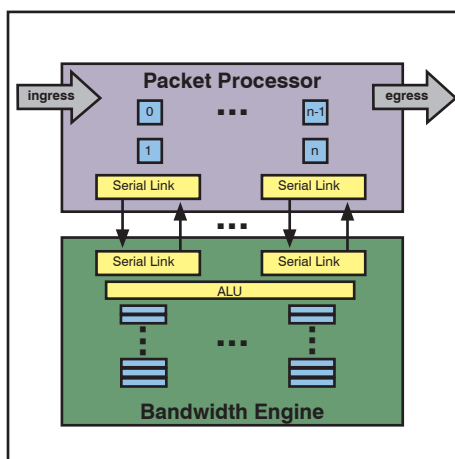


Die Zukunft privater Netze: Anwendungsdruck, Switch-Virtualisierung und 100G erfordern Carrier-Class-Kern-Netze und neue Hardware-Konzepte

von Dr. Franz Joachim Kauffels

In naher Zukunft wird die Netz-Kante zunehmend durch Software-Switches gebildet, die im Rahmen flexibler Virtualisierungsumgebungen etabliert werden. Bei Provider-Netzen kann man heute schon sehen, dass dies ein wesentliches Hilfsmittel zur eleganten Umsetzung von Konzepten der anwendungsorientierten Vernetzung (Application Aware Networks, Application Oriented Networks, Application Centric Infrastructure) ist. Entsprechende Produkte werden noch in diesem Jahr für Aufsehen sorgen. Eine unmittelbare Konsequenz ist aber auch, dass durch die gestiegenen Qualitätswünsche für die Anwendungen und die Umgestal-



tung der Netz-Kante die Anforderungen an das Kern-Netz nicht nur hinsichtlich der Netto-Übertragungsrate erheblich steigen werden. QoS, Traffic Shaping und andere Funktionen in dieser Richtung implizieren die Notwendigkeit einer erheblich aufwändigeren Bearbeitung der Pakete als bei einem Netz mit wenig oder gar keiner Steuerung in dieser Hinsicht.

weiter auf Seite 5

Zweitthema

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

von Dr. Franz Joachim Kauffels

Im Oktober 2013 haben die Analysten der Gartner Group wie in jedem Jahr die ihrer Auffassung nach wichtigsten Technologie-Trends für 2014 definiert. Gartner definiert eine Strategie dann als strategisch, wenn sie das Potential hat, innerhalb der nächsten drei Jahre einen signifikanten Einfluss auf ein Unternehmen zu haben. Faktoren, die

die Signifikanz unterstützen, sind ein hohes Potential zur Umwälzung der IT oder des Geschäfts, die Notwendigkeit für erhebliche finanzielle Investitionen oder ein hohes Risiko in dem Fall, dass man die Entwicklung zu spät erkennt und umsetzt. Noch nie gab es einen so starken Zusammenhang zwischen den Trends: die Top 8 sind alle untereinander

verwoben. Es entsteht ein in Teilen neues, durchaus faszinierendes Bild der zukünftigen IT. In diesem Artikel nennen wir nicht nur die Trends, sondern analysieren ihre Bedeutung für privat betriebene Informations-Infrastrukturen von Unternehmen und Organisationen.

weiter Seite 20

Geleit

Server streben Kontrolle über ihre Netzwerke an: „intelligente Dienste – dumme Netzwerke“, ist das die Zukunft?

von Dr. Jürgen Suppan

ab Seite 2

Standpunkt

Software Updates, die Geißel der IT

auf Seite 19

Aktueller Kongress

Netzwerk Forum 2014

ab Seite 4

Aktuelles Seminar

**Intensiv-Tag Video:
Markttrends und Techno-
logien**

auf Seite 18

Zum Geleit

Server streben Kontrolle über ihre Netzwerke an: „intelligente Dienste – dumme Netzwerke“, ist das die Zukunft?

Erinnern Sie sich an ATM und RSVP? Als Tiger abgesprungen, als Maus gelandet. Immer wieder hat es in den letzten 20 Jahren den Versuch gegeben eine Applikations-spezifische Ende-zu-Ende Kontrolle in Netzwerken zu etablieren. So dass eine Applikation eben sicher sein konnte, dass ihr spezifischer Bedarf erfüllt wird. Alle diese Versuche sind an der Komplexität der Aufgabe und der Umsetzung für den Betreiber gescheitert.

Nun ist es wieder so weit. Aus der Server-Welt kommt der Wunsch, Netzwerke durch Applikationen zu kontrollieren und dabei so zu gestalten, dass der Bedarf der Applikationen erfüllt wird. Die Idee ist, dass Applikationen zusammen mit passenden virtuellen Maschinen und einem „Netzwerk-Container“ ausgeliefert werden, der es ihnen erlaubt, ihr eigenes Netzwerk als zusätzliches Layer auf der Basis eines bestehenden IP-Netzwerkes umzusetzen. Was bedeutet „eigenes“ Netzwerk?

Folgende Eigenschaften sind damit verbunden:

- Eine Applikation „sieht“ nur ihre Kommunikationspartner in anderen virtuellen Maschinen, Speichersystemen etc.
- Die konkrete Hardware-Architektur des darunter liegenden realen Netzwerkes ist dabei unerheblich. Im Prinzip können Teile einer Applikation in unterschiedlichen Standorten oder auch in der Cloud liegen
- Dieses Applikations-eigene Netzwerk wird komplett gegen den Rest der Welt abgeschottet. Die Applikation übernimmt die volle Verantwortung für ihre Netzwerk-Sicherheit
- Um dies zu erreichen, werden in dem „Netzwerk-Container“ alle notwendigen Netzwerk-Instanzen wie Firewalls, Load-Balancer als virtuelle Instanzen gleich mit geliefert

Man kann auch sagen, dass eine Applikation ihr eigenes virtuelles Netzwerk mit sich bringt. Hört sich das ziemlich verwegen für Sie an? Nun, Martin Cassado von Nicira/VMware spricht über „a change to networking with the same level of disruption virtualization produced to the server



environment“. Also auf Deutsch gesagt, werden Netzwerke einer genauso umwälzenden Veränderung ausgesetzt sein wie wir das auf der Serverseite durch Virtualisierung hatten.

Hat der Mann vollständig seinen Verstand verloren oder ist da tatsächlich eine reale Gefahr gegeben? Nun, zumindest ist bei den traditionellen Netzwerk-Herstellern eine erhebliche Nervosität zu beobachten. Als Gegenbewegung werden Application-aware Networks ins Leben gerufen (die allerdings in der Tat einen Beigeschmack vom ATM-Desaster aufkommen lassen). Also, wo stehen wir wirklich:

- Im Umfeld von Servern haben wir inzwischen deutlich mehr Software-Ports in virtuellen Switches als in der Hardware
- Wurden den virtuellen Switches zu Beginn mehr Probleme als Lösungen zugeordnet, so hat sich ihr Leistungsvermögen in den letzten Jahren immer weiter ausgedehnt. Nicht nur nähern sie sich der 10Gig-Grenze immer mehr an, auch der sonstige Funktionsumfang nimmt zu. Virtuelle Switches müssen als Netzwerk-Instanzen sehr ernst genommen werden
- Tatsächlich haben die Netzwerker den Bedarf der Server- und Applikations-Fraktion arrogant verschlafen. Seit Jahren wird auf der Server-Seite intensiv daran gearbeitet, eine One-Touch-Installation auch komplexer Anwendungen umzusetzen, so dass der alte Widerspruch der Virtualisierung, eine neue

Applikation in Minuten ausgerollt zu haben, endlich Wirklichkeit wird. Tatsächlich ist das Netzwerk der Bereich, der diesem Ziel am intensivsten im Wege steht.

Natürlich treffen hier auch unterschiedliche Hersteller-Interessen aufeinander. VMware muss sein Geschäftsmodell dringend ausdehnen. Mit einfacher Virtualisierung ist längst kein Blumentopf mehr zu gewinnen. Schon gar nicht zu den Lizenzpreisen von VMware. Und schon gar nicht im direkten Kampf mit Microsoft, der zunehmend bedrohlicher für VMware wird. Wer als erster eine glaubwürdige Automatisierung von Applikationen von der Installation bis zum selbstoptimierenden Betrieb liefert, der hat wieder Argumente, um mehr Geld für Software zu verlangen. Also lohnt es sich aus der Sicht von VMware auch, eine Milliarden-Investition in Nicira zu tätigen und damit auch gleichzeitig ein Spitzenteam von Netzworkern der Stanford-Universität einzukaufen.

Der Konflikt „VMware NSX kontra Cisco“ ist ein Konflikt der Software-Hersteller mit der gesamten Netzwerk-Branche. Allerdings trifft dieser Konflikt keinen so intensiv wie Cisco. Cisco ist der Anbieter, der Netzwerke über Zusatzdienste „intelligenter“ macht und so für seine dicken Schlachtschiffe der Bauart Catalyst 6000 und Nexus deutlich mehr Geld verlangen kann als für ein einfaches Standard-Netzwerk. Während sich Anbieter wie Alcatel-Lucent, Avaya, Extreme und HP darauf konzentrieren, Standard-Komponenten mit Standard-Chips zu liefern, ist Cisco immer auf der Suche nach einer Zusatzleistung, die verkauft werden kann. Zum Teil mit Grund. Verfahren wie LISP und OTV sind ein gutes Beispiel dafür, dass es einen Bedarf über etablierte Standards hinaus gibt.

Wer wird diesen Kampf gewinnen? Wird die Software-Industrie die Hardware-Netzwerke überrollen? Nun, so einfach ist es natürlich nicht. Die schöne heile Welt der Netzwerk-Container und Virtualisierung setzt dummerweise ein leistungsfähiges und funktionierendes Basis-Netzwerk voraus. Und das hat im Zeitalter eines 10/40/100-Gigabit-Designs seine eigenen Herausforderungen. Aber auf jeden Fall wird es weiterhin aus Hardware bestehen.

Server streben Kontrolle über ihre Netzwerke an: „intelligente Dienste – dumme Netzwerke“, ist das die Zukunft?

Die Kernfrage liegt in der Floskel „Intelligente Dienste – dumme Netze“. Hier deutet sich an, dass wir eine Intelligenz brauchen, die es bisher nicht gibt. Und die Frage ist, wer diese in Zukunft liefern wird: die Serverseite oder die Netzwerk-Fraktion. Ein aktuelles Beispiel dazu: virtualisierte Server brauchen bei wandernden virtuellen Maschinen eine ortsneutrale IP-Adressierung, die nicht an lokale Subnetzbereiche gebunden ist. Das kann man natürlich mit den Nicira-Overlays oder einer sogenannten Edge-Provisioning realisieren. Man kann aber auch auf der Netzwerkseite neue Dienste wie LISP schaffen. LISP separiert zwar die Applikationsnetze nicht und schafft somit keine neue Form von VLAN wie das NSX macht, aber es löst das Adressierungsproblem. Genauso könnte man sich eine neue Form der Netzwerk-Virtualisierung und eine neue VLAN-Technik auf Netzwerk-Ebene vorstellen. Tatsache ist, dass die Netzwerk-Entwickler hier seit vielen Jahren gemütlich schlafen. Die Kritik an QoS und VLANs ist so alt wie die Geschichte der Netzwerke. Und immer noch gibt es keine runde Lösung dafür. In diese Unfähigkeit der Weiterentwicklung stößt das Nicira/VMware-Lager. Hoffen wir, dass jetzt die Netzwerk-Entwickler aufwachen und endlich wieder kreativ werden. Tatsache ist, dass Eile geboten ist. Natürlich weiß auch Martin Cassado, dass er sich im Moment noch auf das vorhandene Hardware-Basis-Netzwerk abstützen muss und darauf keinen Einfluss hat. Aber wie es der Zufall so will, ist er auch einer der Gründungsväter von SDN. Ein Schelm, der Böses dabei denkt. Hat VMware wirklich das viele Geld nur für NSX bezahlt oder geht es in Wirklichkeit um mehr? Ist NSX vielleicht in Wirklichkeit nur der erste Baustein einer wesentlich weitergehenden Strategie?

