH*

11:30 Uhr Kaffeepause

12:00 bis 12:45 Uhr

Internet of Things

- Was ist IoT / Industrie 4.0
- Anwendungsbereiche, Einsatz-Szenarien
- Architektur und Protokolle
- Wo steht die Standardisierung?
- IoT Roadmap der nächsten Jahre

*Dipl.-Inform. Petra Borowka-Gatzweiler,
UBN*

12:45 Uhr Mittagspause

14:15 bis 15:00 Uhr

Wandel der Netzwerkarchitekturen in Zeiten von SDN

- Private Cloud und Hybrid Enterprise verändern die Anforderungen
- Von SDN zu SDx, sind Sie bereit dafür?
- Separierung von Underlay und Overlay Netzmanagement
- Warum Layer 3 Underlay Designs
- Für Enterprise braucht man mehr – Layer 2 over Layer 3 mit VXLAN
- VXLAN Control Plane Optionen – mit (SDN) Controller und ohne
- Wie passen VMware und OpenStack ins Bild

*Dipl.-Ing. Markus Nispel,
Extreme Networks GmbH*

15:00 bis 15:45 Uhr

Docker: Fluten Container bald das RZ?

- Was sind Container und wie funktionieren sie?
- Wie unterscheiden sich Container von klassischen Virtualisierungstechniken?
- Was sind die Vor- und Nachteile? Was werden die typischen Anwendungsgebiete von Containern sein?
- Was sind die Konsequenzen für das Netzwerk von Rechenzentren?
- Sind Container in der Cloud?

*Markus Schaub,
ComConsult-Study.tv*

15:45 Uhr Kaffeepause

16:15 bis 17:15 Uhr

**Cloud-Computing:
der rechtliche Rahmen und die Herausforderungen**

*Dr. Fabian Niemann,
Bird & Bird LLP*

ab 18:00 Uhr Happy Hour im Foyer

Dienstag, den 19.04.2016 - LAN-Design: Planung, Betrieb

9:00 bis 10:00 Uhr

SDN/NFV

- Wo steht der SDN-Markt?
- Was ist NFV? (Architektur, Einsatzszenarien, Marktbedeutung)
- Abgrenzung und Überlappung von NFV und SDN
- NFV und Network Services (Service Chaining mit NSH)

*Dipl.-Inform. Petra Borowka-Gatzweiler,
UBN*

10:00 bis 11:00 Uhr

Netzdesign im Vergleich

- Layer 3 Design mit z.B. BGP
- Layer 2 Design mit SPB
- Layer 4 Design mit z.B. QUIC
- Lösungen wie NSX, ACI oder OpenFlow

*Markus Geller,
ComConsult Research GmbH*

11:00 Uhr Kaffeepause

11:30 bis 12:30 Uhr

Erfahrungen mit IPv6 bei BMW

- Motivation
- Vorgehensweise
- Herausforderungen
- Erfahrungen / Probleme
- Status und Ausblick

*Dipl. Ing. Bernhard Haring,
BMW AG*

12:30 Uhr Mittagspause

14:00 bis 14:45 Uhr

Architektur im Rechenzentrum - 25, 50 und 100G

- Einführung in eine neue Generation von offenen und skalierbaren RZ Switchen
- 25G, die neuen 10G? • 50G, die neuen 40G für Storage?
- 100G, der neue 40G Interconnect?
- Anwendungsfälle für Rechenzentren
- Remote Direct Memory Access über Converged Ethernet (RoCE) in der Praxis

*Arne Heitmann,
Mellanox Technologies Ltd.*

14:45 bis 15:30 Uhr

Fabrics kontra Standard-Design an Projektbeispielen

*Heinz Behrens,
Avaya GmbH & Co KG*

15:30 Uhr Kaffeepause

16:00 bis 17:00 Uhr

40 Gigabit-Ethernet und mehr: Auswahl zukunftssicherer Schnittstellen und der optimalen Verkabelung

- Simplizität der alten und Komplexität der neuen physikalischen Schnittstellen
- Schnittstellenvielfalt der Switch-Hersteller
- Unbeachtete Abhängigkeiten zwischen Elektronik und Verkabelung
- MPO war gestern, LC ist heute! Ist das so?
- Unbekannte Modul-Inkompatibilität der verschiedenen Datenraten
- Die universelle Verkabelung für alle Datenraten

*Dipl.-Ing. Hartmut Kell,
ComConsult Beratung und Planung GmbH*

Programmübersicht - ComConsult Netzwerk Forum 2016

Mittwoch, den 20.04.2016 - WLAN-Design: Planung und Betrieb / Sicherheit

9:00 bis 10:00 Uhr

Neue WLAN-Techniken und ihr Einfluss auf Enterprise WLANs

- DCF: „Pest“ oder Segen für die Entwicklung des WLAN?
- Die dritte Welle der WLAN Chips rollt auf uns zu! Wie profitieren Enterprise WLANs davon? • MU-MIMO ist angeblich DER Schlüssel zu höherer Performance. Was ist an dieser Behauptung dran?
- Warum IEEE 802.11ac eigentlich KEIN "Gigabit WLAN" ist und es auch nie werden wird!
- Welche Anwendungen brauchen überhaupt WLAN mit mehr als 1 Gigabit/s?
- Ausblick: Parallelität wird (mal wieder) die Kapazität erhöhen

Dr. Joachim Wetzlar,

ComConsult Beratung und Planung GmbH

10:00 bis 11:00 Uhr

Wireless, aber richtig!

Von echtem Multi-Gigabit zu LTE-Erweiterungen auf dem Weg zu 5G

- Megatrend Mobilität: Status und Wachstum
- IEEE 802ad reloaded: Änderungen gegenüber der Version von 2010
- Echtes Multi-Gigabit mit 802ad/WiGig im 60 GHz-Bereich, Produktlage
- LTE Rel. 13 und LTE Advanced: Carrier Aggregation, HetNets und LTE/WiFi Interworking
- Gefährdungen durch LTE in lizenzfreien Bändern, LAA, LTE-U, MuLTEfire
- Was 3GPP schon heute für 5G vorbereitet

Dr. Franz-Joachim Kauffels,

Technologie- und Industrie-Analyst

11:00 Uhr Kaffeepause

11:30 bis 12:30 Uhr

WLAN in der Praxis: Ein WLAN für alle Umgebungen, Nutzertypen und Anwendungsfälle

- WLAN Infrastruktur: Indoor, Outdoor, Remote Office, Mesh, ZeroTouch-Provisioning, Beacons, Analytics – Was brauche ich wo?
- Beacon-Beispiel im Enterprise: automatische Konferenzraumerkennung
- Sichere Integration verschiedener Nutzertypen: Mitarbeiter, BYOD, Gäste, IoT-Devices

- IT definiert Regeln und Benutzer nutzt Self-Service Abläufe
- Firewall im Perimeter, Datacenter oder direkt im User-Access
- Applikationserkennung und Web Reputation im Access
- Nahtlose Integration von Wired-Access in das System

Reinhard Lichte,

Aruba - a Hewlett Packard Enterprise Company

12:30 Uhr Mittagspause

14:00 bis 14:45 Uhr

Fallstricke und Best Practice bei NAC

- Warum IEEE 802.1X immer noch ein Alptraum sein kann
- Best Practice NAC: Wie NAC erfolgreich umgesetzt und betrieben werden kann
- Welches Sicherheitsniveau mit NAC überhaupt geschaffen werden kann
- Ist MACsec eine Alternative?
- Evolution von NAC: Von Advanced Monitoring über Profiling bis hin zur Abwehr zielgerichteter Angriffe

Dipl.-Inform. Daniel Prinzen,

ComConsult Beratung und Planung GmbH

14:45 bis 15:30 Uhr

Sichere Kommunikation im Netz mit Zertifikaten: Alptraum oder etablierte Technik?

- Von NAC über Web-Anwendungen bis zum SSL-VPN: Anwendungen von Zertifikaten zur sicheren Kommunikation
- Fallstricke Schlüsselmanagement und Vertrauensketten: Welche Sicherheitsvorfälle es gab und was wir dagegen tun können
- Certificate Pinning und Certificate Transparency: Warum das Konzept der Vertrauensketten dringend renoviert werden musste

Dr. Simon Hoff,

ComConsult Beratung und Planung GmbH

15:30 Uhr Ende der 3-tägigen Veranstaltung - Kaffeepause für Teilnehmer der 4-tägigen Veranstaltung

Donnerstag, den 21.04.2016 - Optionaler Zusatztag "IPv6"

ab 9:00 den ganzen Tag

IPv6 Migration: Projektvorbereitung und Umsetzung

- Organisation eines IPv6 Rollouts (Planung des Vorgehens, was wann entschieden werden, welche Abteilungen sind in welcher Projektphase gefordert, wo existiert Schulungsbedarf)
- Adresskonzept (Welche Alternativen stehen zur Verfügung, was sind die Vor- und Nachteile)
- Zuweisung von IPv6 Adressen (Welche Verfahren stehen zur Verfügung, wie integriert man Komponenten, die kein DHCPv6 unterstützen)
- Anforderungen an Netzwerk- und Infrastrukturkomponenten (Erstellung von Anforderungsprofilen für einzelne Komponenten, Testdurchführung, ausgewählte Testergebnisse)
- LAN-Architektur (Redundanzverfahren: VRRP, HSRP, Routing von IPv6, Umgang mit QoS bei IPv6)

- Migration der Internetpräsenz
- Migration von Anwendungen und Appliances
- Erstellung eines Anforderungskataloges für die Anschaffung von Hard- und Software
- Externe Anbindungen (WAN, Internet, Internet-VPN, Externe Partnerunternehmen)
- Security (Ergebnisse von Proxy-Tests, Firewalls & IDS, First-Hop-Security)

Markus Schaub,
ComConsult Study.tv

10:30 Uhr Kaffeepause

12:45 Uhr Mittagspause

15:30 Uhr Ende der 4-tägigen Veranstaltung

Folgende Aussteller nehmen bisher an der Ausstellung teil:



Schwerpunktthema

Container: Hire and Fire im Rechenzentrum

Fortsetzung von Seite 1



Markus Schaub ist seit 2009 Leiter von ComConsult-Study.tv. Er verfügt über umfangreiche Berufserfahrung in den Bereichen Netzwerken und VoIP. Seine Schwerpunkte liegen im Netzwerk-Design, IP-Infrastrukturdiensten und SIP, zu denen er viele Vorträge auf Kongressen hielt, erfolgreich Seminare durchführte und zahlreiche Veröffentlichungen schrieb.

Bevor man sich mit den Auswirkungen des Einsatzes von Containern beschäftigen kann, muss man zunächst die Grundzüge der Technik verstehen.

Im Grunde sind Container nicht so neu, wie viele womöglich meinen. Die zugrundeliegenden Funktionen gibt es bereits länger und auch Container selbst sind nicht neu. Dass sie nun in aller Munde sind, liegt an dem Unternehmen Docker. Mittels eines guten Marketings und immensen Summen Risikokapitals hat es die Container aus der Nerd-Ecke der Linux-Entwickler geholt und salonfähig gemacht. Das geht so weit, dass man heute Docker und Container oft gleichsetzt. Darum beschäftigt sich dieser Artikel auch nicht mit Containern im Allgemeinen, sondern mit Docker-Containern im Besonderen.

Die Idee

Hat man beispielsweise eine Webseite entwickelt, so kann man die nicht ohne weiteres von einem Webserver zu einem anderen Webserver überspielen. Vielmehr muss der „neue“ Webserver die von der Anwendung benötigten Module unterstützen. Benötigt man bspw. eine Verbindung zu einer MySQL Datenbank, so muss das entsprechende Servermodul installiert und aktiviert sein. Viel Zeit und noch mehr Nerven sind notwendig, um eine Anwendung von einem Entwicklungssystem auf ein Produktivsystem zu übertragen.

Die Idee hinter Container war es, Anwendungen unabhängig vom installierten Betriebssystem zu entwickeln, damit man sie einfach von einem auf ein anderes System portieren kann. Dafür ist es notwendig, dass neben der Anwendung auch deren Abhängigkeiten von ihrer Systemumgebung mit portiert werden. Damit sind neben benötigten Diensten auch Libraries,

Module und ähnliches gemeint. Im Beispiel eben wären das neben der Webanwendung eben auch der Webserver (Apache, nginx, Tomcat...) und die benötigten Webserver-Module. Genau das leisten Container, sieht man mal (vorläufig!?) von der Unabhängigkeit vom Betriebssystem ab. Doch dazu später mehr.

Nun könnte man auf die Idee kommen, dass virtuelle Maschinen das auch können und somit Container nichts Neues leisten. Doch das stimmt nicht: Container enthalten zwar die Anwendung und deren Abhängigkeiten, nicht jedoch das gesamte Betriebssystem (siehe Abbildung 1). Damit werden sie schneller portierbar und sind weniger Ressourcen-hungrig: anstatt ein komplettes Betriebssystem auf einen Server zu überspielen und zu starten, wird der Container gestartet. D.h. der Container wird „geöffnet“ und die darin befindlichen Dienste werden gestartet. Ein Betriebssystem muss nicht gebootet werden, da sich Container und Host denselben Kernel teilen und der Host ja bereits aktiv ist.

Abgesehen davon, dass sich Host und Container denselben Kernel teilen, sind sie jedoch gegeneinander abgeschottet. Das gilt auch für Container untereinander, die

auf demselben Host laufen. Sollen Container miteinander kommunizieren, so wird das über die Netzwerkwerkschnittstelle realisiert.

Durch diese Abschottung ist es möglich, dass man auf demselben Host Container laufen lassen kann, die auf unterschiedlichen Linux-Distributionen laufen, denn – wie gesagt – „nur“ der Kernel muss kompatibel sein. Auch ist es so möglich, Anwendungen auf demselben Host laufen zu lassen, deren Anforderungen sich gegenseitig ausschließen.

Um das zu ermöglichen, wurde keine neue Technik erfunden, sondern bestehende kombiniert, so z. B. Linux Namespaces, cgroups (control groups) und aufs (advanced multi layered unification filesystem).

Kommen wir auf den Vergleich mit der Vollvirtualisierung ala vmware zurück, so werden einige Vor- und Nachteile von Containern nun klar:

• Vorteile

Mangels eigenem Betriebssystem sind Container kleiner und schneller zu starten als virtuelle Maschinen. Auch sind sie weniger ressourcenhungrig, da sie

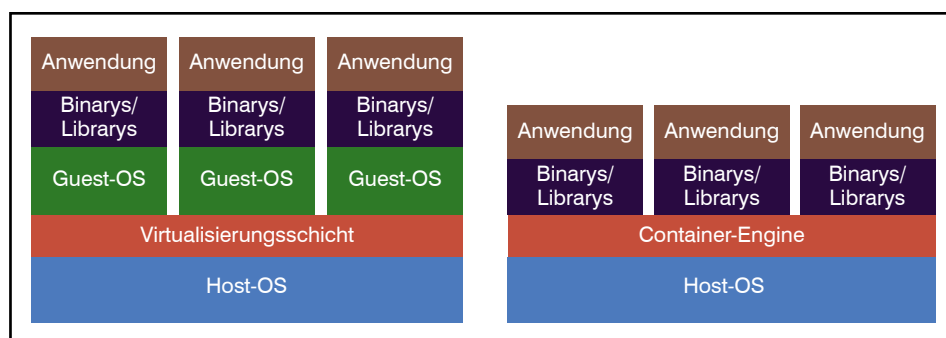


Abbildung 1: Vergleich Container - Vollvirtualisierung

Container: Hire and Fire im Rechenzentrum

sich die Kerneldienste teilen. Weil sie alle Abhängigkeiten der Anwendung beinhalten und nicht nur die Anwendung selbst, können sie schnell von einem Host auf einen anderen übertragen werden. Auch können die containerisierten Anwendungen auf mehreren Hosts gestartet werden.

- Nachteile

Zumindest vorläufig gilt noch, dass Container nicht unabhängig vom Betriebssystem sind. Da sich alle Container den Kernel mit dem Hostsystem teilen, kann man keine Linux-basierte Anwendung auf einem Windows-System entwickeln. Allerdings wird auch daran bereits gearbeitet. Auch gibt es die Möglichkeit die Anwendung in einer virtuellen Linux-Maschine zu entwickeln, die auf einem Windows oder OS X läuft. Dafür gibt es bei Docker gleich die komplette Entwicklungsumgebung auf Basis von VirtualBox.

Ein weiterer Nachteil ist, dass die Abschottung Container-Host und Container-Container natürlich nicht so komplett sein kann, wie das bei der Vollvirtualisierung der Fall ist. Solange man sich denselben Kernel teilt, gibt es nun mal Berührungspunkte.

Container und Images

Um zu verstehen, was es mit Containern auf sich hat, muss man den Unterschied zwischen Containern und Images kennen.

Ein Image ist das unveränderliche Abbild eines Containers. Dieses enthält die Anwendung und ihre Abhängigkeiten, jedoch keine Systemzustände.

Um aus einem Image einen Container zu machen, lädt man das Image auf sein System und startet es. Dabei wird das Image kopiert und mit einer Laufzeitumgebung ausgestattet (siehe Abbildung 2). Während das Image selbst somit unverändert bleibt, kann man mit dem laufenden Container arbeiten. Alle Änderungen werden in der Laufzeitumgebung gespeichert, die im Zweifelsfall Vorrang vor den Daten des originalen Image hat. Lädt man beispielsweise ein Image mit einem Apache-Webserver, startet auf dessen Basis einen Container und aktiviert das SSL-Modul, so ist zu dem Apache des laufenden Container eine HTTPS-Verbindung möglich. Nutzt man das Image erneut für einen anderen Container, so ist in dem kein SSL-Modul aktiv.

Auch Container kann man anhalten und neu starten. Dabei wird - anders als bei Images - jedoch auch die Laufzeitumgebung mitgespeichert. Startet man einen

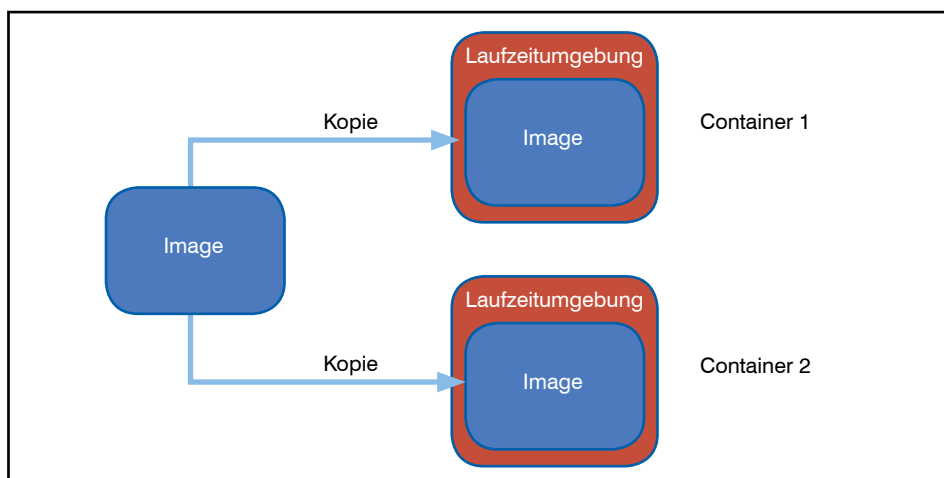


Abbildung 2: Vom Image zum Container

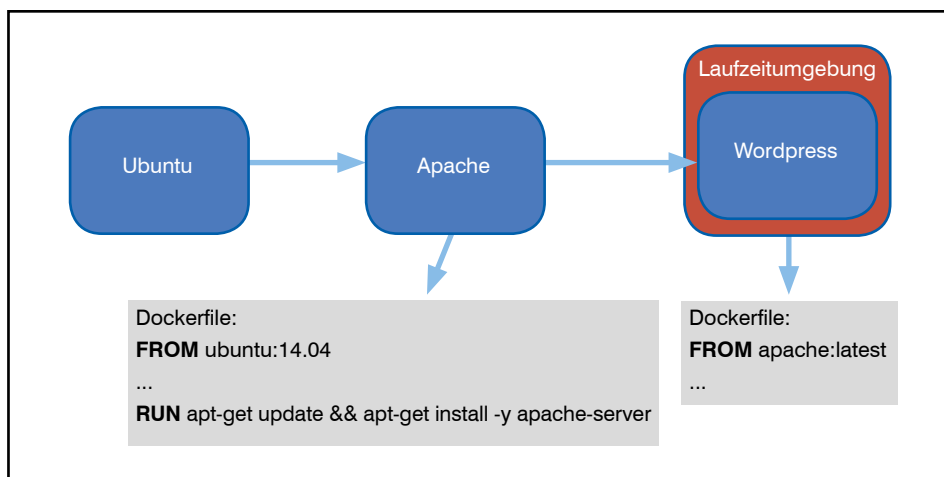


Abbildung 3: Image-Kaskade mit Dockerfile-Auszügen

angehaltenen Container, so befindet er sich in dem Zustand, in dem er angehalten wurde. Im Beispiel von oben, enthält der Container immer noch das aktivierte SSL-Modul.

Eine Besonderheit der Images ist, dass sie auf anderen Images beruhen können. Nehmen wir an, wir wollen eine Plug-In für Wordpress entwickeln. Dann laden wir ein Image, das Wordpress bereits enthält. Dieses Image wiederum beruht auf einem Image, das einen Webserver enthält, welches seinerseits auf einem Image mit den notwendigen Ubuntu-Diensten und Libraries beruht. So entsteht eine Kette von Images, auf denen die spätere Eigenentwicklung beruhen wird. Das spart viel Zeit und Arbeit bei der Entwicklung, denn der Entwickler braucht nur das Wordpress-Image zu installieren, die Abhängigkeiten werden automatisch aufgelöst und mit installiert (siehe Abbildung 3).

Natürlich ließen sich diese Abhängigkeiten auch mittels einer Paketverwaltung wie apt

auf Linux realisieren. Vorteil der virtualisierten Container-Umgebung ist jedoch, dass auf einem physikalischen System auf diese Weise auch Anwendungen laufen können, die sich widersprechende Anforderungen haben.

Die Frage ist: Wie entstehen Images? Dazu gibt es zwei Möglichkeiten:

1. Container zu Image

Man kann jeden Container in ein Image umwandeln. Das klingt plausibel, praktisch und einfach. Letzteres stimmt sogar, der Rest nicht. Denn wandelt man einen Container in ein Image um, so friert er quasi die Laufzeitumgebung inklusive aller genutzten Images ein. Wird im Beispiel von oben das Apache Image aktualisiert, so wirkt sich das auf das neu erstellte Image nicht mehr aus. Das kann gewollt sein, aber auch unerwünscht.

2. Dockerfile

Eleganter als das „Einfrieren“ eines

Container: Hire and Fire im Rechenzentrum

Containers ist es jedoch mit Dockerfiles zu arbeiten. Mittels Dockerfile ruft man das Image auf, auf das man sich bezieht und nimmt per Skriptsprache die notwendigen Änderungen vor. Das können Anpassungen des Containers sein, Updates des Betriebssystems, die Definition von Schnittstellen zwischen Host-OS und Container (bspw: TCP Ports) oder eben auch das Nachladen von Anwendungen, wie dem gerade selbst entwickelten Wordpress-Modul.

Eine weitere Eigenschaft von Containern ist es, dass sie zu jederzeit gestoppt und kopiert werden können. So kann man auf sehr einfache und effiziente Weise eine Versionshistorie bei der Anwendungsentwicklung realisieren. Hat sich der Entwickler verrannt, braucht er nur zu einer früheren Version zurückkehren. Dabei ist gewährleistet, dass eigene Änderungen aber auch vorgenommene Updates, die zur Fehlfunktion geführt haben, restlos beseitigt werden und nicht, wie sonst oft üblich noch irgendwelche „Reste“ im System verbleiben oder im Code vergessen werden. Auch hier schlägt der Wegwerfgedanke wieder durch.

Es wundert also nicht, dass viele Entwickler sich zunehmend für Container begeistern können:

- Ein passendes Image vorausgesetzt, können Sie sich viel Zeit mit dem aufsetzen einer Entwicklungsumgebung sparen.
- Sie können Ihre Anwendung einfach an den Betrieb in Form eines Images übergeben, ohne dass es durch Versionskonflikte oder fehlenden Abhängigkeiten zu Anpassungsschwierigkeiten zwischen Entwicklungs- und Produktionsumgebung kommt.
- Eigene Entwicklungen können einfach gesichert werden.
- Ein Neuaufsetzen und -anfangen ist kaum noch nötig, da man fatale Änderungen schnell „rückgängig“ machen kann.
- Images oder gestoppte Container können Entwickler untereinander tauschen, so dass sie bei Projekten auf einem gemeinsamen Stand sind, sich aber gegenseitig nicht stören, wenn mal was danebengeht.
- Folgt man der Grundidee von Containern, kann sich der Entwickler somit auf seine Anwendung konzentrieren.

Auch Betreiber entdecken den Vorzug von Containern für sich:

- Geht es um die Skalierbarkeit, so benötigt man von einer Anwendung oft nur einen Teilaspekt. Beispiel: viele Anfragen in kurzer Zeit benötigen oft nur mehr Cache-Server, nicht mehr replizierte Datenbanken. Einen Loadbalancer vorausgesetzt, kann man das Image des Cacheservers auf mehrere Hosts verteilen und dort bei Bedarf starten und später wieder stoppen. Ggf. sogar löschen, um den Platz frei zu geben.
- Redundanz ist ein weiterer Vorteil für Betreiber. Da Container auf Images beruhen, die man auf einem Hub lagert, können sie bei Ausfall eines Systems schnell von dort auf ein neues System kopiert und noch schneller gestartet werden. Natürlich ist das eher eine Hot-Standby-Lösung, da die Laufzeitumgebung fehlt und somit alle Verbindungen neu aufgesetzt werden müssen, aber für viele Anwendungen reicht das.

Zusammenfassend kann man sagen, dass sich Container zwar (noch?) eher in der Entwicklerecke tummeln, aber durchaus Potential haben, sich einen signifikanten Marktanteil bei der Servervirtualisierung zu sichern.

Darum macht es Sinn, sich als Netzwerker mit ihnen vertraut zu machen. Werfen wir darum einen Blick auf die Netzwerkschnittstellen.

Netzwerkschnittstellen

Ein Prinzip der Container ist, dass sie gegeneinander abgeschottet sind, auch wenn sie auf derselben Maschine laufen. Die Idee von Containern ist hingegen, dass komplexe Anwendungen in ihre Dienste zerlegt werden, die in unterschiedlichen

Containern laufen. Daraus folgt, dass Container über die Netzwerkschnittstelle miteinander kommunizieren. Dieser kommt also bei containerisierten Anwendungen eine erhebliche Bedeutung zu.

Installiert man Docker, so bringt es von Hause aus drei Schnittstellen mit:

- None
Ein Container, der mit dieser Option gestartet wurde, hat keine Verbindung in die Außenwelt, sondern nur die Loopback Adressen für IPv4 (127.0.0.1) und IPv6 (:::1).
- Host
Startet man einen Container mit der Option „Host“, so ist dessen IP Konfiguration identisch mit der des Host Systems.
- Bridge
Der Default, wenn man nichts angibt, ist „Bridge“. Allerdings ist „Bridge“ nicht, was man als Netzwerker jetzt vielleicht glaubt, der Host fungiert nicht als Bridge und startet auch keinen virtuellen Switch zur Außenwelt, vielmehr fungiert der Host als NAT-Gateway.