Auf jeden Fall eine spannende Situation. Wir greifen das Thema auf dem ComConsult Netzwerk Forum 2014 auf und kombinieren das mit der Diskussion um neue Design-Ansätze für den Campus und die sichere und wirtschaftliche Integration mobiler Teilnehmer.

Auf dem Forum diskutieren wir Fragen wie:

- Brauchen wir das alles? Reicht nicht ein gutes 10/40/100 Bandbreiten-Design aus, um auch wirklich jeden Bedarf abzudecken?
- Wie sieht denn ein modernes 10/40/100 Design denn zum Beispiel für den Campus und den Access-Bereich aus? Wir haben in den letzten Jahren viel über das Rechenzentrum diskutiert,

aber wo stehen wir in den anderen Netzwerk-Bereichen? Lassen sich die neuen Fabric-Ansätze aus dem RZ zum Beispiel sinnvoll auf den Campus übertragen?

- Das universelle Netzwerk für alles und jeden existiert in dieser Form sicher nicht mehr so wie früher. Nichts unterstreicht das so stark wie die Diskussion um die RZ-Netzwerke. Der Ost-West-Verkehr zwischen Servern und Speicher-Systemen in Kombination mit der potenziellen Ablösung von Fibre Channel hat hier ganz spezielle Anforderungen geschaffen, die nur mit spezi-

ellen Verfahren wie DCB zu erfüllen waren. Aber kann trotzdem der Anspruch eines Dienst-neutralen Netzwerks aufrecht erhalten werden, auch wenn immer mehr Spezialfälle in unsere Netzwerke Einzug halten?

- Wie sieht es dann mit der Integration mobiler Teilnehmer aus? Hier besteht ein erheblicher Sicherheitsbedarf auch im Netzwerk. Wie decken wir denn „Dienstneutral“ ab?

Ihr
Dr. Jürgen Suppan

Kongress

ComConsult Netzwerk Forum 2014 24.03. - 26.03.14 in Königswinter

- **Dienstneutralität kontra Applikations-Integration: Kampf um die Steuerung**
- **Netzwerk-Design in Zeiten von 10/40/100: wie sieht die Zukunft aus?**
- **Mobile Endgeräte im Netzwerk: teuer, unsicher, unbeherrschbar?**

Traditionell ist die Dienstneutralität das Hauptmerkmal aller Netzwerke. Alle bisherigen Versuche, eine Applikations-spezifische Ende-zu-Ende Kontrolle und Garantie einzuführen sind in den letzten 20 Jahren gescheitert. Auch Prioritäten und Bandbreitenreservierungen für Verkehrsklassen sind trotz diverser Standards umstritten. Nun haben sich Netzwerke auf der Server-Seite in den letzten 5 Jahren deutlich verändert. Hier dominieren inzwischen die Software-Ports in den virtuellen Switches gegenüber den Hardware-Ports. Aus Sicht der Applikationen und virtuellen Maschinen spielt sich das Netzwerk mehr und mehr in Software ab. Verbindet man das mit dem zunehmenden Bedarf nach „One-Touch-Installationen“ und der schnellen Inbetriebnahme auch komplexer Applikationen innerhalb von Minuten, dann hat sich in der Server-Welt der Netzwerk-Schwerpunkt von der Hardware in die Software verschoben.

Die Fragen der Zukunft sind:

- Wie sieht ein Konzept zur schlüssigen Integration mobiler Endgeräte aus?
- Was bedeutet das für unsere Netzwerke?
- Wie sehen Gesamtarchitekturen von Endgerät bis zum App-Shop aus?
- Wie kann die notwendige Sicherheit wirtschaftlich umgesetzt werden?

Das ComConsult Netzwerk Forum 2014 stellt sich den herausragenden Fragen der Entwicklung der Unternehmensnetzwerke. Nach 10 Jahren des Konzept-Stillstands ist das traditionelle Netzwerk unter Druck sowohl aus der Server- als auch aus der Endgeräte-Welt. Kombiniert man dies mit dem notwendigen Bandbreitenwechsel auf 10/40/100, dann steht das Unternehmensnetzwerk der Zukunft auf dem Prüfstand.

Wir analysieren für Sie wohin die Netzwerk-Entwicklung geht, wer die Kontrahenten sind, welche Vor- und Nachteile die verschiedenen Alternativen für Ihr Unternehmen haben und wie relevant einzelne Technologien sind. Investieren Sie zukunftsicher in Ihre Netzwerke, das ComConsult Netzwerk Forum 2014 unterstützt Sie dabei.

Moderation: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr.-Ing. Behrooz Moayeri
Preise: € 2.390,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Aktueller Kongress

ComConsult Netzwerk Forum 2014

Praxis und Zukunft von Unternehmensnetzen**24.03. - 26.03.14 in Königswinter**

Die ComConsult Akademie veranstaltet vom 24.03. bis 26.03.2014 ihr "ComConsult Netzwerk Forum 2014" in Königswinter.

Traditionell ist die Dienstneutralität das Hauptmerkmal aller Netzwerke. Alle bisherigen Versuche, eine Applikations-spezifische Ende-zu-Ende Kontrolle und Garantie einzuführen sind in den letzten 20 Jahren gescheitert. Auch Prioritäten und Bandbreitenreservierungen für Verkehrsklassen sind trotz diverser Standards umstritten. Nun haben sich Netzwerke auf der Server-Seite in den letzten 5 Jahren deutlich verändert. Hier dominieren inzwischen die Software-Ports in den virtuellen Switches gegenüber den Hardware-Ports. Aus Sicht der Applikationen und virtuellen Maschinen spielt sich das Netzwerk mehr und mehr in Software ab. Verbindet man das mit dem zunehmenden Bedarf nach „One-Touch-Installationen“ und der schnellen Inbetriebnahme auch komplexer Applikationen innerhalb von Minuten, dann hat sich in der Server-Welt der Netzwerk-Schwerpunkt von der Hardware in die Software verschoben.

Dies wird eine Reihe von Fragen auf:

- Liegt die Zukunft der Server-Vernetzung in der Software?
- Hat das dienstneutrale Netzwerk damit ausgedient?
- Wie kann eine automatische Netzwerk-Konfiguration bei der Installation von Servern erreicht werden?
- Wie kann der Bedarf dynamischer virtueller Umgebungen am besten erfüllt werden?

eller Umgebungen am besten erfüllt werden?

- Wer wird die Zukunft bestimmen: Anbieter wie VMware mit NSX oder die traditionellen Netzwerker mit Application Aware Networks?

Gleichzeitig sind wir mitten im Übergang von 1/10 Gigabit-Netzwerken zu 10/40/100 Gigabit. Dies geht einher mit einer Vielzahl sehr verschiedener Gestaltungsmerkmale, die unter dem Begriff Fabric-Design laufen. Auch im Design wird dabei die Neutralität des Netzwerkes immer mehr mit einem Fragezeichen versehen. Hier ist nicht nur das RZ-Netzwerk eine treibende Kraft, sondern die hier entwickelten Technologien dehnen sich immer mehr auf den Campus aus. Abgerundet wird das mit der sich ausdehnenden Problematik der Integration mobiler Endgeräte.

Auch hier haben sich brisante Fragen entwickelt:

- Wie sieht ein übergreifendes 10/40/100 Netzwerk-Design in Zukunft aus?
- Wie sieht das Campus-Netzwerk der Zukunft aus, wird es zunehmend von Design-Prinzipien des Rechenzentrums beeinflusst?
- Ist das Ende modularer Switches nahe?
- Brauchen wir wirklich mehr als 10 Gig parallele Verkehrsströme?

Der Endgerätebereich wird sich in den meisten Unternehmen in den nächsten 5 Jahren deutlich verändern. Mobile Endge-

räte und neue Endgerätevarianten werden die Zukunft bestimmen. Das führt nicht nur zur Frage des besten physikalischen Anschlusses, sondern vor allem zur Frage eines schlüssigen Gesamtkonzepts von der Physik über die Architektur bis zur Sicherheit.

Die Fragen der Zukunft sind:

- Wie sieht ein Konzept zur schlüssigen Integration mobiler Endgeräte aus?
- Was bedeutet das für unsere Netzwerke?
- Wie sehen Gesamtarchitekturen von Endgerät bis zum App-Shop aus?
- Wie kann die notwendige Sicherheit wirtschaftlich umgesetzt werden?

Das ComConsult Netzwerk Forum 2014 stellt sich den herausragenden Fragen der Entwicklung der Unternehmensnetzwerke. Nach 10 Jahren des Konzept-Stillstands ist das traditionelle Netzwerk unter Druck sowohl aus der Server- als auch aus der Endgeräte-Welt. Kombiniert man dies mit dem notwendigen Bandbreitenwechsel auf 10/40/100, dann steht das Unternehmensnetzwerk der Zukunft auf dem Prüfstand.

Wir analysieren für Sie wohin die Netzwerk-Entwicklung geht, wer die Kontrahenten sind, welche Vor- und Nachteile die verschiedenen Alternativen für Ihr Unternehmen haben und wie relevant einzelne Technologien sind. Investieren Sie zukunftsicher in Ihre Netzwerke, das ComConsult Netzwerk Forum 2014 unterstützt Sie dabei.

Fax-Antwort an ComConsult 02408/955-399

Anmeldung

ComConsult Netzwerk Forum 2014

Ich buche den Kongress **ComConsult Netzwerk Forum 2014**

☐ vom 24.03. - 26.03.14 in Königswinter zum Preis € 2.390,-- netto

☐ vom 24.03. - 25.03.14 in Königswinter zum Preis € 1.990,-- netto

☐ am 26.03.14 in Königswinter zum Preis € 990,-- netto



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

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift

Schwerpunkthema

Die Zukunft privater Netze:

Anwendungsdruck, Switch-Virtualisierung und 100G erfordern Carrier-Class-Kern-Netze und neue Hardware-Konzepte

Fortsetzung von Seite 1



Dr. Franz-Joachim Kauffels ist Technologie- und Industrie-Analyst und Autor. Seit über 30 Jahren unabhängiger, kritischer und oft unbequemer Bestandteil der Netzwerkszene. Verfasser von über 20 Büchern in über 70 Ausgaben sowie über 2000 Artikeln, Videos und Reports.

Leider werden damit beim Übergang auf 100G die Grenzen bestehender Hardware-Konstrukte offenbar. Der Tribut für erheblich erhöhte Qualität und Flexibilität bei den Anwendungen ist die Notwendigkeit völlig neuer Hardware-Konzepte. Die so entstehende Vision zukünftiger privater Netze räumt Unfug zur Seite, über den heute noch gerne diskutiert wird und gibt Richtlinien für die Planung.

1995 haben wir heftig über 100 Megabit-Netze diskutiert. Der damalige Platzhirsch FDDI wurde von Neuentwicklungen wie 100 VG AnyLAN, 100 BASE-T oder ATM herausgefordert. Damals ging es sehr schnell mit Produkten und Einführung, gewonnen haben die unter dem Begriff „Fast Ethernet“ standardisierten Varianten.

1998, nur drei Jahre später, kam die erste 1 Gigabit-Produktwelle. Zum zweiten Male wurde das Versprechen „zehnfache Leistung zum dreifachen Preis“ eingelöst.