Was nämlich passiert ist folgendes: die Container bekommen vom Host automatisch IPv4 Adressen auf dem privaten Bereich zugewiesen (e.g. 172.17.X.X) und als Gateway eine IP Adresse aus demselben Subnetz, die dem Host selbst zugeordnet ist. (siehe Abbildung 4)

Will man eine Anwendung innerhalb eines Containers erreichen, muss bei dessen Start auf dem Host ein Port-Mapping erfolgen. Startet man beispielsweise eine MySQL Datenbank im Container A und gleichzeitig in Container B, so kann man beide errei-

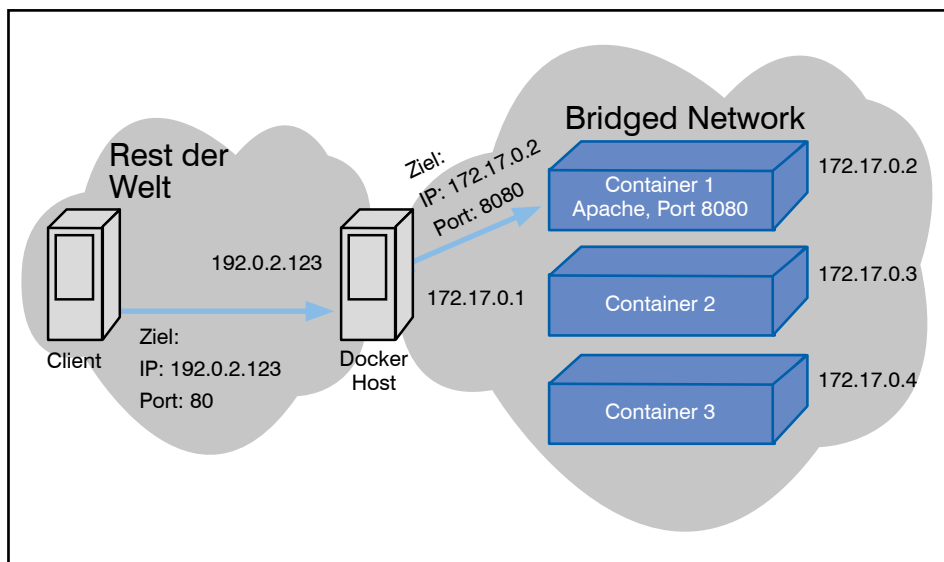


Abbildung 4: Dockers Default-Bridge mit Portmapping und NAT

Container: Hire and Fire im Rechenzentrum

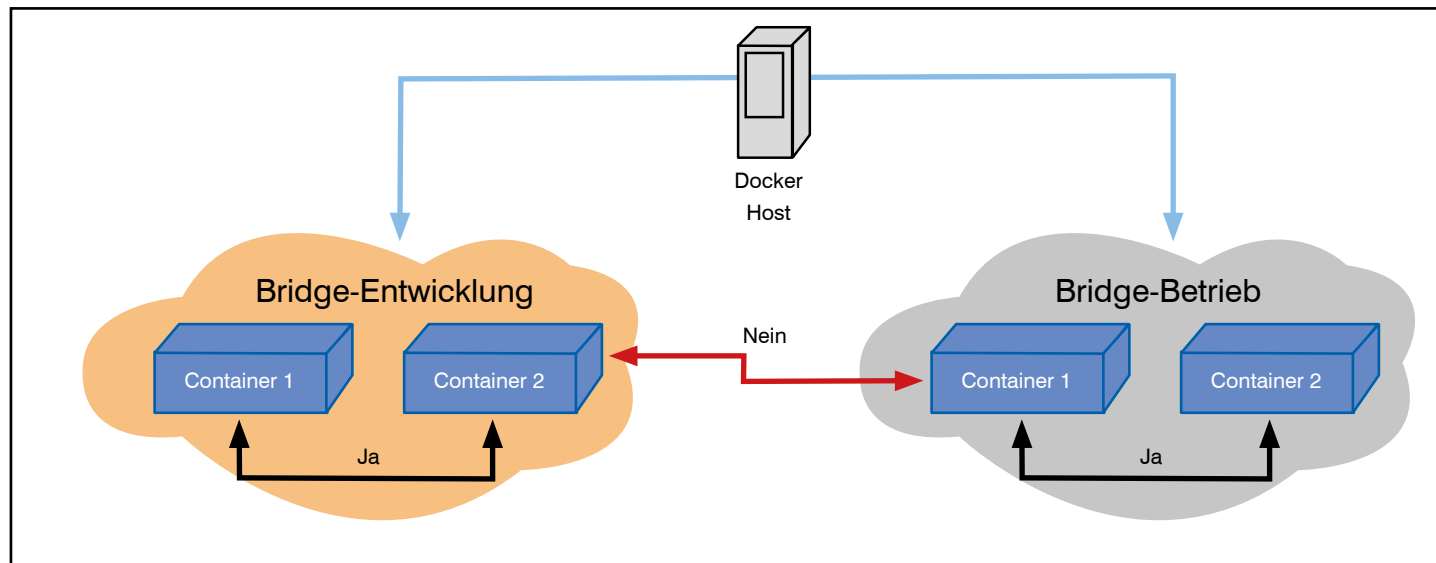


Abbildung 5: Zwei User-Defined Bridges auf einem Host

chen, wenn der Default Port 3306 von MySQL, der innerhalb beider Container genutzt wird, vom Hostsystem auf die Ports 13306 für Container A und 23306 für Container B genutzt wird.

Diese Out-of-the-Box-Varianten betrachtet selbst Docker als „historisch“, was bei einem drei Jahre alten Unternehmen schon merkwürdig anmutet. Natürlich reichen sie nicht für Anwendungen, die später mal eine Cloud bilden sollen. Darum sind weitere Netzwerk-Schnittstellen geschaffen worden. Diese heißen „User-Defined Networks“.

Die Einfachste heißt „User-Defined Bridge“. Es ist eine Erweiterung der „historischen“ Bridge Funktion. Während sich bei der ursprünglichen Bridge Funktion alle Container im selben Subnetz befinden und ungehindert miteinander kommunizieren können, schränkt die User-Defined Bridge diese Möglichkeit ein. Dazu legt man auf dem Host verschiedene Bridge-Networks an und gibt ihnen Namen. Bei Start eines Containers ordnet man diesen einem – oder mehreren – dieser Bridged Networks zu. Jetzt macht der Name „Bridge“ plötzlich Sinn, denn Container im selben Bridged Network können miteinander kommunizieren, Container, die in verschiedenen sind, hingegen nicht. (siehe Abbildung 5)

Aber auch wenn Bridged Networks auf unterschiedlichen Hosts dieselben Namen haben, können Container auf diese Weise nicht Host-Übergreifend kommunizieren. Dazu benötigt man Overlay Netzwerke.

Overlay Netzwerke gehören mit zur Grundausstattung von Docker und basieren auf VXLAN. Um eine gesicherte Kommunikation zwischen den Hosts aufbauen zu können, benötigt man neben den Hosts

noch einen Key-Value Store. Dafür können Consul, Etcd und ZooKeeper genutzt werden. Mit einem Docker Tool namens Docker Machine wird dann per VXLAN ein Overlay Netzwerk etabliert, das die Container nutzen können, um miteinander zu kommunizieren. Wie bei den User-Defined Bridges werden diese per Namen identifiziert. Hosts innerhalb eines Overlay Netzwerkes lassen sich mit Orchestrierungstools Docker Swarm administrieren. Dazu später noch mehr.

Out-of-the-Box ist Docker IPv4-only. Jedoch kann die Software grundsätzlich auch IPv6. Das muss jedoch mittels Parameter beim Start extra gefordert werden, sonst werden nur die Link-Lokalen IPv6 Adressen automatisch konfiguriert. Über einen weiteren Parameter wird ein Netzwerkpräfix an den Container übergeben.

Hat man Docker mit globalen IPv6 Adressen gestartet, so wird automatisch das forwarding aktiviert.

Es stehen verschiedene Möglichkeiten für das IPv6 Handling zur Verfügung, die den Rahmen dieses Artikels sprengen würden.

Das charmante an der Nutzung von IPv6 ist, dass man auf jegliches NAT verzichten kann. So wird es möglich auf dem selben Host in unterschiedlichen Container Netzwerk-Anwendungen zu starten, die auf dem selben TCP oder UDP Port horchen. Grundsätzlich ist das auch mit IPv4 möglich, allerdings benötigt man ausreichend Adressen dafür.

Swarm

So richtig spannend wird die Arbeit mit Containern, wenn man sie clustert. Auch

dafür stellt Dockers ein eigenes Tool namens Docker Swarm zur Verfügung. Mittels Docker Swarm werden Hosts zu einem virtuellen System zusammengefasst, das sich nach außen hin wie ein einziges Hosts-System verhält. Ein Swarm stellt dabei dieselben APIs wie ein einzelnes Host-System zur Verfügung. Das ermöglicht es, dass dieselben Tools für das Management und die Entwicklung von Containern genutzt werden können.

Darüber hinaus kann die Cluster-Technologie genutzt werden um Skalierbarkeit und Hochverfügbarkeit zur Verfügung zu stellen. Das erreicht man zum einen dadurch, dass man schlicht jederzeit weitere physikalische Systeme in das Cluster einbringen kann und so mehr Kapazitäten erhält und natürlich auch mehr Hardware, die als Backup für bestehende Systeme dient. Dockers gibt zurzeit an, dass ihre Schwarmtechnologie bis zu 1000 physikalischen Systeme und 50.000 Container ohne Performance-Verluste skaliert.

Aber was nutzen die besten physikalischen Server, wenn es keine ordentliche Verteilstrategien gibt? Von Hause aus unterscheidet man zwei Klassen von Verteilern, Filter genannt: Node Filter und Container Filter. Bei der Generierung eines Containers kann der Entwickler für jede der beiden Klassen Vorgaben machen.

Node Filter

Bei Node Filtern unterscheidet Docker zwei Typen: „constraint“ und „health“.

Der Node Filter constraint fasst Host Systeme zu Gruppen zusammen, die bestimmte Eigenschaften gemeinsam haben, die für den Entwickler oder die Anwendung wich-

Container: Hire and Fire im Rechenzentrum

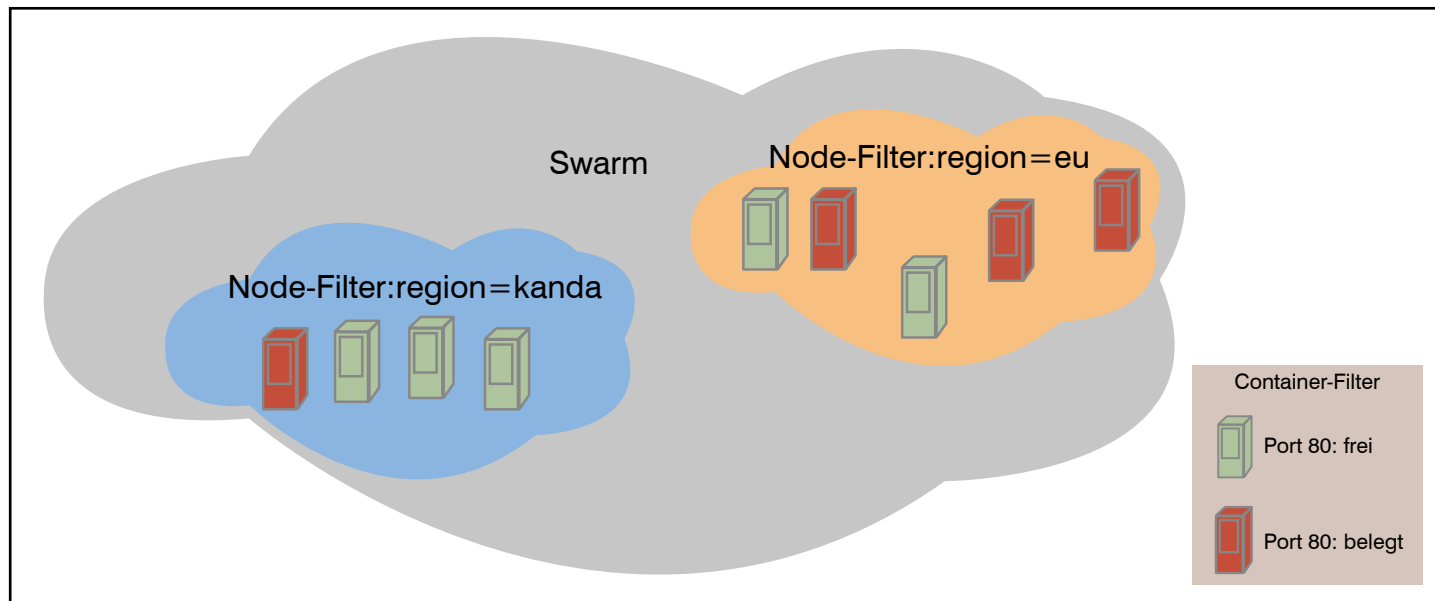


Abbildung 6: Note- und Container-Filter

tig sind. Neben den von Hause aus vorhandenen Default-Gruppen können auch eigene definiert werden. Zu den Default-Gruppen gehören bspw. Storage-driver, Betriebssystem oder Kernelversion. So kann sichergestellt werden, dass ein Container nur auf Systemen läuft, die eine bestimmte Kernelversion haben oder über SSD Platten verfügen.

Neben den Default-Gruppen können auch eigene Gruppen definiert werden. Will man z. B. sicherstellen, dass ein Container innerhalb der EU läuft, können die Nodes, die in Europa aufgestellt sind, den Gruppen-Tag „region=eu“ bekommen, der bei der Container-Erstellung dann als Filter genutzt werden kann. Ein anderes Beispiel wäre es den Gruppen-Tag „environment“ zu nutzen um Entwicklung, Test und Betrieb voneinander zu trennen.

Der Node Filter „health“ gruppiert nach gesunden und ungesunden Hosts. Ungesund ist ein Host, wenn er wahlweise gar nicht erreichbar ist oder seinerseits anderweitige Probleme bei der Kommunikation mit dem Cluster existieren.

Container Filter

Anders als Node Filter, die physikalische Systeme gruppieren, werden Container Filter genutzt, um Container zusammen zu fassen. Dafür stehen drei Varianten zur Verfügung: Affinity, Dependency und Port.

Mit dem Affinity Filter wird, wie der Name schon sagt, eine Affinität zwischen Containern definiert. Dafür gibt es wiederum verschiedene Möglichkeiten:

- **Feste Zuordnung**
Das ist die einfachste und zwingendste Methode: bei der Anlage eines Containers wird der Name oder die ID des Containers angegeben, der zwingend auf demselben Host laufen muss. Welcher physikalische Host das ist, muss man nicht wissen, das Swarm Tool sucht anhand des Namens/der ID den richtigen Host heraus.
- **Vorgabe des Images**
Um unnötige Imageverteilungen im Cluster zu verhindern, gibt man mittels der Image-Affinity an, dass ein Container nur dort gestartet werden darf, wo ein bestimmtes Image bereits vorhanden ist.
- **Label**
Das Label ist frei wählbar, um Container zu gruppieren, die zusammen gehören und auf einem Host laufen sollen.

Anders als der Affinity Filter, der Container nach logischen Kriterien gruppiert, definiert der Dependency Filter Abhängigkeiten. Aktuell werden drei Abhängigkeiten genutzt: geshareter Speicherplatz, Link-Aliase und Netzwerk-Stacks. Will man bspw. einen Container nur dann installieren, wenn dort die Platten von Container A geshared sind, kann man diesen Filter nutzen. Ist das nicht der Fall, wird der Container nicht installiert.

Anders als die Filter Affinity und Dependency, die Container gruppieren, um sie auf demselben Host zu starten, ist der Port Filter das genaue Gegenteil, ein Ausschlussfilter. Mit Port sind TCP/UDP Ports gemeint: bei der Inbetriebnahme eines Containers wird ein Host ge-

sucht, auf dem die geforderten Ports noch frei sind. Da das Default-Verhalten von Docker Containern bzgl. der Netzwerkverbindung mittels NAT und Port-Forwarding realisiert ist, kann man auf demselben physikalischen System keine zwei Container laufen lassen, die nach außen denselben Port benötigen. Bspw. ist es nicht möglich auf demselben Host einen Container mit einem nginx und einen weiteren mit einem Apache zu betreiben, die beide Port 80 benötigen. Darum kann mittels des Port Filters sichergestellt werden, dass ein Container nur auf solch einer Cluster-Node läuft, deren Port 80 noch nicht genutzt wird.