2001, wiederum nur drei Jahre später, kamen die ersten Produkte für 10 GbE. Zum ersten Male war die Standardisierung schneller als die Hersteller, so dass uns Diskussionen zwischen verschiedenen Basistechnologien erspart blieben. Zum ersten Male wurde aber auch der Providerbereich ernsthaft mit einbezogen und hier gab es auch die ersten Produkte. In Verbindung mit den Carrier Ethernet Standards hatte sich Ethernet in allen Netzwerk-Kategorien durchgesetzt, was in Folge dazu führte, dass es heute der Welt-Standard für die Datenübertragung in Netzwerken ist.

2004 war erkennbar, wie die neuen Standards für 40 und 100 GbE aussehen würden.

Aber dieses Mal hat die Technik offenbar Probleme, das Ziel „zehnfache Leistung zum dreifachen Preis“ einzulösen.

Jetzt haben wir schon 2014 und bis auf eine Hand voll Lösungen, die ausschließlich den Providern vorbehalten sind, gibt es nur wenig 100G-Geräte. Für den Bereich der privaten Unternehmensnetze, hier besonders die Rechenzentren, hat sich noch kein wirkliches Angebotsszenario entfaltet.

Zehn Jahre Stillstand auf der physikalischen Übertragungsebene hat es in der Netzwerkwelt noch nie gegeben. Eine Provider-Lösung für 100G und Vielfache kann man im Grunde genommen so komplex und teuer machen, wie man möchte, der Provider wird sie am Ende doch einsetzen, wenn es keine Alternative gibt. Grade die Aufrüstung der Mobilfunknetze von 3G auf LTE hat das eindrucksvoll gezeigt. Jeder Provider, der heute LTE anbietet, hat einen hohen zweistelligen Milliardenbetrag dafür investiert. Und es funktioniert.

Für private Unternehmensnetze ergibt sich das Problem, dass die Technik kompakt, preiswert und energiesparend sein muss. Und hier gibt es echte Probleme, die wir in diesem Artikel darstellen.

Seit jeher wurde die Einführung einer neuen Geschwindigkeitsstufe von der Diskussion begleitet, ob man nicht durch die Kombination einer etwas intelligenteren Netzwerk-Steuerung und der alten Technik in etwa das Gleiche erreichen kann wie mit der neuen Evolutionsstufe. Bei Fast Ethernet versus Gigabit Ethernet war dies die Diskussion um die (einfache) Priorisierung. Bei Gigabit Ethernet versus 10 GbE wurde die Diskussion um differenzierte Priorisierung geführt. In beiden Fällen konnte sowohl durch Warteschlangen als auch durch praktische Messungen der simple Satz: „Auf die Dauer hilft nur Power“ bestätigt werden. Jedes Priorisierungsverfahren basiert darauf, die begrenzte Ressource dadurch zu verwalten,

dass Datenströme z.B. von Anwendungen unterschiedlich behandelt werden. Da die Ressource aber insgesamt deutlich begrenzt ist, geht eine Bevorzugung bestimmter Datenströme immer zu Lasten anderer, die unter der Priorisierung deutlich schlechter laufen als mit einem ungesteuerten Verfahren.

Natürlich ist der Betreiber eines privaten Rechenzentrums verunsichert. Er fragt sich, ob man nicht in Zukunft die meisten Aufgaben der Vernetzung einfach mit Software erledigen kann, anstatt, wie gewohnt, die Hardware immer weiter aufzurüsten. Äußerungen, wie sie von Marktforschungsunternehmen kommen, sind zunächst einmal Wasser auf diese Mühlen. So spricht Infonetics davon, dass „wenn die durchschnittliche Anzahl von VMs pro Server in 2015 30 erreicht, virtuelle Switches, die auf General Purpose Rechnern laufen, die neuen Netzkanten bilden werden.“

In einem früheren Artikel im Netzwerk Redesign Sonderheft 2013 bin ich bereits der Frage nachgegangen, welche Art von mit einem Netzwerk assoziierten Funktionen durch Software günstig implementiert werden können und welche eher nicht. Das Ergebnis war, dass Funktionen, die viele reguläre Ausdrücke verarbeiten, meist nur mit entsprechender Zusatzhardware implementiert werden können, andere Funktionen, wie Systeme zur intelligenten Lastverteilung, sehr gut komplett in Software mit Unterstützung einer oder mehr VMs realisiert werden können. Ob sich dadurch allerdings wirklich wie versprochen Kostenvorteile ergeben, hängt nicht nur von den Kosten für die Server Hardware, sondern vor allem von den Kosten für die notwendigen Lizenzen des Virtualisierungssystems ab. Hier hatte das Lizenzmodell von VMware eine eher Konzept-vernichtende Wirkung.

Die Zukunft privater Netze

In diesem Artikel geht es darum, ähnliche Aussagen auch für die Netze selbst und nicht nur für die assoziierten Funktionen zu treffen.

1. Software: das Ende von ToR-Switches & Co

Wir haben ja schon häufiger darüber diskutiert, ich verweise hier auch auf die entsprechenden Ausführungen von Herrn Dr. Suppan, dass die Anzahl der VMs pro Prozessor aufgrund der technischen Entwicklung immer weiter steigen wird und die von Infonetics genannte Zahl ist sicher keine schlechte Schätzung. Für die üblichen Anwendungen in privaten RZs wird Virtualisierung weiterhin das dominierende Betriebskonzept sein. Die heute eher statischen Lösungen werden sich nach Möglichkeiten und Bedarf zu dynamischen Lösungen weiterentwickeln.

Wir diskutieren seit mindestens drei oder vier Jahren über die Anbindung von VMs und es gibt eine Reihe von marktgängigen Alternativen, wer die noch nicht kennt, sollte sich schleunigst wenigstens das Video von Herrn Höchel-Winter dazu anschauen. Diese Alternativen erfahren laufend Zuwachs, aber wie ich bereits in meinem Artikel zur hoch leistungsfähigen VM-Migration im Insider 12/2013 geschrieben habe, wird sich dieses Spektrum noch erweitern.

Maßgeblich für den Erfolg der Implementierung eines Rand-Switches in Software ist die Hardware-Unterstützung. Es ist eigentlich nur ein kleiner Schritt notwendig: die Adapterkarten und Line Cards, die es natürlich auch bei virtualisierten Switches geben muss, müssen in der Lage sein, als unmittelbare Kommunikationsziele differenziert VMs ansprechen zu können und nicht mehr physikalische Server als Ganzes. Dafür gibt es schon seit vielen Jahren den Standard SR-IOV. Er ist zu Beginn oftmals falsch interpretiert worden. Das primäre Ziel von SR-IOV ist es weniger, irgendwelche vorsortierte Warteschlangen zu definieren, die dann von den Hypervisoren sozusagen mit den VMs verbunden werden, sondern vielmehr, den VMs direkte physikalische Kommunikationsziele zu geben, die vermöge der Vorsortierung und Warteschlangenverwaltung dann auf einer Leitung einer höheren Kapazität wie z.B. 10 GbE konzentriert werden. Mellanox ist der erste Hersteller, der in seinen Adapterkarten eine derart unmittelbare Nutzung des Netzes durch VMs anbietet, er wird aber nicht der letzte sein.

Ergänzend zu einer solchen direkten Hardware-Unterstützung der Kommunikation benötigt man dann für die Implementierung

eines Rand-Switches nur noch eine Software, die die üblichen Funktionen eines solchen Switches nachbildet. Diese Software kann man sich in etwa so vorstellen wie eine verschärfte Variante eines Cisco Nexus 1000v. Der Unterschied besteht prinzipiell nur darin, dass die Software nicht selbst die Pakete switcht, sondern vermöge einer geeigneten Schnittstelle wie SR-IOV direkt mit den physikalischen Switches in der unterliegenden Adapter- oder Line-Card Ebene kommuniziert.

Im Zuge der technischen Entwicklung sind natürlich auch andere technische Alternativen denkbar, um die Ebene der ToR- oder EoR-Switches samt ihrer Verkabelung und sonstiger Verwaltungsprobleme komplett zu entsorgen. Die Möglichkeiten, die sich ergeben, wenn man kleine Switch-ASICs direkt auf die Server-Boards bringt, hatte ich ja schon vor Jahren ausführlich diskutiert.

Die möglichen Vorzüge einer Software-Realisierung der Randfunktionen sind fast schon so überwältigend, dass sie einen eventuellen Kostenvorteil verblässen lassen:

- unbegrenzte Skalierbarkeit in der Leistung durch Hinzufügen von VMs
- unterbrechungsfreier Betrieb mit VM-basierten Verfügbarkeitskonzepten
- optimale Ausnutzung der Ressourcen bei dynamisierten Umgebungen

Bei der Implementierung von Application Gateways durch VMs konnte man sehen, dass es einen Punkt gibt, an dem sich die Leistung nicht mehr weiter steigern lässt, wenn man mehr VMs hinzufügt. Das ist aber ein konstruktives Problem der Virtualisierungssoftware und der Implementierung der VM-Kommunikation. Ist die Kommunikation zwischen den VMs langweilig, bekommt man bei zu vielen VMs ein Performance-Problem. Vergleichbares gilt auch für alle Systeme zur Realisierung eines unterbrechungsfreien Betriebes. Sie sind sehr davon abhängig, wie die Qualität der VM-Migration ist.

Große RZs z.B. von Providern haben als wesentliches Betriebskonzept eine möglichst bipolare Auslastung der Server: ein Server ist entweder ganz ausgelastet oder im Ruhezustand. Sobald die Last steigt, werden Server aus dem Ruhezustand aktiviert, fällt sie, dürfen sie wieder schlafen. Das senkt die Kosten, vor allem beim Strom, enorm. Mit Software-basierten Randswitches könnte man natürlich das Gleiche machen: sie werden überhaupt nur dann aktiviert, wenn der Server, den sie versorgen, auch aktiv wird.

Jetzt werden viele noch Bedenken haben, ob das denn wirklich so funktioniert wie beschrieben oder ob es nicht irgendwo schwere Haken gibt.

Ich möchte es einmal anders herum ausdrücken: es gibt bereits heute Bereiche, in denen die Implementierung von Switching-Funktionen in Software nach heutigem Stand existentiell ist. Selbst ich war davon überrascht, eine in dieser Hinsicht sehr weit gehende konstruktive Struktur ausgerechnet im absoluten Höchstleistungs-Segment zu finden: DWDM-Kernetze. In der Produktlinie der CRS-Router von Cisco gibt es zwar jede Menge aufregender Hardware, wie z.B. integrierte Schnittstellenoptik mit kohärenter Decodierungstechnik für die Übertragung auf Distanzen bis zu 3000 km, aber die gesamte Kontrollstruktur basiert auf Virtualisierung, die letztlich von UCS-Servern implementiert wird. Das aktuelle Spitzenmodell hat Line-Card Slots mit einer Leistung von 400 G pro Slot und kann nach Aussagen des Herstellers bis zu einem Petabit/s skaliert werden. Übrigens arbeiten diese Systeme nicht mit irgendwelchen proprietären Verfahren, sondern primär mit IPoDWDM, also der standardisierten direkten Verarbeitung von IP-Paketen über DWDM-Systeme. Gesteuert wird das Ganze mit OAM&P-Software nach G.709.