Verteil Strategien

Hat man bspw. ein Rechenzentrum in Kanada und eines in Deutschland und möchte einen Apache Container starten, so könnte der Gesamtfilter lauten: node-filter:region=eu&container-filter:Port=80 (vgl. Abbildung 6). Nun gibt es drei Möglichkeiten:

1. Es gibt kein System, das die Anforderungen erfüllt
2. Es gibt genau ein System, das die Anforderungen erfüllt
3. Es gibt eine Reihe von Systemen, die genutzt werden könnten

In ersten Fall muss man entweder seine Filter aufweichen oder die Voraussetzungen schaffen. Der zweite Fall ist trivial.

Spannend ist der dritte, der – ein ausreichend großes Cluster vorausgesetzt – der

Container: Hire and Fire im Rechenzentrum



Abbildung 7: Prozentuale Verteilung auf Basis von 501 befragten Usern, mehrfach Antworten waren möglich

Normalfall sein sollte. Hier bietet der Docker Swarm drei Verteilstrategien an:

1. Spread
2. BinPack
3. Random

Die ersten beiden berücksichtigen neben der Anzahl installierter Container auch die Auslastung von CPU und RAM, **nicht aber der Netzwerkschnittstelle(!)**. Der Unterschied zwischen den beiden ist, dass die Spread Strategie die Container möglichst breit streut, wohingegen BinPack möglichst viele Container auf einen Host packt. Beide haben offensichtlich ihre Vor- und Nachteile: Spread verteilt die Last optimal auf die physikalischen Ressourcen, läuft jedoch Gefahr, dass nirgends mehr Platz ist, wenn ein „große“ Container gestartet werden soll. Bei BinPack ist es genau umgekehrt.

Random macht, wonach es klingt, es verteilt die Container wahllos auf die Hosts.

Alternativen zu Docker Swarm

Da ein Cluster von Docker Containern sich wie eine einzelne Node verhält und somit die Standard APIs zur Verfügung stehen, können neben dem Docker eigenen Tool Docker Swarm auch andere zum Management und zur Orchestrierung verwendet werden. Wie man an Abbildung 7 erkennt, sind Kubernetes von Google und ECS bei Docker-Usern noch sehr beliebt.

Fazit

Wie eingangs geschrieben, ist die Missi-

on von Dockers nichts Geringeres als mit Containern einen Paradigmenwechsel bei der Entwicklung und dem Betrieb alt hergebrachter Anwendungen zu erreichen. Was ist von dieser „Mission“ erfüllt, wenn man den den Pathos mal außer Acht lässt?

Im Grunde sind Container „kastrierte virtuelle Maschinen“: man nimmt ihnen ihr Herz – den Kernel. Aber macht sich das zu Zombies? Nein: zwar sind sie einerseits weniger flexibel, da der Host das Betriebssystem definiert – wenn auch nicht die Distribution –, andererseits werden die Anwendung selbst flexibler. Vorausgesetzt man kann eine Anwendung in eine Reihe parallelisierbarer Aufgaben zerlegen. Denn dann können diese Aufgaben jeweils in einem eigenen Container entwickelt und als Image für den Betrieb bereitgestellt werden. Wird eine dieser Aufgaben nun stärker beansprucht als eine andere, können für diese eine Aufgabe weitere Container gestartet werden, ohne gleich die gesamte Anwendung mittels virtueller Maschinen zu vervielfältigen. Da Container sich gut automatisieren lassen, kann das bei Bedarf auch schnell geschehen und ebenso schnell wieder gestoppt werden, wenn der Bedarf entfällt. Ein Beispiel wäre der eigene Webaufruf, der unvorhergesehener Weise in den populären Medien genannt wird. Ein Vielfaches der Normalen Userzahl greift in diesem Moment auf den Webaufruf zu, ohne dass jedoch irgendwelche Hintergrundprozesse neue Daten berechnen müssten. In diesem Fall reicht es kurzfristig die Anzahl der Cacheserver zu erhöhen, ohne die Datenbank aufblasen zu müssen oder mehr Batchjobs bereit zu hal-

ten. Ist der Ansturm verebbt, kann man die Cacheserver wieder löschen und die Ressourcen für sinnvolle Dinge nutzen.

Was dieses Beispiel aber auch zeigt, ist, dass Container einen erheblichen Impact auf die Netzwerke haben: startet ein Administrator mit einem Tool wie Docker Swarm oder Googles Kubernetes mal eben 10 bis 20 neue Webserver-Instanzen, so sorgt das Orchestrierungstool dafür, dass diese Webserver über eine Reihe von Host verteilt werden, die gerade wenig RAM- und CPU-Ressourcen benötigen und bei denen der Port 80 nach außen auch noch frei ist. Gleichzeitig werden – ebenfalls automatisch – die notwendigen Overlay-Netze in Betrieb genommen, damit die Cache-Server die Anwendungsserver, Datenbanken, Netzwerklaufwerke etc. erreichen können. Stand heute, beachtet Docker Swarm dabei nicht mal die Netzwerkauslastung der eigenen Schnittstelle. Von den Gefahren von Overlay-Netzwerken, die die Physik missachten, ganz zu schweigen. Hinzu kommen ggf. noch Sicherheitsregeln, die auf Firewalls geschaltet werden müssen.

Container dynamisieren Anwendungen, da sie jederzeit kommen und gehen. Auf der GlueCon 2014 gab Google Senior Staff Software Engineer Joe Beda an, dass sie alles in Containern laufen lassen: mehr als eine Millionen Server, auf denen über 2 Milliarden Container pro Woche gestartet werden – wie viele wegfallen oder nur neugestartet werden, sagte er nicht.

Damit sollte jedem Netzwerker klar sein, dass er sich mit dem Thema Container auseinandersetzen muss, wenn das Un-

Container: Hire and Fire im Rechenzentrum

ternehmen auch nur ansatzweise darüber nachdenkt, diese Technik einführen zu wollen. Lösungen aus Sicht des Netzwerks gibt es verschiedene: eine grundlegende Frage wäre z.B., ob man sich auf die Overlay Technik einlassen möchte oder lieber mit gerouteten Netzen arbeitet, auch eine Form von SDN ist denkbar. Wie das Beispiel von IPv6 von vorhin zeigte, können Container grundsätzlich auch geroutet erreicht werden, was ein Tunneling überflüssig machen würde. Welche Voraussetzungen für welche Lösung vorliegen müssen und welche Vor- und Nachteile diese haben, wird in einem Vortrag auf dem ComConsult Netzwerk Forum erörtert werden und später im Jahr als Video erscheinen.

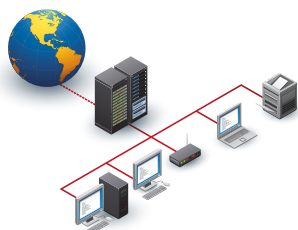
Ohne hier bereits auf die Details einzugehen, sieht der Autor die Container Technologie grundsätzlich als aussichtsreichsten Kandidaten, der auf Tunnel verzichten kann.

Das erscheint auf den ersten Blick merkwürdig, setzt doch die hauseigene Swarm-Technik auf VXLAN und die Default-Praxis über NAT+Portmapping zwingt einen gerade dazu, Host-Systeme in unterschiedlichen IP Netzen mittels Overlays zu verbinden. Das stimmt zwar, jedoch gilt für Container auch, dass zum einen auf das Portmapping und NATting kann verzichtet werden kann und man stattdessen den Containern IP Adressen zuweist, die zumindest im eigenen Netz geroutet werden können. Hinzu kommt, dass die Idee hinter Containern doch gerade daraufsetzt, dass komplexe Anwendungen in kleine „Dienste“ aufgeteilt werden, die über IP kommunizieren. Und noch ein weiterer Punkt kommt hinzu: die Containerphilosophie ist „Hire and Fire“ nicht „Move on“. Sprich: fällt ein Container aus, wird er woanders neu gestartet, nicht verschoben. Eine Layer 2 Verbin-

dung – virtuell oder physikalisch – muss also gar nicht vorhanden sein, eine geroutete reicht völlig aus, vorausgesetzt die Umschaltzeiten im IP Netz reichen, um die Anwendung als ganzes nicht abstürzen zu lassen. Dafür jedoch muss die Anwendung entsprechend entwickelt sein.

Genau darin liegt nun die Hoffnung zumindest für künftige Anwendungen: da Container sich bei den Entwicklern zunehmend großer Beliebtheit erfreuen, darf man davon ausgehen, dass sie zunehmend auf IP als Standardschnittstelle zwischen Prozessen setzen. Wenn wir Netzkern ihnen im Gegenzug geroutete Verbindungen zur Verfügung stellen, die in Sekundenschnelle umschalten können, würden in Zukunft Layer 2 Verbindungen zwischen Systemen überflüssig und damit auch die Tunnel. Ganz ehrlich: ich würde ihnen nicht nachweinen.

Kongress



ComConsult Netzwerk Forum 2016

18.04. - 21.04.16 in Königswinter

Aus der Sicht von ComConsult Research werden die folgenden Fragen die Entwicklung unserer Netzwerke in den nächsten Jahren dominieren:

1. Brauchen wir mehr Intelligenz und Service-Orientierung in unseren Netzwerken als bisher? Wann ja, wo kommt sie her?
2. Müssen wir dynamisch wachsende Rechenzentren mit mehr dynamisch skalierenden Netzwerken begleiten? Wenn ja, wie können Netzwerke wirtschaftlich dynamisch skalieren?
3. Welche Bandbreiten brauchen wir wo, warum und wann? Ist die Zeit von Scale up vorbei und brauchen wir intelligentere Konzepte um ein Optimum aus Preis, Verfügbarkeit und Leistung zu erreichen? Welche Auswirkungen haben die neuen WLAN-Standards im Access Bereich?

Diese Fragen sind ebenso dominant wie komplex. Ihre Beantwortung erfordert die Auseinandersetzung mit Technologien wie Fabrics, SDN, ACI, NSX, IPv6, WLAN, 25, 50, 100, 400 Gbit/s Ethernet. Die Bewertung des Bedarfs erfordert eine genaue Analyse der neuesten Anwendungs-, Speicher- und Server-Architekturen. Und das ganze muss zudem immer sicherer werden. Sicherheit wird zum Schlüsselkriterium, kann aber nicht losgelöst von der Netzwerk-Architektur gesehen werden. So erfordern dynamisch skalierende Netzwerke eine dazu passende dynamisch skalierende Sicherheits-Lösung.

Die Anforderungen an Netzwerk-Infrastrukturen waren noch nie so komplex und gleichzeitig mit großen Fragezeichen der Wirtschaftlichkeit versehen.

Hier setzt das ComConsult Netzwerk Forum 2016 an:

- wir analysieren wo der Bedarf herkommt und wer davon betroffen ist
- wir bewerten die dominanten Lösungstechnologien
- wir geben Empfehlungen zu den anstehenden Investitionen und deren Wirtschaftlichkeit

Moderatoren: Dipl.-Inform. Petra Borowka-Gatzweiler, Dipl.-Math. Cornelius Höchel-Winter, Dr.-Ing. Behrooz Moayeri

Preis: € 2.590,- netto 4 Tage
 € 2.390,- netto 3 Tage
 € 990,- netto 1 Tag



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Standpunkt

Manchmal sind es die einfachen Dinge...

Der Standpunkt von Dr. Joachim Wetzlar greift als regelmäßiger Bestandteil des ComConsult Netzwerk Insiders technologische Argumente auf, die Sie so schnell nicht in den öffentlichen Medien finden und korreliert sie mit allgemeinen Trends.

Der Bagger hatte die Glasfaser durchtrennt. Und natürlich waren die Fasern beider Provider davon betroffen. Der Standort war also von der Welt abgetrennt, obwohl man bei der Planung auf höchste Verfügbarkeit geachtet hatte. Es gab zwei Provider-Übergabepunkte, zwei getrennte Trassen, eine im Norden, die andere im Süden des Standortes und zwei separate Hauseinführungen. Redundanz schützt eben grundsätzlich nicht vor Schaden. Soweit – so schlecht.

Die Provider versprochen, ihre Trassenführung zu überarbeiten. Bis dahin würden einige Monate ins Land ziehen. In der Zwischenzeit sollte eine zusätzliche Anbindung bereitgestellt werden. Eine Richtfunkstrecke hoher Bitrate wurde aufgebaut. Dort hinauf würde kein Bagger kommen.

Das Routing wurde entsprechend konfiguriert. Auch bestimmte Firewall-Regeln und Access Lists waren den Sicherheitsrichtlinien entsprechend einzurichten. Zuletzt wurde ein Test durchgeführt. Man schaltete die (inzwischen reparierten) Glasfaserstrecken ab und prüfte, ob alle externen Ziele nach wie vor erreichbar waren. Vorbildlich!

Leider konnte ein bestimmter externer Geschäftspartner nicht erreicht werden. Und wie so oft, konnte niemand einen Fehler finden. Der Provider sah in seinem Managementsystem alle Komponenten der Richtfunkstrecke „grün“. Das Routing am Standort funktionierte offensichtlich, denn andere Geschäftspartner wurden erreicht. Die Firewall-Regeln und Access Lists waren von verschiedenen Fachleuten geprüft und für korrekt befunden worden.

Schließlich wurde ich gebeten, mich des Problems anzunehmen und entsprechende Messungen mit Protokollanalyatoren durchzuführen. Informationshalber erhielt ich verschiedene E-Mails der genannten Fachleute in Kopie, aus denen ich Informationen über die Topolo-



gie und die IP-Adressierung entnehmen konnte.

Als erstes nahm ich mir die Access Lists vor, die aus der Konfigurationsdatei eines Routers herauskopiert worden waren. Man hatte mehrere IP-Netze eingetragen, derer zwei ich Ihnen hier aufschreibe:

```
permit ip any 192.168.11.0 0.0.0.255
permit ip any 192.168.22.0 0.0.0.128
```

Das IP-Netz des Geschäftspartners, der nicht erreicht werden konnte, findet sich in der zweiten Zeile. Die Adresse des Netzes lautet also 192.168.22.0/25.

Haben Sie das Problem bereits erkannt? Zugegeben, ich habe einen Moment gebraucht, bis ich es plötzlich sah. Betreiben wir also ein wenig IP-Mathematik: Die Subnetzmaske zer-

teilt bekanntlich die IP-Adresse in einen Anteil für das Netzwerk und einen Anteil für das Endgerät. Der Netzwerk-Anteil wird durch eine Folge von 1-Bits gekennzeichnet, der Endgeräte-Anteil mit 0-Bits. Die Subnetzmaske des betroffenen Netzes besteht somit aus 25 1-Bits und sieben 0-Bits, in der gebräuchlichen Dezimalschreibweise ist das 255.255.255.128.

Der hier eingesetzte Router möchte in Access Lists stattdessen das Einerkomplement der Subnetzmaske sehen, also 25 0-Bits und sieben 1-Bits. Das entspricht der Zahlenfolge 0.0.0.127 (ich erspare Ihnen einen Bandwurm aus Nullen und Einsen). Kleine Ursache, große Wirkung: Die Zahlen 128 und 255 sind uns aus der Subnetzmaske so vertraut, dass unser Gehirn sie auch in den Access Lists als korrekt wiedererkennt.

Eigentlich verdient es der Router-Hersteller, dafür gerügt zu werden. Ein Computer kann nämlich das Einerkomplement einer 32-Bit-Zahl schnell und fehlerlos errechnen. Warum um alles in der Welt darf man in den Access Lists nicht die übliche Subnetzmaske schreiben?

Warum erzähle ich Ihnen diese Geschichte? Weil ich (wieder einmal) erlebt habe, dass es Probleme gibt, die eine so einfache Ursache haben, dass niemand darauf kommt. Anders ausgedrückt: Lassen Sie sich nicht von Fachleuten beirren, die behaupten, sie hätten das alles schon mehrfach geprüft. Schauen Sie lieber selbst noch einmal darauf!

Seminar

Trouble Shooting in vernetzten Infrastrukturen 10.05.-13.05.16 in Aachen

Dieses Seminar vermittelt, welche Methoden und Werkzeuge die Basis für eine erfolgreiche Fehlersuche sind. Es zeigt typische Fehler, erklärt deren Erscheinungsformen im laufenden Betrieb und trainiert ihre systematische Diagnose und die zielgerichtete Beseitigung. Dabei wird das für eine erfolgreiche Analyse erforderliche Hintergrundwissen vermittelt und mit praktischen Übungen und Fallbeispielen in einem Trainings-Netzwerk kombiniert.

Referenten: Markus Geller, Markus Schaub, Dr.-Ing. Joachim Wetzlar

Preis: € 2.290,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Aktuelle Sonderveranstaltung

Voice und Video im WAN

**All-IP, B2B, B2C und Daten kämpfen um die Kapazität,
wie kann das beherrscht werden?**

30.05.16 in Köln

Die ComConsult Akademie veranstaltet am 30.05.16 ihre Sonderveranstaltung "Voice und Video im WAN" in Köln.

Wie kann die Übertragung von Sprache und Video im WAN optimiert werden ohne andere Anwendungen zu gefährden? Wie gehen wir mit einem immer größeren Anteil und vor allem der Integration von Video in wesentliche Geschäftsprozesse um? Neue Kollaborations-Lösungen erhöhen zudem den Druck auf die Infrastrukturen. Verkehrslasten werden dabei immer dynamischer und einfache statische Regeln wie traditionelles QoS stoßen an ihre Grenzen.

Die Sonderveranstaltung beleuchtet den Status Quo, die Zukunftsaussichten, mögliche Optionen und eventuellen Investitionsbedarf für den sicheren Betrieb aller Anwendungen mit folgenden Inhalten:

Basistechnologien für Weitverkehrsnetze und Anwendungen im WAN

- Welche WAN-Technologien existieren, welche sind praxisrelevant?
- Welche Multimedia-Anwendungen werden typischerweise im WAN übertragen?
- Welche Anwendungen werden zukünftig verstärkt über das WAN transportiert?
- Welche Anforderungen stellen diese an die Qualität des WAN?
- Wie unterscheiden sich WAN-Technologien in Hinblick auf Multimediakommunikation?
- Wie wichtig sind private Weitverkehrsnetze in Zeiten der Public Cloud?



- Welche effektiven Datenraten stehen im WAN zur Verfügung?
- Welche Mechanismen für den Lastausgleich im WAN existieren?
- Welche Topologie-Varianten existieren?

Voice- und Video-CODECs für das WAN

- Welche Voice- und Video-CODECs spielen in der Praxis eine Rolle?
- Wie funktionieren adaptive CODECs?
- Welche Mechanismen zur Fehlerkorrektur existieren?
- Welche CODECs eignen sich für den Einsatz im WAN?
- Welche CODECs werden von den Herstellern präferiert?

Quality of Service und Call Admission Control im WAN

- Wie einfach ist eine WAN-Leitung auszulasten?
- Welche Auswirkungen hat Überlast im WAN auf Multimediakommunikation?
- Wie wird die Qualität von Multimediakommunikation im WAN sichergestellt? Welche QoS-Mechanismen existieren und wie funktionieren sie?
- Was ist Call Admission Control (CAC) und wie spielt es mit QoS zusammen?
- Welche Mechanismen verbergen sich hinter CAC?
- Wie implementieren die TK-, UC- und Video-Hersteller CAC?
- Welche Topologie-Varianten existieren?

Troubleshooting von Voice und Video im WAN

- Welche Fehlerbilder treten typischerweise bei Voice und Video im WAN auf?
- Wie analysiert man Fehlerbilder in Weitverkehrsnetzen?
- Welche Tools benötigt man für das Troubleshooting?

SIP-Trunking, NGN und das WAN

- Was bedeutet die All-IP-Transformation der Sprachnetze für das WAN?
- Wie bindet man VoIP- und UC-Lösungen heute an Carrier-Netze an?
- Was versteht man unter On-Net- und Off-Net-Routing?
- Was ist bei der multinationalen Sprachübertragung zu beachten?
- Welche Angebote existieren heute am Markt?
- Welche Anforderungen werden zukünftig an ein NGN gestellt?

Fax-Antwort an ComConsult 02408/955-399

Anmeldung

Voice und Video im WAN

Ich buche das Seminar
Voice und Video im WAN

☐ 30.05.16 in Köln
zum Preis von € 1.090,-- netto

☐ Bitte buchen Sie mir ein Hotelzimmer

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Neuer Report, Seminar und kostenloses Video

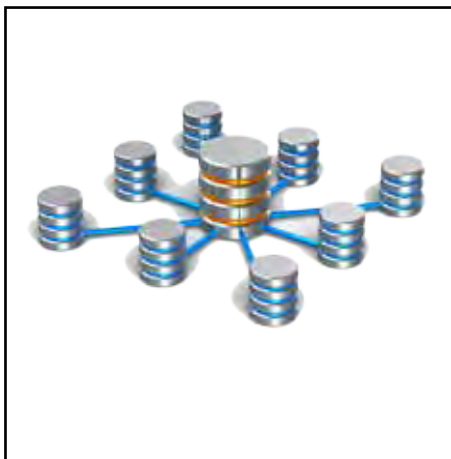
Storage-Lösungen in Unternehmen



Report-Neuerscheinung: Design von Storage-Lösungen in Unternehmen

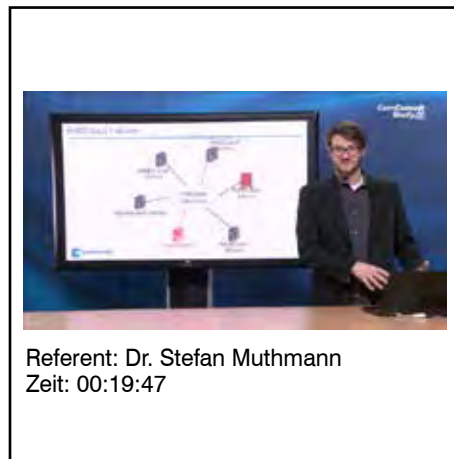
Der Storage-Markt befindet sich durch neue Technologien im Bereich der Virtualisierung, der Hardware und der Speichernetzwerke im Umbruch. Dieser Report beleuchtet die Neuerungen und ordnet sie im Vergleich zu etablierten Technologien ein. Dabei werden alle wichtigen Aspekte einer modernen Speicherumgebung, von sicheren, hochverfügbaren und leistungsfähigen Online-Speichern mit zentraler Verwaltung bis hin zur revisionssicheren Archivierungslösungen betrachtet und die jeweiligen Grundlagen detailliert beschrieben.

Damit liefert der Report eine breite Entscheidungsgrundlage für die Planung einer optimalen, unternehmensspezifischen Speicherlösung.



Planung moderner Speicherlandschaften - 30.05.-31.05.16 in Köln

Hohe Lese-/Schreibraten, niedrige Latenz, revisionssichere Bereitstellung, Skalierbarkeit, Hochverfügbarkeit und nicht zu Letzt niedrige Kosten sind nur einige der teils gegensätzlichen Anforderungen, die an Speicherlandschaften gestellt werden. Moderne Speicherprotokolle und Medien, die Konvergenz von Speicher- und Clientnetzen und die Virtualisierung von Speicher z.B. im virtuellen SAN bieten die Möglichkeit, Speicherlösungen zu entwerfen, die den individuellen Ansprüchen zentraler Rechenzentren oder von Filialen unterschiedlicher Größe und Bedeutung gerecht werden. Im Seminar werden die unterschiedlichen Technologien vorgestellt und, basierend auf einem pragmatischen Best-Practice-Ansatz, Szenarien beschrieben um das persönliche Speicher-Optimum zu erreichen.



Referent: Dr. Stefan Muthmann
Zeit: 00:19:47

Kostenloses Video: Preiswerte I/O-Performance durch die Virtualisierung des DAS

Hochverfügbarer, skalierbarer, hochleistungsfähiger Direct Attached Storage? Liefern virtuelle SANs das Allheilmittel für immer weiter steigende Storage-Budgets und Performance-Ansprüche? Das klassische SAN dient aus guten Gründen seit langem als zentraler Speicherort für Unternehmensdaten. Der Aufbau einer, den Unternehmensansprüchen genügenden, SAN Infrastruktur ist allerdings mit hohen Kosten verbunden. Außerdem bringen schnelle, über NVMe angebundene, SSD Speicher traditionelle SANs an ihre Leistungsgrenze.

Im Video werden virtuelle mit klassischen SANs verglichen und die Relevanz für verschiedenen Einsatzszenarien beleuchtet.

Fax-Antwort an ComConsult 02408/955-399

Anmeldung

Planung moderner Speicherlandschaften

Ich buche das Seminar
Planung moderner Speicherlandschaften

☐ 30.05.-31.05.16 in Köln
zum Preis von € 1.590,- netto

☐ inkl. Report "Design von Storage-Lösungen für Unternehmen"
zum Sonderpreis von nur 338,- €

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Zweitthema

Session Border Controller: Die Perimeter-Komponente für All-IP

Teil 4

Fortsetzung von Seite 1



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.

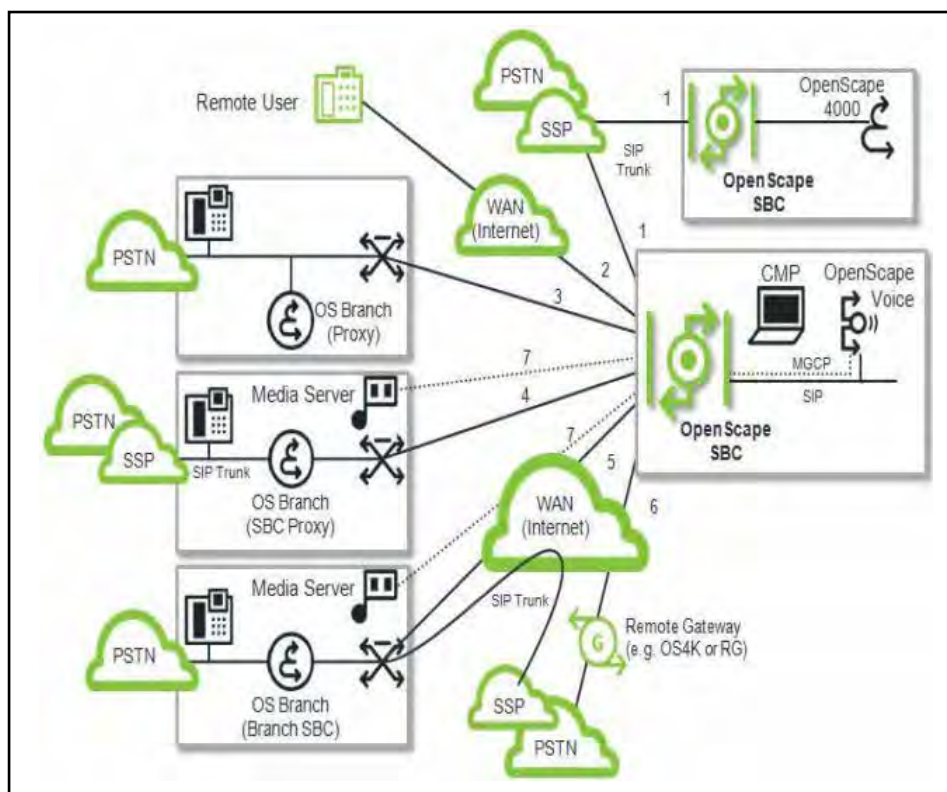


Abbildung 5.17: Einsatz-Szenario des Unify/Atos OpenScope Session Border Controllers

Wie die Tabelle in Abbildung 5.18 zeigt, ist der OpenScope SBC für Unternehmensbedarfe zwischen 1000 und 8000 Sessions ausgelegt. Die virtualisierte Lösung (siehe Abbildung 5.19) ist in der Leistung eingeschränkt auf 250 bis maximal 3500 unverschlüsselte oder maximal 2800 verschlüsselte Sessions. Für die Leistung und Skalierbarkeit werden vielfältige Annahmen vom Hersteller zugrunde gelegt. Sie sind in den Produkt-Informationen dokumentiert und von dort für diesen Beitrag übernommen.

Anmerkungen von Atos/Unify:

1. Die Schaltgeschwindigkeit der Netzwerkschnittstelle ist auf 1 Gigabit Ethernet eingestellt.
2. Bei Keysets wird die Anzeige jeder Keyset-Leitung als ein registrierter Benutzer gezählt.
3. Teilnehmerregistrierungsintervall 3600 Sekunden. Rechnen Sie die folgenden Kapazitätzuschläge (zusätzliche Kapazitätserfordernisse) mit ein, um zu be-

stimmen, wie viele Benutzer max. bei OpenScope SBC registriert sein können (tatsächliche Kapazitätsgrenze), wenn die nachfolgenden Funktionen aktiviert sind:

- a. Digestauthentifizierung: 25 %
- b. Drosselung** (Drosselungsintervall vom SBC 600 s): 60 %
- c. TLS** (Keep-Alive-Intervall 600 s; keine Drosselung): 50 %

* Um kumulative Kapazitätzuschläge zu bestimmen, wenden Sie Zuschlag1 und auf die neue Nummer Zuschlag2 an.