Noch deutlicher wird die Tendenz beim neuen Network Convergence System 6000 des gleichen Herstellers. Das ist eine Familie von integrierten Routing- und Transportsystemen für den Providerbereich mit vielen Innovationen in Übertragungstechnik, Organisation und Software. Ziele sind neben der deutlichen Senkung der Betriebskosten die erhebliche Erweiterung der Skalierbarkeit im Hinblick auf die Anforderungen des IoT und eine five9-Verfügbarkeit. Zum ersten Mal wird die Software bis auf das Niveau von Line-Cards hinunter virtualisiert. Die technische Basis bilden spezielle Kommunikationsprozessoren, die Cisco selbst entwickelt hat, und an den geeigneten Stellen UCS-Systeme. Natürlich ist hier nicht VMware die Virtualisierungsbasis, sondern eine Linux-Erweiterung. (siehe Abbildung 1)

Es ließe sich noch viel über diese extrem leistungsfähigen Systeme sagen, aber da ich weiß, dass der durchschnittliche Leser daran eher weniger interessiert ist, belassen wir es mit folgendem Statement:

- *die Virtualisierung von Edge-Switching-Funktionen ist nicht nur technisch möglich, sondern in den kommenden Jahren sogar eine notwendige Voraussetzung für die Schaffung hoch leis-*

Die Zukunft privater Netze

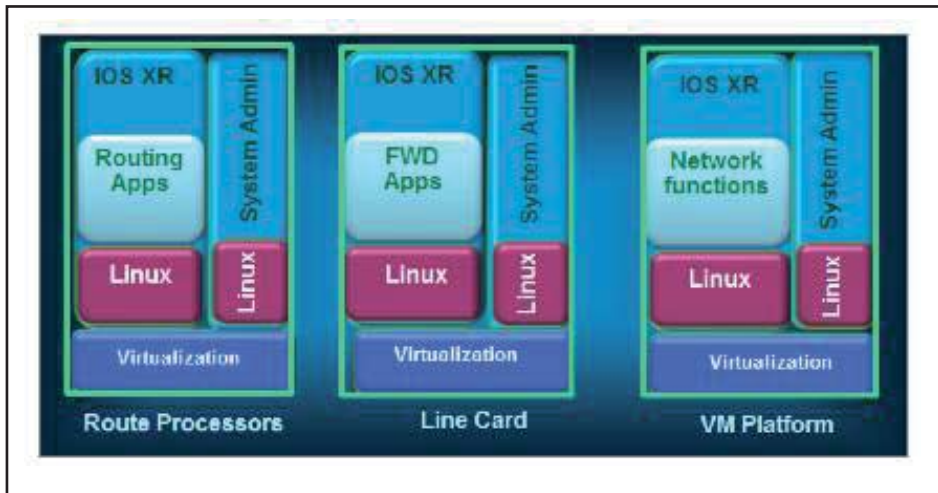


Abbildung 1: Vollständig virtualisierte Software in den NCS-Systemen

Quelle: Cisco

tungsfähiger weit skalierbarer zuverlässiger Kommunikations-Infrastrukturen

höheren Grad von Konvergenz in den Kernnetzen einher

Seit mindestens zwei Jahren überziehen uns die Hersteller mit dem Gedanken, Netze mit Software zu steuern (SDN) und letztlich zu sog. „Anwendungs-bewussten“ Netzwerken (SOA, ACI) zu kommen.

Man sollte sich aber klar machen, dass derartige Techniken die Datenströme aus den Anwendungen zwar „besser“ sortieren, das Kern-Netz dadurch aber insgesamt nicht leistungsfähiger wird. Das kennt jeder von einem alten PC, der mit neuer Software nur selten schneller wird, während die Wahrscheinlichkeit eines Total-Kollapses immer weiter steigt. Völlig pervers finde ich dabei persönlich, dass sich auch vermehrt Hersteller in die Diskussion einmischen wollen, die in der Vergangenheit nichts zu Netzen beigetragen haben, allen voran VMware. Der Gedanke, Netze auf eine Basis-Funktionalität zu reduzieren, dieses dann mit einem Overlay in Art des Mäntelchens der christlichen Nächstenliebe zu überziehen und den Rest mit Software zu optimieren, kann nur von jemandem kommen, der die innere Natur von Netzen überhaupt nicht begriffen hat.

Die Leistung von Kern-Netzen wird definitiv durch Hardware bestimmt und durch sonst rein gar nichts. Ist die Übertragungstechnik schnell, können Datenströme schnell bedient werden. Ist die Übertragungstechnik langsam, helfen auch keine Berge von Steuerungs- und Optimierung-Software, ganz im Gegenteil.

Deshalb kommen wir zu einer weiteren These:

- *die Virtualisierung von Edge-Switching-Funktionen geht zwangsweise mit einem höheren Konzentrationsgrad und einem*

Das ist eigentlich auch ohne ausschweifende Erklärung zu verstehen. Gebe ich jeder der durchschnittlich 30 VMs auf einem einzigen Prozessor eine grundsätzliche Kommunikationsleistung von nur einem einzigen Gigabit/s, muss ich ohne wilde Überbuchung für einen einzigen Prozessor bereits eine 40 GbE-Leitung vorsehen. Ein virtualisierter Switch versorgt aber mehr als einen Prozessor. Also wird es schnell darauf hinauslaufen, dass im Kernnetz 100 GbE-Verbindungen benutzt werden müssen. Aber damit nicht genug:

- *Um das Kern-Netz nicht ausufern zu lassen, werden Methoden zur Verkehrs-Steuerung wenigstens im Leistungsumfang der DCB-Funktionen eine absolute Notwendigkeit.*

Ungesteuerte Netze sind genau so wie fixe Überbuchungsraten nicht mehr Stand der Technik. Auf einer konvergenten Leitung versammeln sich Flows unterschiedlicher Anwendungen. Damit das geordnet geschehen kann, sind QoS-Definitionen und deren effektive Durchsetzung absolut elementar.

Praktisch alle Leser werden jetzt sagen: „Wir brauchen aber kein 100GbE.“ Abgesehen davon, dass diejenigen, die so etwas am Lautesten sagen, erfahrungsgemäß die sind, die es als erste einsetzen, ist das auch gar nicht der Punkt. Wichtig ist mir darzustellen, warum der Betreiber eines durchschnittlichen Corporate Networks sich weder in fragwürdige proprietäre Experimente einlassen sollte noch den über Jahrzehnte erfolgreichen Weg einer Standard-basierten Lösung verlassen muss. Aber dazu später.

2. Zu Risiken und Nebenwirkungen in der Hardware

In den vergangenen Jahrzehnten sind wir immer dadurch verwöhnt worden, dass eine Verzehnfachung der Ethernet-Leistung nach einigen Ausschlägen recht zügig etwa zum dreifachen Preis der einfachen Lösung zu haben war und anschließend immer preiswerter wurde.

Beim Übergang zu 100 G haben wir aber leider eine unangenehme Besonderheit: rechnet man es auf die Jahre zurück, ist die Übertragungsrate in RZ-Netzen deutlich schneller gewachsen als Moore's Law. Dieses Gesetz sagt aber nur, dass sich die Anzahl der auf einer Fläche realisierbaren Transistoren alle 18 Monate verdoppelt. Es sagt aber nichts über die Geschwindigkeit, die letztlich vom Herstellungsprozess abhängt. Und in den letzten 30 Jahren hat der für praktisch alle Massen-Elektronik maßgebliche CMOS-VLSI-Herstellungsprozess die maximale Taktrate in einem Bereich von rund 3 GHz festgefroren. Das hat zur Konsequenz, dass praktisch alle Designs in die Breite gehen. Das kann man bei Prozessoren und Speichern schön sehen. Ein Prozessor hat immer mehr Cores (zu deren Beschäftigung wir die Virtualisierungssoftware benötigen), aber der einzelne Core ist nicht wesentlich schneller geworden.

Die Industrie hängt deshalb so liebevoll am VLSI-CMOS-Prozess, weil er zu den mit Abstand preiswertesten Lösungen führt. Für Kommunikationskomponenten bedeutet dies, dass ihre jeweiligen Aufgaben möglichst parallelisiert werden müssen, soweit es geht. Und genau dazu dient das Multi-Lane Konzept in den IEEE-Standards zu 10, 40 und 100 GbE. Ein 10 GbE-Signal wird innerhalb der Schaltung sehr schnell in vier parallele Lanes von je 2,5 GHz aufgespalten, die dann jeweils bequem mit CMOS-Taktrate bearbeitet werden können. Erst ganz am Ende werden die Signale der einzelnen Lanes wieder zusammengefügt. So können über 90% einer Schaltung in billiger CMOS-Logik ausgeführt werden, nur weniger als 10% müssen mit weit weniger stark integrierter, aber dafür erheblich schnellerer Technik wie GaAs gebaut werden.

Heutige 40 GbE-Lösungen sind wenn man es genau betrachtet nichts anderes als zusammengepackte 10 GbE-Lösungen, die man eben viermal nebeneinander gesetzt hat. Das hat auch in anderen Bereichen eine lange Tradition, IEEE 802.11n WLAN-Transceiver sind auch nur aus je vier 11a-Transceivern zusammengestellt und die neuen 11ac-Transceiver bestehen primär aus je vier 11n-Transceivern.

Die Zukunft privater Netze

vern und somit letztlich 16 11a-Transceivern. Das ist die Gewähr dafür, dass eine neue Generation prinzipiell genau die gleichen Schwächen hat wie eine ältere, sie werden systematisch vererbt.

Natürlich benötigt man bei dieser Zusammenfügerei hier und da auch mal neue Funktionen, das hält sich aber in Grenzen und es gibt praktisch für alles Schaltungs-tricks.

Bei 100 G hört der Spaß aber auf. Mit aufeinandergepappter 10 G-Technik wären 40 2,5 G-Lanes erforderlich. Aber selbst dann, wenn man das hinkommt, gibt es durch die pure brutale Geschwindigkeit der aufeinander folgenden Pakete erhebliche Schwierigkeiten. Es gibt ja schon seit Jahren 100G-Lösungen für Provider, die auch in den neuen Produkten Richtung 400 G pro Schnittstelle gehen. Die sind aber komplex, haben relativ große Formfaktoren und wären für die Anwendung in einem normalen RZ viel zu teuer. Glücklicherweise gibt es aber auch wieder neue Entwicklungen, die in den nächsten zwei Jahren dazu führen werden, dass auch die Betreiber normaler privater RZs sie wirtschaftlich einsetzen können.

Das Ganze ist für jemanden, der sich nicht dauerhaft intensiv mit der IC-Technik auseinandersetzt, kaum nachzuvollziehen. Deshalb greife ich hier zwei relativ simple Funktionen heraus, die sich im Rahmen herkömmlicher Technologie für 100 G nur noch extrem aufwändig implementieren lassen und für die man einfach völlig neue Ansätze benötigt. Um auch dies übersichtlich zu gestalten, greife ich nur einen einzigen der neuen Chips heraus, die Bandwidth Engine von MoSys.

2.1 Netzwerk-Speicher: Puffern von Datenströmen aus den Anwendungen

Um mit der wachsenden Zahl von Internet- und anderen Benutzern und deren wachsenden Ansprüchen klar zu kommen, müssen Netzwerk-Geräte sowohl für leitungsgebundene als auch für drahtlose Netze hinsichtlich der Leistung problemlos skalieren. Gleichzeitig benötigt man zur Unterstützung einer robusten Benutzer-Erfahrung die Freiheit von Delays und Down-Times sowie die Implementierung neuer Techniken wie Service Level Agreements, IPv6 und Intrusion Detection. Dadurch steigt nicht nur die fundamentale Paketrate deutlich, sondern unseligerweise auch die notwendige Anzahl von Lookups pro Paket zur Unterstützung der erweiterten Dienstleistungen. Die nächste Generation von Netzwerk-Geräten und Line Cards für Metro Ethernet, Carrier- und Core Router basiert vollständig auf ASICs und FPGAs, teilweise sogar mit der

Integration optischer Komponenten. Das ist ja bei den neuen Cisco Core Routern recht ausgeprägt. Mit einem gewissen Zeitversatz, erfahrungsgemäß ca. 2-3 Jahre, kommen diese Systeme auch in einen für „normale“ private Netze und RZs sinnvollen Rahmen hinsichtlich der Anschaffungs- und Betriebskosten.