**Kapazitätzuschläge für Drosselung und TLS lassen sich nicht auf gehostete Branch-Fernanwender anwenden.

4. Eine SBC-Sitzung wird als signalisierter SIP-Anruf mit benutzerzugewandter Signalisierungsverbindung und systemzugewandter Signalisierungsverbindung definiert. Ein typischer Sprachanruf zwischen einem lokalen OpenScope Voice-Benutzer und einem Benutzer an einem anderen Standort, der über den SBC registriert ist, oder einem SIP Anschluss, der über den SBC angeschaltet ist, erfordert eine SBC-Sitzung. Ein typischer Videoanruf erfordert zwei SBC-Sitzungen; eine für die Video- und eine zweite für die Audioverbindung. Außerdem sollten im Gegensatz zu einer Audioverbindung bei einer Videoverbindung weitere 20 % der Kapazität von OpenScope SBC eingerechnet werden, da während eines Videogesprächs zusätzliche SIP INFO-Meldungen ausgetauscht werden.

5. Diese Nutzdatenströme werden über den SBC weitergeleitet, wenn eine direkte Medienverbindung zwischen den Endpunkten nicht möglich ist; das ist z.B. der Fall, wenn der SBC für die Medienpakete NAT durchführen muss, weil sie sich in verschiedenen Teilnetzen

Session Border Controller: Die Perimeter-Komponente für All-IP - Teil 4

	IBM x3250 M3 ¹	IBM x3250 M5 ¹	IBM x3550 M3/ M4 oder Fujitsu RX200 S6/S71
Max. registrierte gehostete OpenScape Branch-Fernanwender ² (ohne Digestauthentifizierung oder TLS; Drosselung entfällt)	6,000 ³	6,000 ³	50,000 ³
Max. registrierte SIP-Fernanwender ² , z.B. Telemitarbeiter (ohne Digestauthentifizierung, Drosselung, oder TLS)	6,000 ³	6,000 ³	32,000 ³
Max. Anzahl gleichzeitiger SIP-Signaisierungsanrufe/SBC-Sitzungen ⁴	1,600	2,700	8.000
Max. Anzahl gleichzeitiger RTP-Nutzdatenströme, die über den OpenScape SBC weitergeleitet werden (ohne Medien-Transcoding) ⁵	1,600	2,700	8.000
Max. Anzahl gleichzeitiger sicherer SRTP-Medienströme (MIKEY0 oder SDES), die vom SBC terminiert/vermittelt werden (ohne Medien-Transcoding)	1,280	2,160	6.400
Max. Anzahl von Medien/Standort-Bereichsgruppen	1,024	1,024	1,024
Max. Anzahl eindeutiger Fernanwenderprofile (d.h. Notruf-Standortinformationen, Media Anchoring und Sicherheit etc.)	255	255	255
Anzahl gleichzeitiger SIP Service Provider (SSP)	10 ⁷	10 ⁷	10 ⁷
Anzahl der Belegungsversuche je Stunde (Full-Calls ⁸)	27.000	27.000	79,200
Max. Anzahl Half-calls ⁹ pro Sekunde (ohne Digestauthentifizierung, Drosselung, oder TLS)	15 ⁹	15 ⁹	44 ⁹
Anforderungen zur Registrierungsaktualisierung pro Sekunde (zufällige Registrierung in stabilem Betriebszustand)	5	5	26
Rate der aufgebauten Gespräche im stabilen Betriebszustand	99,99%	99,99%	99,99%
Zeit bis zur Wiederherstellung des stabilen Betriebszustands (99,99% Gesprächsaufbau) nach einem simultanen Neustart aller Endgeräte ¹⁰	<15min.	<15min.	<15min.

Abbildung 5.19: Skalierbarkeit des Unify/Atos OpenScape SBC als virtualisierte Lösung

Quelle: Unify

befinden. Jeder half-call umfasst zwei Nutzdatenströme, die jeweils in die entgegengesetzte Richtung fließen. Beispielsweise werden zwei „Half Calls“ verwendet, wenn ein Benutzer an einem anderen Standort, der über den SBC registriert ist, mit einem anderen Benutzer an einem anderen Standort, der über den SBC registriert ist, verbunden wird, oder mit einem über den SBC verbundenen SIP-Anschluss. Ein einzelner „Half-Call“ wird verwendet, wenn ein lokaler Teilnehmer, der direkt beim OpenScape Voice Server registriert ist, mit einem Benutzer an einem anderen Standort, der über den SBC registriert ist, oder mit einem SIP-Anschluss, der über den SBC angeschaltet ist, verbunden wird.

6. Von den 10 nicht zum VLAN gehörenden IP-Adressen, die an der WAN-Schnittstelle konfigurierbar sind, können zwei das Protokoll UDP unterstützen; die anderen müssen die Protokolle TCP oder TLS unterstützen.

7. Es werden bis zu 10 gleichzeitige SSP SIP-Trunk-Schnittstellen unterstützt. Diese Schnittstellen können eine Verbindung zum selben oder unterschiedlichen SSP herstellen, vorausgesetzt die IP-Adressen auf SSP-Seite sind unterschiedlich. Die SSP-Verbindung kann

auf dieselbe oder unterschiedliche IP-Adressen auf dem OpenScape SBC verweisen.

8. Bei einem „Half-call“ handelt es sich um einen Anruf von der benutzerzugewandten Seite (WAN) an die systemzugewandte Seite (LAN) oder von der systemzugewandten Seite (LAN) an die benutzerzugewandte Seite (WAN). Ein „Full-call“ besteht aus zwei halben Anrufabschnitten. Das bedeutet, dass ein Anruf von der benutzerzugewandten Seite (WAN) ausgelöst wurde, an die systemzugewandte Seite (LAN) geleitet wird und dann zur benutzerzugewandten Seite (WAN) zurückgeleitet wird.

9. Rechnen Sie die folgenden Kapazitätzuschläge mit ein, um zu bestimmen, wie viele Anrufe pro Sekunde max. möglich sind, wenn folgende Funktionen aktiviert sind:

- a. Digestauthentifizierung: 30 %
- b. Drosselung** (Drosselungsintervall 600 s): 40 %
- c. TLS** (Keep-Alive-Intervall 600 s; keine Drosselung): 50 %

* Um kumulative Kapazitätzuschläge zu bestimmen, wenden Sie Zuschlag1 und auf die neue Nummer Zuschlag2 an.

**Kapazitätzuschläge für Drosselung und TLS lassen sich nicht auf gehos-

tete Branch-Fernanwender anwenden.

10. Beim Neustart müssen SIP-Endgeräte die in RFC3261 und OSCAR Kapitel 11, „Best Practices“, aufgeführten Verfahren einhalten. Bei einem gleichzeitigen Neustart aller Endgeräte muss ein erfolgreich registrierter Benutzer unmittelbar in der Lage sein, Anrufe durchzuführen und zu empfangen, und dies mit einer Erfolgsrate von mindestens 99,99% beim Gesprächsaufbau.

Anmerkungen von Atos/Unify:

- 1. s. Hardware Lösung.
- 2. s. Hardware Lösung.
- 3. s. Hardware Lösung.
- 4. s. Hardware Lösung.

5. Jeder RTP-Nutzdatenstrom (Full-call), der durch den zentralen OpenScape SBC weitergeleitet wird, besteht aus zwei Half-calls, die in entgegengesetzter Richtung verlaufen. Beispielsweise werden zwei „Half Calls“ verwendet, wenn ein Benutzer an einem anderen Standort, der über den SBC registriert ist, mit einem anderen Benutzer an einem anderen Standort, der über den SBC registriert ist, verbunden wird, oder mit einem über den SBC verbundenen SIP-Anschluss. Ein einzelner „Half-Call“ wird

Session Border Controller: Die Perimeter-Komponente für All-IP - Teil 4

verwendet, wenn ein lokaler Teilnehmer, der direkt beim OpenScape Voice Server registriert ist, mit einem Benutzer an einem anderen Standort, der über den SBC registriert ist, oder mit einem SIP-Anschluss, der über den SBC angeschaltet ist, verbunden wird.

6. Die RTP-Paketleistung (z.B. Packet Loss) wird durch verschiedene Faktoren beeinflusst:

- Hardware BIOS-Einstellungen, die sich auf Leistung & Energieeinsparung beziehen,
- Hardware BIOS Hyper-Threading,
- Gasteinstellungen einer virtuellen Maschine (VM) für Hyper-Threaded Core Sharing,
- Arbeitsspeicher, der der virtuellen Maschine auf dem Computer-System, auf dem sie läuft (also bei dem Sie „Gast“ sind), bereitgestellt wird,
- Größe des Ringpuffers der Netzwerkkarte des Computersystems, der für den Empfang von Daten für die virtuelle Maschine, innerhalb der z.B. Windows als Gastbetriebssystem installiert ist, bereitgestellt wird,

7. RTP-Paketisierungszeit/-größe. Schalten Sie, zum Erhalten einer besseren Leistung, in den BIOS-Einstellungen Ihres Systems das Hyper-Threaded Core Sharing aus. Mehrere aktive virtuelle Maschinen und kleinere vRAM-Zuordnungen können den RTP Packet Loss verringern.

8. Es werden bis zu 10 gleichzeitige SSP SIP-Trunk-Schnittstellen unterstützt. Diese Schnittstellen können eine Verbindung zum selben oder unterschiedlichen SSP herstellen, vorausgesetzt die IP-Adressen auf SSP-Seite sind unterschiedlich. Die SSP-Verbindung kann auf dieselbe oder unterschiedliche IP-Adressen auf dem OpenScape SBC verweisen.

- s. Hardware Lösung, Anm. 8.
- s. Hardware Lösung, Anm. 9.
- s. Hardware Lösung, Anm. 10.

Ebenso detailliert ist die Feature-Übersicht des OpenScape SBC vom Hersteller angegeben (wir veröffentlichen einen relevanten Auszug)

- Managebar mit CMP Integrated Management, (separate) Managementschnittstelle für OpenScape SBC (seit V8)
- Alarme
- Ethernet Bonding auf LAN/WAN-Schnittstellen
- Separate Ethernet-Schnittstelle für Management und Verwaltung (seit V8)

- Separate IP-Adresse für Signalisierung und Medien
- Einarmige LAN/WAN-Schnittstelle innerhalb der DMZ (seit V8)
- Unterstützung mehrerer WAN-Schnittstellen und Netzwerke
- Überwachung und Auswertung der Dienstgüte (QoS) (seit V8)
- RADIUS-Support (seit V8)
- SBC-Serverredundanz im selben Teilnetz (SBC-Cluster)
- Unterstützung der SBC-Redundanz für geografisch verteilte L3 OpenScape Voice-Knoten
- DNS-Unterstützung
- NTP-Unterstützung
- Call Admission Control
- Drosselung der Datenrate
- Traffic Control per QoS
- Paket-Verfolgung (Tracing)

- Notrufunterstützung
- Media Anchoring
- Media-Pass-Through
- Media-Transcoding (seit V8: G.711, G.729, G.722, G.722.1, iLBC, iSAC)
- Unterstützung von Skype Connect
- SRTP Termination und Mediation
- Konformität mit und Zertifizierung für SIPconnect v1.1 (seit V8)
- TLS / SRTP
- TLS-Unterstützung für SIP Service Provider
- VLAN-Unterstützung für die Verbindung zu Zweigstellen
- Sprach- und Videounterstützung
- Videotelefonie – Dual-Video Content Sharing (seit V8)
- VPN-Unterstützung
- Reduzierung von Dienstblockaden (DoS-Mitigation)
- SIP Firewall-Funktionalität

Hersteller	Unify OS SBC			
Kriterium	Max	Gew.	Roh-P.	gew. P.
Architektur	158	0,62	115	72
Authentisierung, Zugangskontrolle	21	1,13	15	17
Topology Hiding, Angriffsschutz	133	1,01	65	65
SIP Trunking, Leistungsmerkmale	82	1,68	59	99
GESAMTSUMME	394		254	253
Prozent	100%		64,5%	64,3%

Hersteller	Unify OS SBC
Kriterium	
Architektur	72,8%
Authentisierung, Zugangskontrolle	71,4%
Topology Hiding, Angriffsschutz	48,9%
SIP Trunking, Leistungsmerkmale	72,0%

Abbildung 5.20: Evaluierungs-Tabellen des Atos/Unify OpenScape SBC

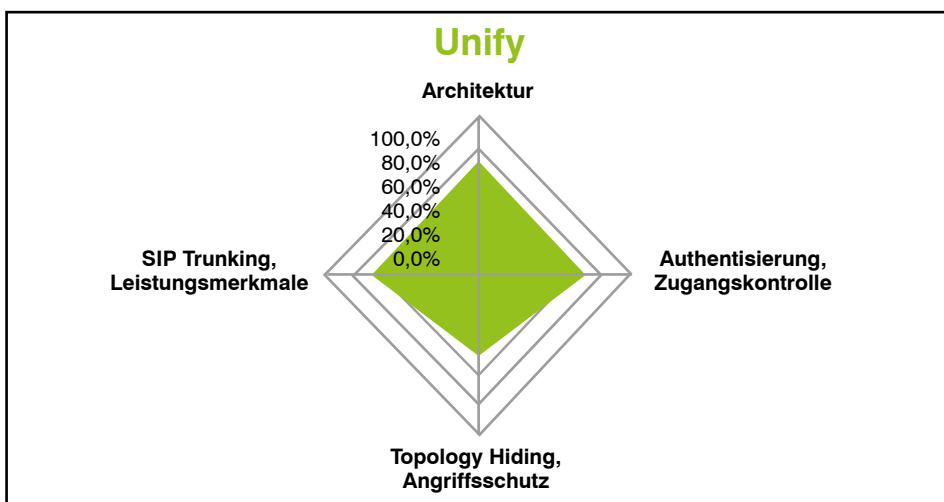


Abbildung 5.21: Evaluierungs-Grafik des Atos/Unify OpenScape SBC

Session Border Controller: Die Perimeter-Komponente für All-IP - Teil 4

- Unterstützung für dynamische NAT-Geräte in Zweigstellen
- Unterstützung der Far-End NAT Firewall
- Unterstützung der Near-End NAT Firewall
- Hosted-Unterstützung für mehrere OSB hinter der gleichen Remote-NAT
- Protokollierung
- Unterstützung von mehreren Geräten hinter einem remote NAT-Router
- Sichere remote Media Server-Kommunikation für Hosted OpenScape Branch
- Überprüfung und Manipulation von SIP-Nachrichten
- Reduzierung von Spam over Internet Telephony (SPIT)
- Synchronisation von Teilnehmer- und Digestauthentifizierung (seit V8)
- Unterstützung von Teilnehmerleistungsmerkmalen

Die Evaluierung des OpenScape Session Border Controllers zeigt einen Gesamt-Erfüllungsgrad von 64,3 Prozent und liegt damit etwa gleichauf mit dem Cisco CUBE. Allerdings verteilen sich die Erfüllungsgrade ganz anders als bei letzterem: Die Bereiche Architektur, Authentisierung / Zugangskontrolle und SIP Trunking / Leistungsmerkmale haben alle einen Erfüllungsgrad von über 70 Prozent, der Bereich Topology Hiding / Angriffsschutz fällt jedoch mit unter knapp 49 Prozent stark ab (siehe Abbildung 5.20 und Abbildung 5.21). In der grafischen Übersicht ist dies durch die deutliche Verkürzung der Raute im unteren Bereich erkennbar.

5.5 Produktbeispiel: Session Border Controller von Mitel

Der Session Border Controller von Mitel heißt passend zum TK/UC Branding MiVoice Border Gateway (früher Mitel Border Gateway). Er arbeitet nach dem typischen Application Proxy Prinzip (Terminieren und Neuaufsetzen der Sessions). Als Application Proxy unterstützt er auch die Mitel UC Applikationen (MiCollab für Audio-, Web- und Video-Konferenzen, MiVoice Business für Telefonie, Mobilnutzer-Integration und Unified Messaging und Contact Center).

Das MiVoice Border Gateway läuft unter Linux und ist sowohl als Hardware als auch virtualisiert unter VMware und Hyper-V betreibbar. Im letzteren Fall unterstützt der SBC virtuelle 1 Gbit Schnittstellen zum Netzwerk. Bei Nutzung mehrerer Netzwerk-Interfaces wird Bonding (Link Aggregation) unterstützt.

Die Produktivnetz-Anbindung zum Netzwerk kann redundant ausgelegt werden, eine separierte Management Schnittstelle ist nicht vorgesehen. Die automatische Fehlerumschaltung auf einen alternativen SIP Trunk bei einer Ausfall-Situation des

primären SIP Trunks regelt bei der Mittel Lösung nicht der SBC, sondern der TK-Server.

Bis zu 5 Session Border Controller (zuzüglich 1 Redundanz-SBC) können im Verbund betrieben werden. Insgesamt können sich maximal 25.000 Telefone registrieren, je SBC werden maximal 2500 Sessions unterstützt.

Mitel gliedert die Funktionalität insbesondere in die nachfolgend beschriebenen sechs Bereiche, die als Funktionsmodule auf der SBC-Plattform integriert sind.

Der Teleworker Dienst unterstützt den Telefonie-Zugriff von remote Nutzern und Contact Center Agenten, insbesondere auch über MiVoice Softclients oder Mitel IP Telefone, die im SOHO installiert sind. Der Nutzer registriert sich als Nebenstelle am Mitel

TK-Server und arbeitet außerhalb des Unternehmens wie mit seinem internen Telefon / Softclient. Als Softclients werden das MiCollab Softphone, das MiContact Center Softphone und der Couterpath Bria SIP Softclient unterstützt.

Der SIP Trunk Proxy-Dienst nimmt die eigentliche SIP Firewall Funktion wahr. Dieses Modul leistet Authentisierung und 128-Bit Verschlüsselung für Signalisierung und Mediatstreams.

Der Application Web Proxy Dienst sichert den remote Zugriff für UC-Anwendungen von Mitel für alle Sessions, bei denen die Kommunikation zwischen dem unternehmens-internen LAN und dem öffentlichen Internet verläuft. Im Grunde ist dies die Erweiterung des Teleworker Dienstes auf UC als Web-Anwendung.

Hersteller			Mittel	
Kriterium	Max	Gew.	Roh-P.	gew. P.
Architektur	158	0,62	94	59
Authentisierung, Zugangskontrolle	21	1,13	14	16
Topology Hiding, Angriffsschutz	133	1,01	75	76
SIP Trunking, Leistungsmerkmale	82	1,68	50	84
GESAMTSUMME	394		233	234
Prozent	100%		59,1%	59,4%

Hersteller	Mittel
Kriterium	
Architektur	59,5%
Authentisierung, Zugangskontrolle	66,7%
Topology Hiding, Angriffsschutz	56,4%
SIP Trunking, Leistungsmerkmale	61,0%

Abbildung 5.22: Evaluierungs-Tabellen des Mitel MiVoice SBC

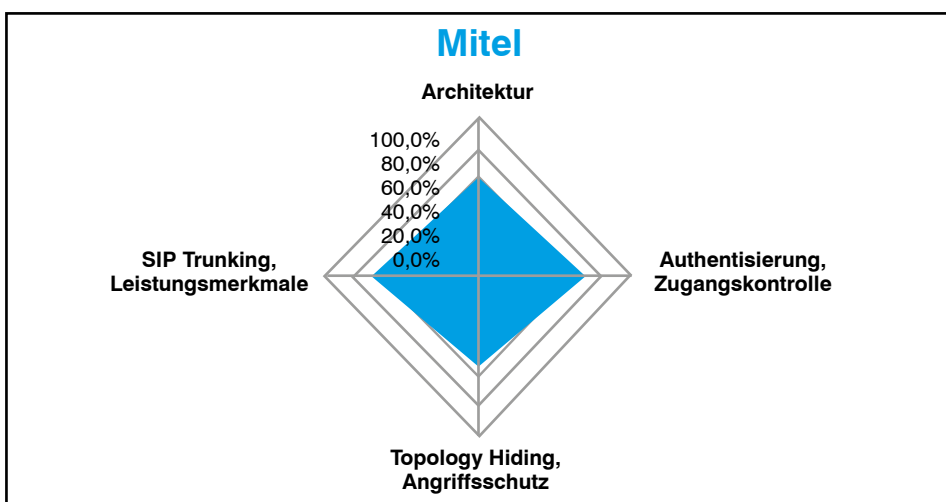


Abbildung 5.23: Evaluierungs-Grafik des Mitel MiVoice SBC

Session Border Controller: Die Perimeter-Komponente für All-IP - Teil 4

Der Remote Management Dienst ermöglicht den administrativen Zugriff von außerhalb des Unternehmens über ein Web-Interface, was insbesondere auch für externe Wartungsnehmer erforderlich ist.

Der Aufzeichnungs-Dienst ermöglicht die Anschaltung von 3rd Party Aufzeichnungs-Lösungen und unterstützt die gesicherte Aufzeichnung von IP Telefonen und Teleworker Telefonen aus. Er leistet auch die gesetzlich vorgeschriebene Abhörfunktion (lawful interception).

Das IPv6 Gateway adaptiert die Kommunikation zwischen IPv4 und IPv6 Geräten.

Zu den Sicherheitsfunktionen gehört das erwartete Minimum im Verschlüsselungsbereich:

- Verschlüsselung und Authentisierung der Signalisierung mit TLS
- Verschlüsselung und Authentisierung der Media Streams für Sprache und Video mit SRTP
- AES 128 Bit Verschlüsselung

Die Evaluierung des Mittel MiVoice Session Border Controllers zeigt einen Gesamt-Erfüllungsgrad von 59,4 Prozent und liegt damit im Rahmen unserer Evaluierung im unteren Bereich. Die Bereiche Authentisierung/ Zugangskontrolle und SIP Trunking / Leistungsmerkmale haben immerhin einen Erfüllungsgrad von über 60 Prozent, Architektur Hiding / Angriffsschutz fällt jedoch mit unter knapp 56 Prozent ein Stück weit ab (siehe Abbildung 5.22 und Abbildung 5.23). In der grafischen Übersicht ist die schwächere Gesamtfunktionalität durch die Verkürzung der Raute in allen Bereichen erkennbar.

5.6 Produktbeispiel: Session Border Controller von Innovaphone

Der Session Border Controller von Innovaphone ist erst seit kurzem auf dem Markt und wird noch einige Weiterentwicklung erfahren. Der SBC kann als Appliance und als virtualisierte Lösung zum Einsatz kommen (evaluiert wurde die Appliance). Als Appliance ist er auf der TK-Plattform IP0010 implementiert und besteht aus den drei Komponenten SBC, Reverse Proxy und Edge (ähnlich wie die Alcatel-Lucent Enterprise Lösung). Nach dem Innovaphone Design-Vorschlag ist er in der DMZ positioniert, kann also durch die üblicherweise vorhandenen IT Firewalls mit geschützt werden.

Der IP0010 hat als SBC die nachfolgenden Funktionsmodule:

- Sicherheit
- Reverse Proxy

- TURN Server
- NAT Erkennung

Eine Übersicht der Gesamt-Architektur zeigt Abbildung 5.24. Der Fokus des Innovaphone SBC liegt eindeutig bei der Zugangskontrolle: Auf dem Reverse Proxy lassen sich Zugriffs-Beschränkungen konfigurieren, durch die Brute Force Angriffe, SIP Dialer oder unberechtigte Zugriffe auf die Administrations-Schnittstelle erkannt und blockiert werden. Da der SBC in der DMZ positioniert ist, wurde von Innovaphone Wert auf einen geringen Administrationsumfang beim gelegt.

Die implementierten Sicherheits-Funktionen schützen gegen unberechtigten Zugriff, DoS Angriffe und Übernahme der Kontrolle durch einen Angreifer.

Der Reverse Proxy unterstützt die Protokolle:

- H.323
- SIP
- HTTP, HTTPS
- LDAP, LDAPS

Da der Session Border Controller ausschließlich die genannten Protokolle durchlässt, können Angriffe auch nur über diese Protokolle geführt werden, was die Schützbarkeit des SBC selbst erhöht.

Die Innovaphone Session Border Controller Evaluierung zeigt einen Gesamt-Erfüllungsgrad von 65,2 Prozent. Sie hat damit einen deutlichen Abstand zu Alcatel-Lucent Enterprises und Avaya und liegt in der Nähe des Cisco CUBE. Allerdings verteilen sich die Erfüllungsgrade genau entgegengesetzt zu letzterem: Einen vergleichsweise hohen Erfüllungsgrad haben

die Bereiche SIP Trunking / Leistungsmerkmale mit knapp 77 Prozent und Authentisierung / Zugangskontrolle mit gut 85 Prozent, einen vergleichsweise niedrigen Erfüllungsgrad haben die Bereiche Architektur mit 57 Prozent und Topology Hiding / Angriffsschutz mit gut 55 Prozent. In der grafischen Übersicht werden die Defizite in diesen Bereichen durch die asymmetrische Form als "quer gestauchte Raute" deutlich (siehe Abbildung 5.25 und Abbildung 5.26).

Vergleich verschiedener Enterprise Session Border Controller

Abschließend ist eine Vergleichsgrafik der evaluierten Produkte mit farblich unterschiedlichen, übereinandergelegten Erfüllungsgraden zu den vier betrachteten Funktionsbereichen in Abbildung 5.27 gezeigt.

Abkürzungen

AES	Advanced Encryption Standard
BIOS	Basic Input Output System
CMP	Common Management Platform (Unify)
CUBE	Cisco Unified Border Element
HTTP	HyperText Transfer Protocol
HTTPS	Secure HyperText Transfer Protocol
IP	Internet Protocol
LAN	Local Area Network
LDAP	Lightweight Directory Access Protocol
LDAPS	Secure Lightweight Directory Access Protocol
NAT	Network Address Translation
OSB	OpenScape Branch (Unify)
OSCAR	Open System for Communication in Realtime

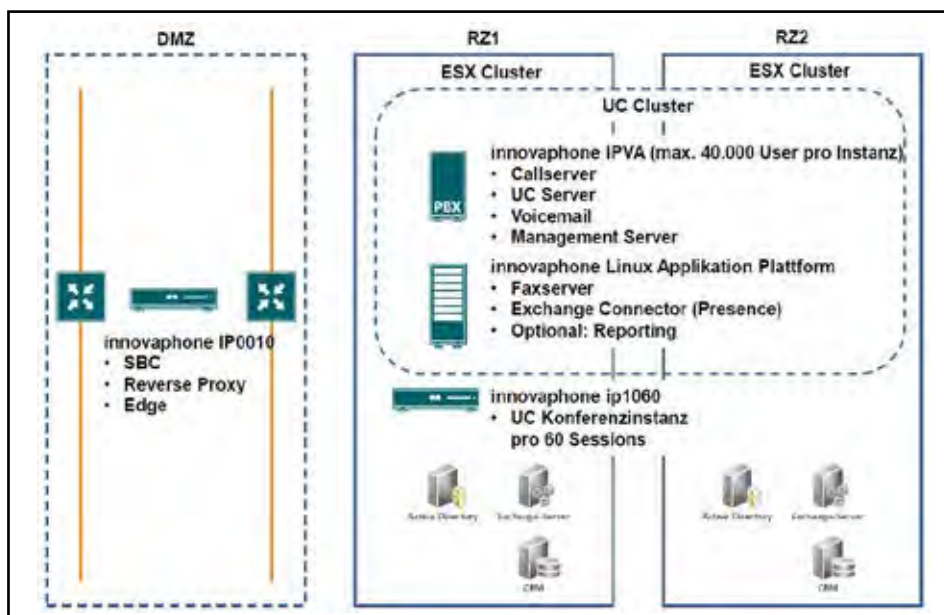


Abbildung 5.24: Architektur-Übersicht der Innovaphone SBC-Lösung

Quelle: innovaphone

Session Border Controller: Die Perimeter-Komponente für All-IP - Teil 4

Hersteller	Innovaphone			
Kriterium	Max	Gew.	Roh-P.	gew. P.
Architektur	158	0,62	90	56
Authentisierung, Zugangskontrolle	21	1,13	18	20
Topology Hiding, Angriffsschutz	133	1,01	74	75
SIP Trunking, Leistungsmerkmale	82	1,68	63	106
GESAMTSUMME	394		245	257
Prozent	100%		62,2%	65,2%

Hersteller	Innovaphone
Kriterium	
Architektur	57,0%
Authentisierung, Zugangskontrolle	85,7%
Topology Hiding, Angriffsschutz	55,6%
SIP Trunking, Leistungsmerkmale	76,8%