In den nächsten Abschnitten betrachten wir die diskreten Funktionen und Lösungen, die für die Leistungsansprüche notwendig sind und gleichzeitig durch Network Operations, Administration und Management (OAM&P) so unterstützt werden können, dass die erforderliche QoS erreicht wird.

2.1.1 Das eigentliche Speicherproblem

Line Rates im Networking haben sich in den Jahren 1995 bis 2005 alle 18 Monate mehr als verdoppelt (ausgehend von 10 Mbps in 1994 hätten wir ab 2005 erst 1,28 GbE und nicht 10 GbE) und bei Fernstrecken ist dieser Trend ungebrochen, um den immer weiter steigenden Anforderungen gerecht werden zu können. Das war somit schneller als Moore's Law. Mit der reinen Geschwindigkeit nicht genug, sind im Laufe der Zeit viele neue Funktionen, Standards und Features primär zur Verbesserung der Qualität hinzugekommen. Beides zusammen erzeugt signifikante Herausforderungen an den Speicherdurchsatz und die Zugriffsrates. Auf FPGAs basierende Line Cards erlauben heute im Zusammenhang mit Hochleistungs-Speicherarchitekturen die Kombination von Bandbreite, Intelligenz und Flexibilität, die benötigt wird, um Pakete in Echtzeit zu bearbeiten und eine robuste Benutzererfahrung ohne Delay und Jitter zu ermöglichen.

Während normaler Datenverkehr eigentlich relativ unempfindlich ist, gibt es grade im Umfeld neuerer Anwendungen wie z.B. solchen zur Kollaborationszwecken einen nachvollziehbaren Mehrwert durch eine bessere Qualität im Netz. Benutzer akzeptieren solche Dienste umso besser, desto geringer die Störungen im Realzeitver-

halten sind. Das ist für einen Provider sehr wichtig, aber mit der Zunahme von BYOD- oder ähnlichen Anforderungen müssen sich auch normale Unternehmen mit derartigen Problemen auseinandersetzen. Dazu sprechen wir eigentlich seit Jahren über Konvergenz und in diesem Zusammenhang wird „Lossless“ Ethernet benötigt. Es geht also nicht mehr wie früher, dass im Fall einer kurzzeitigen Überlastsituation Pakete einfach verworfen werden. Solche Ereignisse führen zur Verschlechterung der Benutzererfahrung. Betrachtet man es fundamental, muss die Leistungssteigerung in Netzen mit deterministischer Paketverarbeitung in Echtzeit einher gehen. Die Zeiten der „Best Effort“ Services sind vorbei.

Vielleicht hierzu noch ein anschauliches Beispiel. In der Universitätsstadt Aachen war der Betreiber E-Plus besonders erfolgreich mit dem Verkauf von günstigen Flatrates, hatte aber sein Netz nicht entsprechend ausgebaut. Das führte zu einer erheblichen Verschlechterung der Übertragungsqualität bei Sprache, über Halbduplex bis hin zu permanenten Abbrüchen. Proteste waren zwecklos, also haben Bestandskunden gekündigt. Die kommen auch nicht mehr wieder. Ein schlechtes Netz zieht also ggf. unmittelbar wirtschaftliche Konsequenzen nach sich.

Innerhalb von Netzwerk-Geräten wird Speicher für drei fundamentale Operationen benötigt:

- Zwischenpuffern von Paketen
- Bearbeitung des Paket-Headers für Switching und Routing (Entscheidungsprozess)
- Aufzeichnung der Entscheidungen für Netzwerk-Management und Accounting

Die Leistungsanforderungen an Speicher in einer 100 GbE-Line Card sind erschreckend. In einer gewischten Umgebung werden ausschließlich Vollduplex-Leitungen unterstützt. Wir müssen uns also ansehen, was ein Speicher in einer solchen Line Card können muss, wenn letztlich auf

Anwendung	Dichte	Leistung	Traditionelle Lösung
Paket Puffer	20 Gb (mit ECC)	600 Gbps	20 x 1Gb DDR3 DRAM @ 1866
Statistik	0,5 Gb	2,4 G Zugriffe/sec	4 x 144M x18 QDR SRAM @ 600 MHz
Deskriptoren	5 Gb	660 M Zugriffe/sec	10 - RLDRAM oder 7+ 1 Gb DDR3 DRAM

Abbildung 2: FPGA-Anforderungen für 200 GbE Line Rates

Die Zukunft privater Netze

die Konstruktion 200 GbE zukommen. Die Paketankunftsrate von einem Paket pro 3,3 Nanosekunden schafft in Verbindung mit der Anforderung, die Pakete in Echtzeit zu bearbeiten, ein ganz neues Problemuniversum. Eine Implementierung mit traditioneller Technik hätte die in Abbildung 2 aufgezeichneten Rahmenbedingungen.

2.1.2 Puffer-Anwendungen

Puffern kann in einem Switch oder Router je nach seiner Architektur und den Leistungsanforderungen an bis zu drei unterschiedlichen Stellen stattfinden: Ingress-Ports, Egress-Ports und Packet Processor. So kann z.B. ein Ingress Port bei einem insgesamt mit Überbuchung ausgeführten System im Falle eines Bursts, der die Bearbeitungskapazität des Packet Prozessors überschreitet, die Spitze abfangen und Pakete zwischenspeichern. Ein Egress Port kann einen zu schnellen Datenfluss aus dem Packet Prozessor zwischenspeichern, wenn aktuell die Leistung des Port degradiert ist. Mittlerweile hat sich das Prinzip der speicherbasierenden Switch-ASICs für den Packet Processor durchgesetzt. Man muss aber hier zwischen dem Speicher, in dem das Switching bearbeitet wird und dem einem Paket Prozessor immer zugeordneten Puffer-Speicher unterscheiden. In Größe und Reaktionsgeschwindigkeit muss er natürlich genau auf den Prozessor abgestimmt sein. Normalerweise macht man diese Speicher aus preiswertem DRAM, vor allem deshalb, weil man die Kapazität auf diese Weise leicht und billig steigern kann. Untersuchungen zeigen jedoch, dass im Bereich von Hochgeschwindigkeitsanwendungen eine bessere Speicher-Leistung wichtiger als eine wachsende Speicher-Größe ist. Betrachtet man die Gesamt-Ökonomie auf der System-Ebene (Leistungsverbrauch, Wärme, Platzverbrauch) ist eine spezialisierte auf Leistung optimierte Lösung mit speziellen Komponenten letztlich günstiger.

Bei einer Übertragungsrate von 10 Gbit/s kann man z.B. einen Puffer für 100 Millisekunden an einem Ingress- oder Egress-Port mit zwei DDR2-800 DRAMs und einem 1 MB On-Die-Cache SRAM aufbauen. Wenn man darauf achtet, dass die I/O-Rate des DRAMs grob das Dreifache der Line Rate ist, skaliert diese Lösung auch für 100 GbE. Der Leser möge immer im Kopf behalten, dass wir über eine Umgebung sprechen, in der die Nominal-Datenrate ohnehin im Rahmen eines Multi-Lane-Konzepts mit entsprechender Parallelisierung für die individuellen, eben mehrfach parallel vorhandenen Komponenten umgesetzt wird. Insgesamt würde man für eine maximale Puffer-Zeit von 200 Millisekunden bei 100 Gbps für einen

Puffer 10 DDR3-2133 DRAMs und 10 Mb On-Die SRAM benötigen. In einem solchen Fall würden die DRAMs an ihrem Limit laufen, 10 Watt Leistung verbraten und knapp 10 Quadratzentimeter Platz auf einem Board benötigen.

Abgesehen davon, dass dies schon zu einem Gesamt-Design führt, welches in keinsten Weise den Verhältnissen in einem RZ Rechnung tragen würde, ist dieses Prinzip nicht beliebig skalierbar.

Vor allem wegen des in früheren Artikeln hinreichend beschriebenen Problems der Microbursts hat sich als konstruktive Daumenregel eingebürgert, einen Puffer so zu gestalten, dass er ein Paket etwa für die Dauer eines „Round-Trips“ zwischenspeichert. Mit der Zeit sind die Übertragungsraten gestiegen und somit die Round-Trip-Zeiten entsprechend gefallen und liegen jetzt im Markt bei rund 100 msec. Da die Datenraten jetzt die I/O-Raten der Speicherkomponenten erheblich übersteigen, ist eine immer größere Ansammlung paralleler Komponenten notwendig. Es gibt aber schon eine Reihe wissenschaftlicher Arbeiten, die zeigen, dass so große Puffer gar nicht erforderlich sind und negative Auswirkungen auf Jitter und Antwortzeiten haben. Bei Verkehrsströmen ohne Congestions werden die Puffer so gut wie nicht benutzt, aber im Rahmen von Burst Traffic können sie pathologische TCP Timeouts und Antwortzeiten in inakzeptablen Bereichen verursachen. Es gibt heute im Markt bereits latenzarme Switches, die vollständig verlustfreien Datenverkehr bei 10 GbE demonstrieren können, aber pro Port weniger als 40 msec. Puffer-Kapazität haben. Bei 40 GbE wird die Zeitanforderung unter 20 msec. und bei 100G unter 15ms fallen.

Außerdem gibt es mittlerweile unterschiedliche Optimierungsziele. Betrachtet man nur Puffer-Anwendungen, ist das

primäre Leistungsmaß der Durchsatz, betrachtet man aber die Bearbeitung von Paket-Headern, ist das kritische Leistungsmaß, möglichst schnell auf Tabellen und Entscheidungsbäume zugreifen zu können. Schon bei 10 Gbps übersteigt die notwendige Paket-Bearbeitungsrate die Random Access Rate eines generischen DRAMs. Ab diesem Punkt kann DRAM nur noch mit zunehmend komplexer werdenden Speichersteuerungen eingesetzt werden. Seither haben sich die Speicher auch weiterentwickelt, es gibt z.B. Low Latency DRAM oder QDR SRAM, aber nur mit weiterentwickelten Steuerungen können sie für höhere Übertragungsraten eingesetzt werden. Diese Lösungen haben heute bis 100 Gbps skaliert, aber nicht ohne erheblichen Aufwand und Kompromisse.

Um die hohen Line Rates handhaben zu können, benutzen Paket-Prozessoren lange, deterministische Pipelines. Das kann man z.B. bei den Intel Switch ASICs der Reihen 6000 und 7000 sehr schön sehen. Diese Pipelines entschärfen zwar die Anforderungen an Speicher mit ultra-geringer Latenz, aber nicht die Anforderungen an die Zugriffsraten für die Bearbeitung von Paket Headern in Echtzeit. Mit dreistelligen Gigabit Datenraten werden die notwendigen Bearbeitungsgeschwindigkeiten um den Faktor 15 schneller als die DRAM Zykluszeit. Schlaue Systemtricks oder spezielle Speicherprodukte haben immer weniger Wirkung und die Lücke in der Leistung wächst immer weiter.