Abbildung 5.25: Evaluierungs-Tabellen des Innovaphone IP0010 SBC

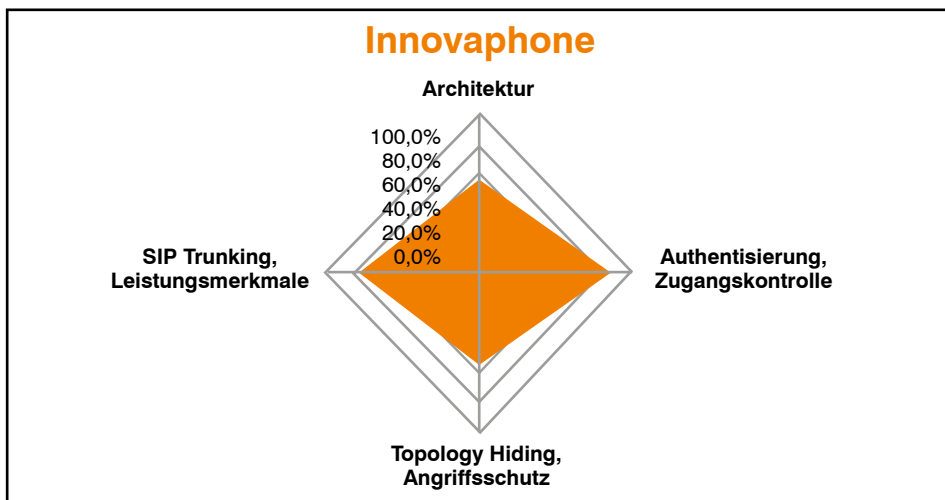


Abbildung 5.26: Evaluierungs-Grafik des Innovaphone IP0010 SBC

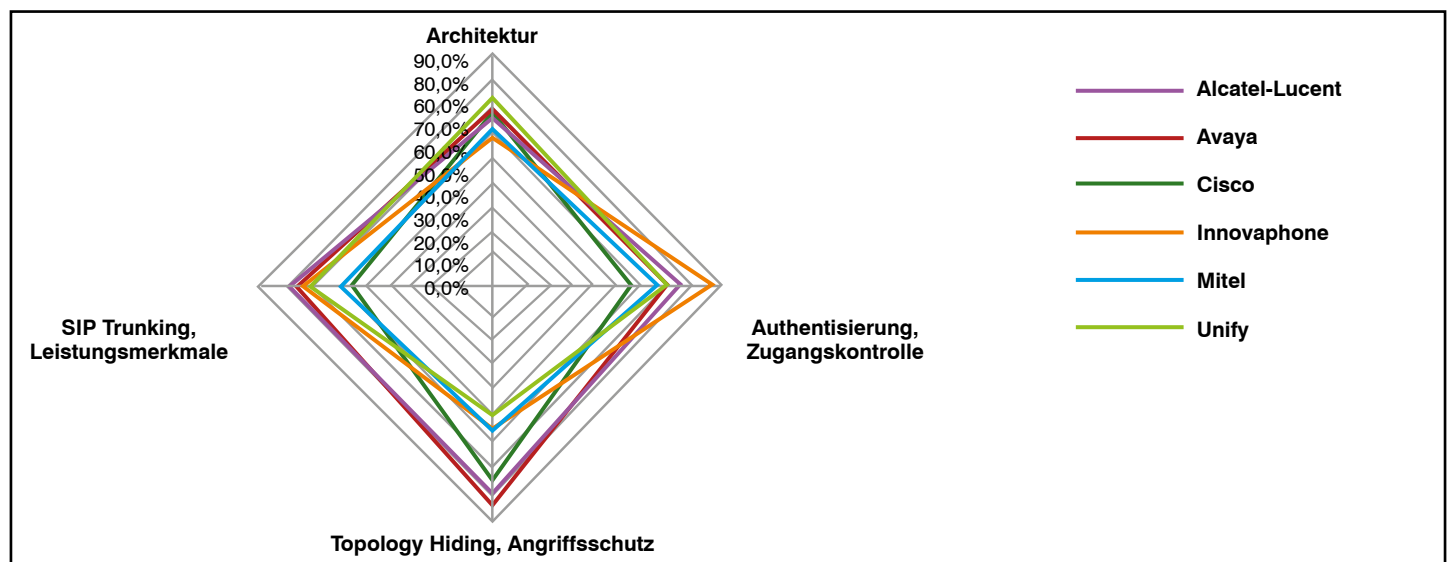


Abbildung 5.27: Evaluierungs-Grafik zum Vergleich der Enterprise SBCs verschiedener Hersteller

RTP	Real-time Transport Protocol
RZ	Rechenzentrum
SBC	Session Border Controller
SIP	Session Initiation Protocol
SOHO	Small Office, Home Office
SRTP	Secure RTP
SSP	SIP Service Provider
SPIT	Spam over Internet Telephony
TCP	Transmission Control Protocol
TK	Telekommunikation
TLS	Transport Layer Security
TURN	Traversal Using Relay NAT
UC	Unified Communications
UM	Unified Messaging
VLAN	Virtual LAN
VM	Virtuelle Maschine
vRAM	virtual Random Access Memory
VPN	Virtual Private Network
WAN	Wide Area Network

Links

www.ietf.org
www.sipforum.org
www.sonus.net

Literatur

- Pat Hurley: Session Border Controller for Dummies; Wiley & Sons, 2nd Edition 2013
- Session Border Controllers: A Primer; Oracle White Paper 2013
- Market Guide for Enterprise SBC; Gartner, Juni 2014
- John Hardwick: Session Border Controllers, Enabling The VoIP Revolution; Data Connection Whitepaper, 2005

ComConsult Veranstaltungskalender

Crashkurs IT-Recht für Nichtjuristen, 11.04.16 in Stuttgart**Garantietermin**

Diese Veranstaltung wendet sich an IT-Leiter, Compliance-Beauftragte und Geschäftsführer, die sich kompakte und praktische Grundkenntnisse zu den rechtlichen Eckpunkten des IT-Projektes verschaffen wollen. Die Inhalte sind insbesondere an Nichtjuristen gerichtet, die sich nicht alltäglich mit rechtlichen Fragestellungen befassen und eine Grundorientierung suchen. In dem Seminar werden auch Praxisfälle erörtert.

Preis: € 1.090,-- netto

Rechenzentrumsdesign - Technologien neuester Stand, 11.04.-13.04.16 in Köln**Garantietermin**

Dieses Seminar analysiert die neuesten Technologie-Trends im Rechenzentrum. Sie lernen von der Verkabelung über die Stromversorgung, die Klimatisierung und den Schrankaufbau, wie ein ausfallsicheres und energieeffizientes Rechenzentrum heute strukturiert wird. Mechanismen für Redundanz im Netzwerk, Lastverteilung und Standort-übergreifende Hochverfügbarkeit werden diskutiert und es wird untersucht wie diese mit dem fortwährenden Trend zur Virtualisierung zusammenspielen. Abschließend werden aktuelle Speichersysteme, deren Anbindung über die am Markt verfügbaren Übertragungsprotokolle sowie Aspekte zur Datensicherung und Disaster Recovery diskutiert.

Preis: € 1.890,-- netto

SIP (Session Initiation Protocol) - Basis-Technologie der IP-Telefonie, 11.04.-13.04.16 in Stuttgart**Garantietermin**

Ziel der Schulung ist die Erläuterung von SIP als den Schlüssel für eine offene, leistungsfähige und Kosten-optimale Kommunikations-Lösung. Es umfasst nahezu alle Dienste, die wir heute für UC benötigen: Sprache, Video, Daten und Präsenz. Lernen Sie was SIP leistet, worin sich wesentliche Hersteller-Lösungen unterscheiden und wie Sie das Beste aus beiden Welten zukunftsorientiert nach Ihrem Bedarf optimieren.

Preis: € 1.890,-- netto

Wireless LAN professionell, 11.04.-13.04.16 in Köln**Garantietermin**

Dieses Seminar vermittelt den aktuellen Stand der WLAN-Technik und zeigt die in der Praxis verwendeten Methoden für Aufbau, LAN-Integration, Betrieb und Optimierung von WLANs im Enterprise-Bereich auf. Die verschiedenen WLAN-Varianten werden analysiert, die Markt- und Produktsituation bewertet, und Empfehlungen für eine optimale Auswahl gegeben.

Preis: € 1.890,-- netto

Recht und Datenschutz bei Einführung von Voice over IP, 25.04.-26.04.16 in Bonn**Garantietermin**

Ziel der Schulung ist es, den Teilnehmern einen Überblick über die aktuelle Situation im Bereich des Datenschutzes im Kommunikationsumfeld zu verschaffen. Datenschutz und Datensicherheit werden zunehmend wichtiger im Umgang mit Kunden und Mitarbeitern. Gerade mit der Einführung von IP basierten Lösungen in den Bereichen Telefonie oder Contact Center, stellen sich neue Herausforderungen in Bezug auf personenbezogene Informationen. Um Ihnen einen Überblick über den rechtlichen Rahmen zu geben beschäftigt sich dieses Seminar u.a. mit Fragen zur Abhörsicherheit, Vorratsdatenspeicherung, Datenverlust und den dazugehörigen Aspekten.

Preis: € 1.590,-- netto

IPv6 Grundlagen - SeminarPlus, 25.04.-26.04.16 in Düsseldorf**Garantietermin**

IPv6 betreiben, bedingt IPv6 verstehen. In diesem Seminar werden die Grundlagen des neuen IP Protokolles verständlich und praxisnah vermittelt. Die Schulung richtet sich gleichermaßen an Planer, Betreiber, Administratoren und Software-Entwickler.

Preis: € 1.790,-- netto

IP-Wissen für TK-Mitarbeiter, 25.04.-26.04.16 in Düsseldorf**Garantietermin**

Dieses Seminar vermittelt TK-Mitarbeitern ohne Vorkenntnisse im Bereich LAN und IP das erforderliche Wissen zur Planung und zum Betrieb von VoIP-Lösungen. Die Inhalte sind so gegliedert, dass Sie die Grundlagen schnell verstehen. Es werden die wichtigsten VoIP spezifischen Aspekte vorgestellt und unter Praxis-relevanten Gesichtspunkten beleuchtet. Die Themen erstrecken sich von IP und LAN Grundlagen hin zu Praxis relevanten Themen wie QoS, Jitter und Bandbreiten Fragen. Ziel ist es dem IP-Unkundigen die wichtigen Grundlagen der Netzwerk Technik kompakt und praxisnah zu vermitteln.

Preis: € 1.590,-- netto

Information Security Management mit ISO 27001 und BSI-Grundschutz, 25.04.-27.04.16 in Bremen**Garantietermin**

Dieses Seminar stellt den Aufbau und die nachhaltige Umsetzung eines standardisierten und zertifizierbaren Information Security Management System (ISMS) auf Basis von ISO 27001 und BSI IT-Grundschutz vor. Es wird dabei aufgezeigt, wie eine praxisgerechte Sicherheitslösung mit optimalem Aufwand erreicht werden kann.

Preis: € 1.890,-- netto

Virtualisierungstechnologien in der Analyse, 25.04.-26.04.16 in Bremen**Garantietermin**

Im Zuge stetig zunehmender Konsolidierung ist Virtualisierung längst zum Standard in jedem Rechenzentrum geworden. Doch der Blick hinter die Kulissen offenbart einen rapide wachsenden Komplexitätsgrad, dessen Beherrschung ein tieferes Verständnis dieser Technologie erfordert. In diesem Seminar werden die Zusammenhänge zwischen Server, Netzwerk und Storage im Umfeld der Virtualisierung analysiert.

Preis: € 1.590,-- netto

Umfassende Absicherung von Voice over IP und Unified Communications, 25.04.-26.04.16 in Bremen**Garantietermin**

Dieses Seminar zeigt die Risiken beim Einsatz von Voice over IP und Unified Communications auf und gibt den Teilnehmern einen Überblick über die zu ergreifenden Sicherheitsmaßnahmen. Auf Grundlage von Best Practices aus dem Beratungsgeschäft sowie den marktrelevanten Standards, wie z.B. der „Technischen Leitlinie für organisationsinterne Telekommunikationssysteme mit erhöhtem Schutzbedarf“ (TLSTK II) des BSI, werden den Teilnehmern die Anforderungen an eine Sicherheitskonzeption für TK und UC vermittelt.

Preis: € 1.890,-- netto

Zertifizierungen

ComConsult Certified Network Engineer

Lokale Netze

09.05. - 13.05.16 in Aachen
19.09. - 23.09.16 in Aachen

TCP/IP-Netze erfolgreich betreiben

14.03. - 16.03.16 in Berlin
20.06. - 22.06.16 in Bonn
24.10. - 26.10.16 in Bonn

Internetworking

04.04. - 08.04.16 in Aachen
04.07. - 08.07.16 in Aachen
14.11. - 18.11.16 in Aachen

Paketpreis für zwei 5-tägige und ein 3-tägiges Intensiv-Seminar € 6.180,-- netto (Einzelpreise: € 2.490,-- netto bzw. 1.890,-- netto)

ComConsult Certified Trouble Shooter

Trouble Shooting in vernetzten Infrastrukturen

10.05. - 13.05.16 in Aachen
27.09. - 30.09.16 in Aachen

Trouble Shooting für Netzwerk-Anwendungen

14.06. - 17.06.16 in Aachen
15.11. - 18.11.16 in Aachen

Paketpreis für beide Seminare inklusive Prüfung € 4.280,-- netto
(Seminar-Einzelpreis € 2.290,-- netto , mit Prüfung € 2.470,-- netto)

ComConsult Certified Voice Engineer

IP-Telefonie und Unified Communications erfolgreich planen und umsetzen

11.05. - 13.05.16 in Bonn
24.10. - 26.10.16 in Frankfurt

Session Initiation Protocol Basis-Technologie der IP-Telefonie

11.04. - 13.04.16 in Stuttgart
20.06. - 22.06.16 in Bonn
09.11. - 11.11.16 in Berlin

Umfassende Absicherung von Voice over IP und Unified Communications

25.04. - 27.04.16 in Bonn
04.07. - 06.07.16 in Stuttgart
28.11. - 30.11.16 in Bonn

Optionales Einsteiger-Seminar:

IP-Wissen für TK-Mitarbeiter
25.04. - 26.04.16 in Düsseldorf
19.09. - 20.09.16 in Frankfurt

Wir empfehlen die Teilnahme an diesem Seminar **"IP-Wissen für TK-Mitarbeiter"** all jenen, die die Prüfung zum ComConsult Certified Voice Engineer anstreben, ganz besonders aber den Teilnehmern, die bisher wenig bis kein Netzwerk Know How, insbesondere TCP/IP, DNS, SIP usw., vorweisen können.

Basis-Paket: Beinhaltet die drei Basis-Seminare
Grundpreis: € 5.100,-- netto statt € 5.670,-- netto
Optionales Einsteigerseminar: Aufpreis € 1.190,-- netto statt € 1.590,-- netto

Impressum

Verlag:
ComConsult Research Ltd.
64 Johns Rd
Christchurch 8051
GST Number 84-302-181
Registration number 1260709
German Hotline of ComConsult-Research:
02408-955300

E-Mail: insider@comconsult-akademie.de
<http://www.comconsult-research.de>

Herausgeber und verantwortlich
im Sinne des Presserechts:
Dr. Jürgen Suppan
Chefredakteur: Dr. Jürgen Suppan
Erscheinungsweise: Monatlich,
12 Ausgaben im Jahr

Bezug: Kostenlos als PDF-Datei
über den eMail-VIP-Service
der ComConsult Akademie

Für unverlangte eingesandte Manuskripte
wird keine Haftung übernommen
Nachdruck, auch auszugsweise
nur mit Genehmigung des Verlages
© ComConsult Research