Bevor die CPU-Architekturen von Single- zu Multi-Core-Systemen entwickelt wurden, waren Taktrate und Leistung eng miteinander verbunden. In ähnlicher Weise gibt es seit je her eine Beziehung zwischen Zugriffsraten und Latenz bei Speichern, wie es in Abbildung 3 zu sehen ist. Die Bandwidth Engine, die wir gleich näher kennenlernen, weist hier erheblich

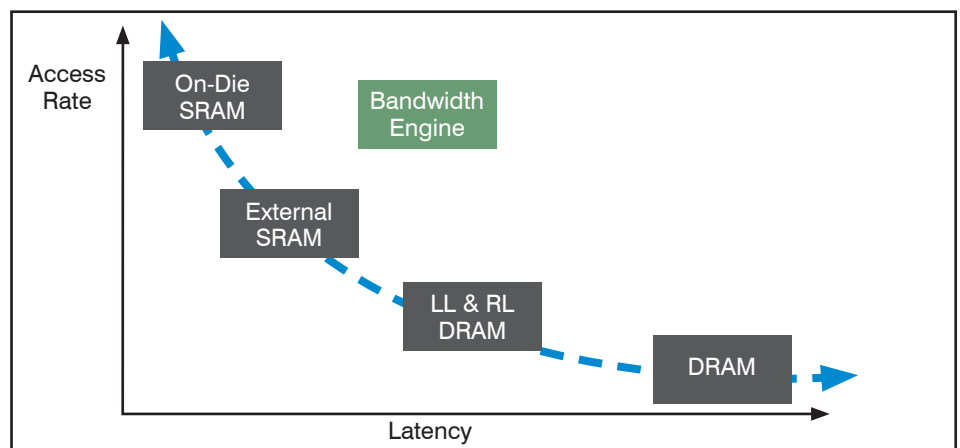


Abbildung 3: Relative Leistung traditioneller Speicher im Vergleich zur Bandwidth Engine

Die Zukunft privater Netze

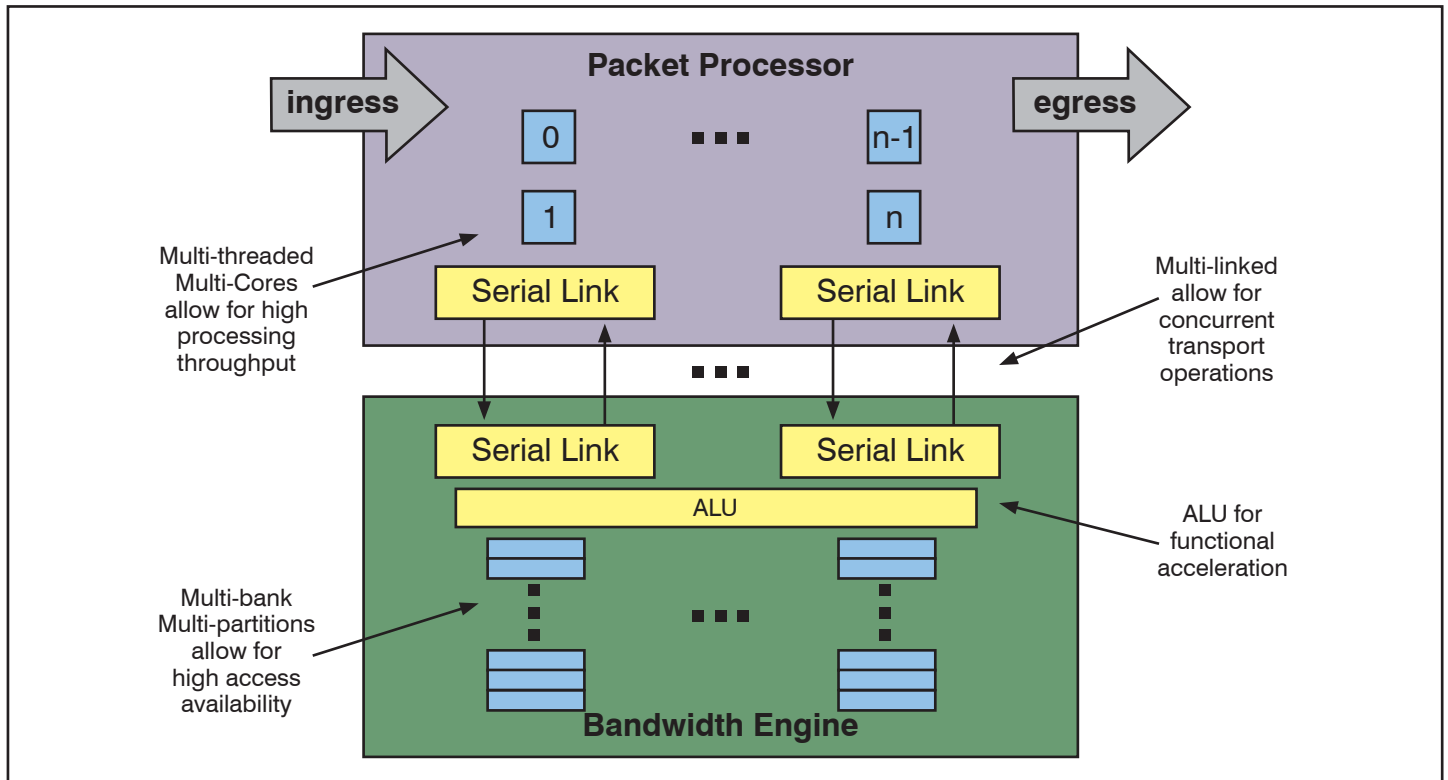


Abbildung 4: Bandwidth Engine und Anbindung an Packet Processor

Quelle: MoSys

günstigere Werte auf.

Da nun CPUs und auch Netzwerk-Prozessoren den Weg in Richtung Multi-Core eingeschlagen haben, ergibt sich ihr Gesamt-Durchsatz durch Aggregation der parallelen Komponenten. Der übliche RZ-Betreiber wird die Multi-Core Architekturen dadurch benutzen, dass er Virtualisierungstechniken einsetzt. Diese führen natürlich wieder zu einem gewissen Leistungsverlust, es gibt aber dazu kaum eine Alternative, wenn man Anwendungen versorgen und verwalten möchte, die mit der Leistung eines oder zweier Cores zufrieden sind. Es gibt allerdings auch Konstruktionen für die unmittelbare Nutzung der Parallelität der Cores zur Erzielung besonders hohen Durchsatzes. Das ist eher eine Konstruktion, die wir für die Anwendungen auf Netzwerk-Prozessoren im Bereich von 100 GbE oder höher benötigen.

2.2 Die Bandwidth Engine

Ein sehr interessanter Vertreter der neuen Netzwerk-Prozessor-Generation ist die sog. Bandwidth-Engine des Herstellers MoSys. Für den Anfang reicht es, sich vorzustellen, dass dieser Zusatz-Prozessor wesentliche Aufgaben bei der Bearbeitung von Paketen übernimmt. Die Architektur ist daher für hohe Leistung beim Speicherzugriff optimiert und implementiert ein Transportprotokoll mit 90% Effizienz auf einer seriellen CEI-11 und XFI-kompatiblen Schnittstelle. Anstatt sich

wie die traditionellen Speicheransätze nur auf Low Latency zu konzentrieren, ist die Bandwidth Engine Architektur hochgradig parallel und erlaubt gepipelineten, deterministischen sowie nebenläufigen Zugriff. Das ist ein wesentlicher konstruktiver Unterschied zu der Funktionsweise anderer Multi-Core-Netzwerk-Prozessoren, die lediglich mit Pipelines arbeiten, aber keine kontrollierte Nebenläufigkeit erlauben. (siehe Abbildung 4)

In Abbildung 5 sehen wir einen grundsätzlichen Überblick über die Konstruktion der Bandwidth Engine.

Möchte man den Unterschied zwischen einem konventionellen Netzwerk-Prozessor und der Bandwidth Engine auf den Punkt bringen, könnte man sagen, dass ein Netzwerk-Prozessor als grundsätzliche Von-Neumann-Maschine vielfach rekursiv arbeitet, was zu einer hohen Varianz der Ende-zu-Ende-Latenz führt. Die Bandwidth Engine ist sehr auf lineare Parallelität ausgerichtet und benutzt Rekursion nur im extremen Ausnahmefall. Dadurch wird sie so schnell.

Auch mit jahrzehntelanger Erfahrung ist es immer sehr schwierig, den Markt hinsichtlich wichtiger IC-Neuheiten zu beobachten. Viele Entwicklungen werden von Herstellern betrieben, die so winzig sind, dass man sie kaum findet. Sie sind viel flexibler als Giganten wie Intel, AMD, TI,

Broadcom oder Motorola. Außerdem gibt es ja im Bereich der Netze immer wieder auch Eigenentwicklungen bei den Herstellern, und sie sind mit gutem Grund längst nicht so freigiebig mit Informationen, wie man das als Beobachter gerne hätte.

Das Prinzip der Bandwidth Engine ist allerdings konstruktiv so elementar, dass ich mir nicht vorstellen kann, dass dies die einzige Entwicklung in dieser Richtung ist.

2.3 Funktionen mit besonderen Anforderungen

Viele Netzwerk-Anwendungen verwalten Zustände. Dazu gehören Anwendungen, die Policies durchsetzen, Network Address Translations-Anwendungen, Stateful Firewalling, TCP-Empfang, Netzwerk-basierte Erkennung von Anwendungen, Server Load Balancing, URL-Switching usw. um nur einige zu nennen. Die Verwaltung eines Zustandes ist eine speicherintensive Lese-Modifiziere-Schreibe-Operation (RMW-Operation), die eine geringe deterministische Latenz und komplett freizügigen Zugriff benötigt. Die komplette RMW-Operation muss innerhalb der Paket-Ankunfts-Zeit (3,2 Nanosekunden bei 100 GbE) erledigt sein und die Teile der Operation beziehen sich immer auf die gleiche Speicherstelle. Traditionell wurde hier QDR SRAM wegen seiner voneinander unabhängigen Schreib- und Leseope-

Die Zukunft privater Netze

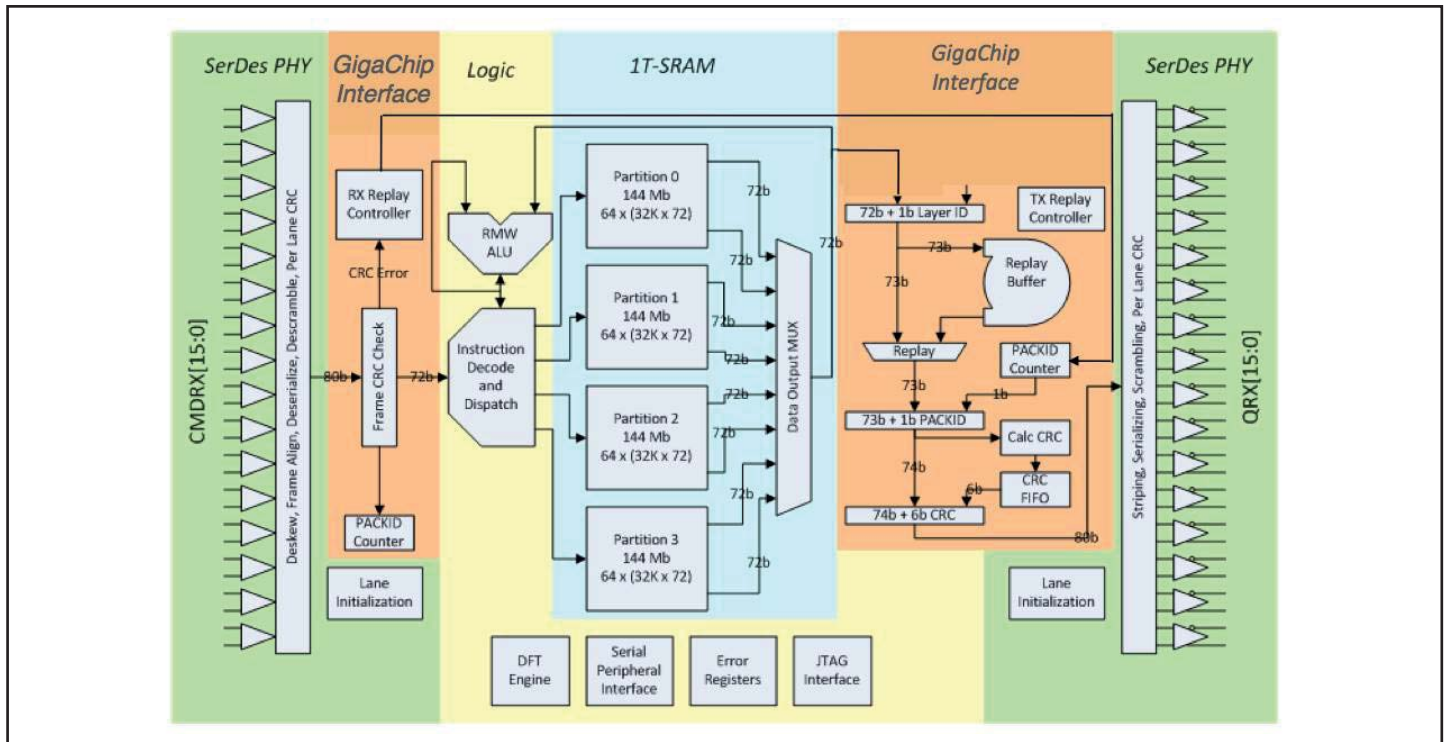


Abbildung 5: Grundsätzliche Konstruktion der Bandwidth Engine

Quelle: MoSys

rationen benutzt. Die Prozessor-Elemente für die Paketbearbeitung haben üblicherweise lange Pipelines und können keine Unordnung durch nicht-deterministische Antwortzeiten vertragen. Es gibt auch Tricks und Kniffe, wie man die benötigten Funktionen mit anderen Speichertypen implementieren kann, aber die haben ihre engen Grenzen und außerdem werden auch immer sehr viele Komponenten benötigt.

Die Implementierung der genannten Funktionen durch eine Bandwidth Engine benutzt je einen Schreib- und einen Lese-Port und erlaubt so das kontrollierte simultane Lesen und Schreiben auf jede Speicherstelle des verwalteten Bereichs. So kann die Kohärenz unter allen Zugriffsmustern gewährleistet werden. Der Effekt ist wie folgt: für eine 100 GbE-Schnittstelle würde man für die genannten Funktionen 70 DRAM-Bausteine mit mehr als 2000 Pins und einem Leistungsbedarf von 70 W benötigen. Die Bandwidth Engine kommt mit einer Komponente mit 64 Pins aus und verbraucht 9 W. Sie ist dabei um einen Faktor 30 schneller.

Andere Anwendungen, die in die Kategorie „Packet Header Processing“ fallen, sind Queuing/Scheduling, Statistik, Zähler und Traffic Management, die alle von der hohen Zugriffsgeschwindigkeit und der deterministischen Latenz profitieren.

Eine andere Anwendung, die in gewis-

ser Weise orthogonal zum Header Processing ist, ist Deep Packet Inspection DPI. DPI ist eine speicherintensive rekursive Operation, die eher die Nutzlast eines Datenpakets in Augenschein nimmt als den Paketkopf. DPI und anschließendes Filtern erlaubt erweitertes Netzwerk-Management, verbesserte Versorgung der Benutzer und natürlich Sicherheitsfunktionen. Leider gibt es keine praktisch nutzbare technische Möglichkeit, diese Funktion auf dem gesamten Netzwerkverkehr bei den hohen Datenraten anzuwenden. Man wird DPI an den Stellen des Netzes etablieren, wo die Inspektion logistisch machbar ist. Dennoch kann DPI von den höheren Zugriffsraten und den geringeren Verzögerungszeiten intelligenter Optionen für den Umgang mit Speichern erheblich profitieren.

Trotz der geringen Materialkosten ist die Anhäufung relativ dummer DRAMs keine wirkliche Lösung. Komponenten der Zukunft müssen über hoch effiziente serielle Schnittstellen und deterministische flexible Speicher-Arrays verfügen, die mittels einer entsprechenden Steuerung intelligent eingesetzt werden können.

Das Problem wurde hier für Switches dargestellt, ist aber nicht auf diese beschränkt. Auch die Adapterkarten benötigen Puffer. In der Vergangenheit wurde gerne der Weg gewählt, die Adapterkarte möglichst schlank zu designen und wesentliche Teile der Last auf den Prozessor

zu verschieben. Das kann man so lange machen, wie der Prozessor billig und gelangweilt ist. Bei Rechenzentren bilden aber Server den größten hardware-orientierten Kostenfaktor. Dynamische Virtualisierungskonzepte dienen hauptsächlich dazu, die Auslastung von Servern zu optimieren. In solchen Umgebungen schaut man primär darauf, wie man einen Prozessor entlasten kann, damit er seine eigentliche Arbeit besser erledigen kann. Ich verweise dazu auch auf meinen Artikel zur leistungsorientierten VM-Migration im Insider 12/13. Für den Betrieb ist es von elementarer Bedeutung, ob eine Kommunikationslösung einen Prozessor zu 70 oder 80% belastet oder nur zu 10%. Jenseits der 100 GbE wird es auch problematisch, die Kommunikation in einem Server so zu organisieren, dass die Übernahme von Funktionen durch den Prozessor noch mit der allgemeinen Paketrate Schritt halten kann und nicht im schlimmsten Fall sogar zu Rückstaus auf das Netz führt. Offload ist hier das Gebot der Stunde. Also werden auch hier entsprechende intelligente Komponenten benötigt.

2.4 400 G App-Beschleunigung und Host-Offload

Nach diesen vielen Vorbereitungen können wir uns jetzt ansehen, was passiert, wenn man den Speicher-Einrichtungen einer erhöhte Intelligenz mitgibt, um den Host von speicherintensiven Operationen zu entlasten. Es gibt ja hier zwei grund-

Die Zukunft privater Netze

sätzliche Fälle: die Entlastung eines Server-Prozessors durch seinen intelligenten Netzwerk-Adapter (das wäre die kleinere Lösung) und die Entlastung eines Steuer-Prozessors in einem Switch (das ist erheblich mehr Aufwand). Da man sich das auf dem Adapter eigentlich leicht vorstellen kann, betrachten wir hier nur die letzte Alternative. Dabei gehen wir für diesen Artikel von der Unterstützung eines Gesamtstromes von 400G in Form von 4 X 100G aus. Die Darstellungen sind aber durchaus leicht auf andere Konfigurationen wie 12 X 40G oder 48 X 10 G übertragbar. Wir sprechen hier also nicht über futuristische aggregate Leistungsbereiche, sondern über solche, die von Switches praktisch aller Hersteller heute erreicht bzw. übertroffen werden. Die intelligente Speichereinheit spielt die Rolle eines kompletten Subsystems, welches in der Lage ist, Instruktionen auszuführen und anzuhalten und dabei Datenaufzeichnungen zu machen, die den Qualitätsanforderungen in Providernetzen entsprechen.

Die Anwendungen für eine intelligente Speichereinheit gehen weit über die Übertragung und allgemeine Verarbeitung von Netzwerk-Verkehr hinaus. Netzwerke sammeln Informationen über Leistung und Nutzung. Diese Informationen können für Operation, Administration und Management (OAM) von Netzwerk-Dienstleistungen und Geräten benutzt werden. OAM ist letztlich die Heimat einer breiten Klasse von Dienstleistungen, die sich auf Überwachung und Verwaltung von QoS, die Sicherung der Compliance mit SLAs beziehen und für Abrechnungszwecke benutzt werden können.

Die Leser von Publikationen des ComConsult-Universums befassen sich ja meist weniger mit der inneren Struktur von Providernetzen. Deshalb sei an dieser Stelle angemerkt, dass die OAM-Funktionen erweitert um das Provisioning (OAM & P) beginnend mit SONET eine fast 30-jährige Geschichte haben, in vielen Standards verankert sind und sich auch permanent entlang der technischen Möglichkeiten weiterentwickeln, so z.B. für die Carrier Class Ethernet. Für Provider ist es untragbar, sich an diesen wichtigen Stellen mit proprietären Sonderlocken auseinanderzusetzen zu müssen. Auch wenn Cisco in diesen Bereich viel Equipment verkauft, käme der Hersteller hier nie auf den Gedanken, zu Gunsten eigener Ideen und Programmpakete auszuscheren. Auch wenn die Betreiber privater Netze es nicht wahrhaben wollen: mit zunehmenden Anforderungen an Netzwerk-Leistungen und deren Qualität kommen wir dem Bereich, an dem OAM&P-Funktionen wirklich lebenswichtig und elementar

sind, immer näher. Und es bedarf keiner besonderen prophetischen Gabe, dass die Anforderungen in Zukunft immer weiter steigen werden.

2.4.1 Advanced Network Management

Funktionen für Statistik, Messungen und Traffic Shaping werden traditionell mit Carrier Class Ethernet assoziiert, welches sich vom gewöhnlichen LAN-Ethernet durch Standardisierte Dienstleistungen, Skalierbarkeit, Zuverlässigkeit und Service-Management unterscheidet. Diese Attribute ermöglichen die Evolution des traditionellen Ethernet zu einer Technologie, die auch für die Belange von Metro- und WAN-Providern geeignet ist.

Bislang konnten sich Betreiber privater Netze (und auch deren Lieferanten) damit herausreden, dass diese Art von Qualität nicht benötigt wird. Mit der Etablierung von RZ- und Cloud-Services benötigen die zugrunde liegenden Netze jedoch zunehmend Carrier Grade Qualität, um eine robuste Benutzererfahrung zu ermöglichen.

Das wird besonders an der aktuellen BYOD-Diskussion klar. In vielen Fällen wird es darauf hinauslaufen, dass die mobilen Benutzer innerhalb der Reichweite eines Unternehmensnetzes z.B. mittels einer WLAN-Infrastruktur durch Leistungen des privaten RZs versorgt werden,

während sie nach Verlassen des unmittelbaren Einzugsbereichs die Leistungen von einem Provider erhalten, z.B. über eine LTE-Anbindung. Damit es keinen logischen Bruch gibt, wird eine hybride Cloud-Lösung erzeugt. Nun hat der Provider auf jeden Fall ein Carrier Class Netz. Das Unternehmen sieht schlecht aus, wenn die lokale Versorgung nicht wenigstens die gleiche Robustheit und Qualität hat. Das private Netz muss in jedem Fall „Carrier Class“-Eigenschaften und -Qualität besitzen.

Jedes Provider-Netz hat eine tief integrierte Multi-Mandantenfähigkeit, man könnte sie als elementares Arbeitsprinzip betrachten. Ein Carrier Class Ethernet liefert den Verkehr nicht wild an alle, die grade zuhören, sondern an spezielle Lokationen. Jede Verkehrsart repräsentiert dabei eine eigene Dienstklasse. Jede Dienstklasse benötigt ihre eigenen Qualitätsgarantien. Sie besitzt eine eigene Markierung, damit sie während des Transports immer identifiziert werden kann um das Ziel, sie mit genau für sie zugelassenen Service-Levels und einer präzisen Leistung zu unterstützen, erreichen zu können. Um wirklich auch unter komplexeren Bedingungen, wie z.B. sich dynamisch ändernden Anforderungen aus den Anwendungen den Service Level auch immer gewährleisten zu können, benötigen

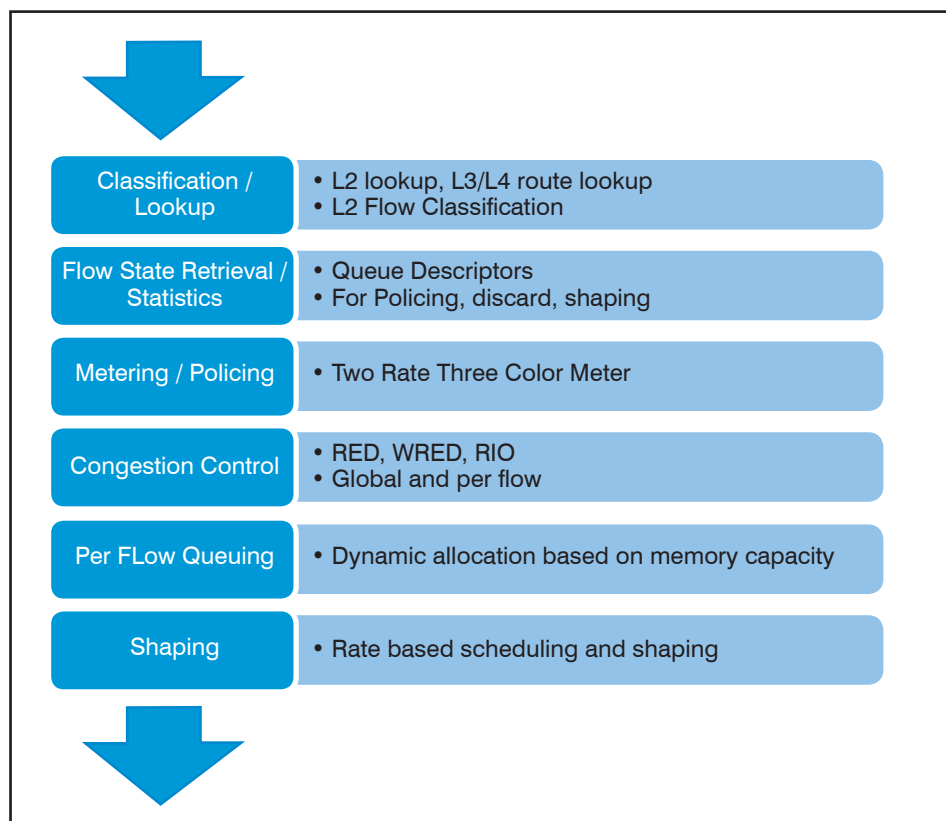


Abbildung 6: Bearbeitungsstufen für Traffic Management

Die Zukunft privater Netze

Operateure immer mehr granulare und an Real-Zeit orientierte Maße wie Frame Delay, Jitter und Frame Loss. Die Abbildung 6 zeigt, welche Bearbeitungsstufen alle notwendig sind, um den Paketfluss durch ein Carrier Class Netz angemessen zu gewährleisten.

Ob nun diese Funktionen als Background Prozesse implementiert werden oder Bestandteil des aktiven Managements des Netzwerk-Verkehrs sind: die tatsächliche Datensammlung, wie Packet Header Processing, muss mit den aggregierten Bearbeitungsraten geschehen. In diesem Artikel betrachten wir zwei Beispiele: den Offload von Statistik-Zählern und eine two rate/three-colour token bucket Implementierung, die für Accounting, Messungen und Trimmung des Netzwerk-Verkehrs benutzt werden können.

2.4.2 Anwendung: Statistik-Zähler

Weil ein Paket durch viele Entscheidungsprozesse läuft, ist es üblich, festzuhalten, welche Aktionen vorgenommen wurden und die Netzwerk-Leistung zu überwachen. Das kann mehrere Zähler pro Line-Interface oder pro Flow im Falle einer Flow-basierten Architektur erforderlich machen. Zähler können mit einem weiten Bereich zusätzlicher Fähigkeiten implementiert werden, wenn die Anwendungs-Umgebung das verlangt. Eine elementare Zählerinformation kann z.B. Bits oder Pakete pro Flow sein. Komplexere Systeme, wie z.B. die Hardware eines Hochgeschwindigkeits-Switches für Carrier mit intelligenterem Flow-Management, können durchaus acht oder mehr Zähler pro Flow besitzen. Komplexe Firewall-Systeme haben heute 26 oder mehr Zähler pro Flow. Counter werden durch eine Read-Modify-Write Speicher-Operation implementiert, die dadurch beschleunigt werden kann, dass man Speicher mit geringer Latenz und schnellerer Zykluszeit oder Dual-Port Speicher verwendet. Die Abbildung 7 zeigt das grundsätzliche Prinzip.

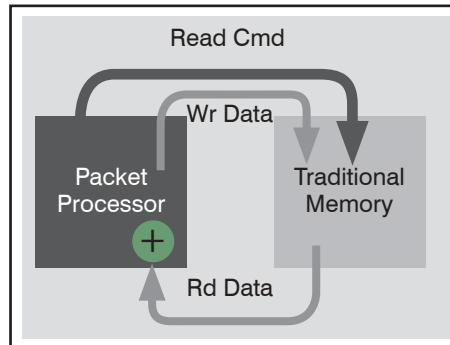


Abbildung 7: Implementierung eines Counters mit traditioneller Speichertechnik

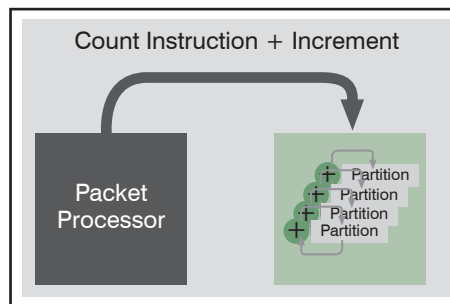


Abbildung 8: Counter Offload Engine

heitswert für jedes Paket

- Additions- und Subtraktions-Operationen bei der Paket-Länge

Betrachtet man die ersten drei Kriterien, kommt man schnell darauf, dass man zu ihrer Implementierung eine hybride DRAM/SRAM-Lösung benötigt, wie sie schon weiter oben bei der Pufferung beschrieben wurde. Ein anderes Konzept ist es, die Sache dadurch wirtschaftlicher zu gestalten, dass man die Adressen in kleinere Speicherbereiche hasht. Derartige Kompressionsstrategien ignorieren generell Kollisionen und sind nur für Zähler brauchbar, die auch mit ungefähren Werten auskommen. Sie sollen hier nicht weiter betrachtet werden. Die Leistung hyb-

riden Zähler-Implementierungen basiert auf dem Verhältnis zwischen SRAM und DRAM Zyklus-Zeit, der Anzahl der DRAM-Bänke, der Cache-Größe und der Queue-Länge. Z.B. würde man für 4 X 100 GbE 144 Mb SRAM und 20 GB/s DRAM-Bandbreite benötigen. Auch mit QDR SRAM ergeben sich hinsichtlich der Wirtschaftlichkeit völlig unbefriedigende Dimensionen.

Machen wir das doch besser mit einem Netzwerk-Prozessor, z.B. der Bandwidth Engine von MoSys. Mit ihren Onboard-Beschleunigern ist sie in der Lage, atomare Counter-Update Operationen, die die Aufzeichnungen komplett intern inkrementieren können, und reduziert damit die Anzahl der notwendigen Speicherbus-transaktionen von sechs auf eine einzige. Die ganze Operation kann in weniger als 20 Nanosekunden durchgeführt werden und stellt damit alle anderen Lösungen auf Basis von Speichern weit in den Schatten. Die Bandwidth Engine entlastet den Packet Prozessor wirkungsvoll von Zähler- und ECC-Operationen und sorgt für eine kohärente Überwachung von „In-Flight“-Transaktionen innerhalb der Counter Pipeline. Durch den kompletten Offload hängt die mögliche Counter Rate letztlich nur noch von der internen Zugriffsrates des parallelen Speicher-Arrays ab. Die Abbildung 8 zeigt das grundsätzliche Prinzip einer Counter Offload Engine.

2.4.3 Anwendung:

Two Rate/Tree Colour Token Bucket

Um QoS- und SLA Anforderungen gewährleisten zu können, muss ein Netz Verkehr mit unterschiedlichen Prioritätsstufen und unterschiedlichen Flows implementieren und dabei Latenz, Jitter und Paket-Übertragungsleistung für jeden Flow realisieren. Das erfordert Geräte im Netz, die grundsätzlich eine Funktion haben, wie in Abbildung 9 dargestellt, nämlich ein „Messinstrument“ für die Unterstützung von Policies sowie einen zwei-

Die traditionelle direkte Implementierung mit High Speed SRAM wird bei höheren Line Rates und einer steigenden Anzahl von Flows, die überwacht werden müssen, wirtschaftlich uninteressant. Es gibt verschiedene Vorschläge, diese Funktion mit DRAM zu implementieren, aber derartige Vorschläge schaffen es häufig nicht, die folgenden Leistungsziele zu erreichen:

- Zählen bei jedem Flow statt Sampling von Stichproben
- Kontinuierliches Zählen, welches genügend Bandbreite benötigt, damit man lesen kann, während man zählt
- Zählen von Paketen und deren Längen
- Paket-Zähler-Inkmente in einem Ein-

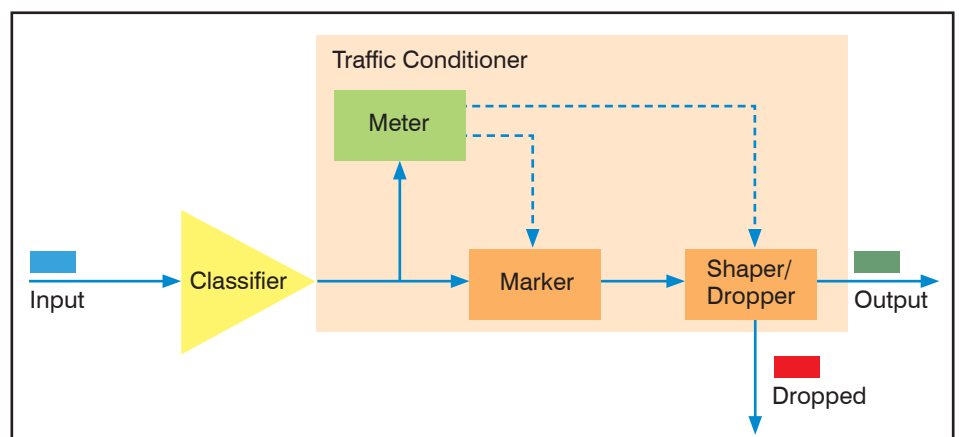


Abbildung 9: Traffic Conditioner

Die Zukunft privater Netze

stufigen Queuing-Mechanismus, der den Verkehrsfluss in Form hält und für vorhersehbare Leistung sorgt und Fairness im Scheduling mit einer besseren Lastverteilung im Netz unterstützt.

Kommt ein Paket an einem Eintrittspunkt des Netzes an, wird es gemäß seiner Verkehrsklasse eingeordnet, die ab jetzt dafür maßgeblich ist, in welcher Weise die Flows des Paketes behandelt werden. Man kann Traffic Shaping auch dazu benutzen, um Bursts zu eliminieren oder Puffer-Überläufe zu vermeiden. Shaping kann durch einen Token Bucket Algorithmus für ein Bandbreitenprofil implementiert werden. Pakete, die den Eingangspuffer einer in Kommunikationsrichtung vorne (downstream) liegenden Einrichtung zu überlasten drohen, werden auf der Senderseite verzögert, bis der Empfangspuffer wieder genügend Platz hat. Alleinige Messungen an den Eintrittspunkten des Netzes sind unzureichend, um Congestion im Netz zu vermeiden. Wenn ein Paket das Netz betritt, wird es mit Paketen aus anderen Flows zu einem schnelleren konvergierten Datenstrom aggregiert. Es ist wichtig, dass auch auf diesen Links höherer Geschwindigkeit Policies implementiert sind, um Ende-zu-Ende-QoS zu realisieren und auf drakonische Maßnahmen üblicherweise verzichten zu können. Metering und Policing kann also an verschiedenen Knoten entlang des Weges eines Paketes vorgenommen werden, um sicherzustellen, dass Realzeit-Eigenschaften und Priorisierung vertragsgemäß eingehalten werden können. Ein Operator kann geforderte Bandbreite-Garantien dadurch einhalten, dass er die passenden Netzwerk-Ressourcen reserviert. Dabei kann er die Two Rate / Three Colour Methode zur Raten-Begrenzung als Teil des Traffic-Engineerings benutzen.

Das Messen von Verkehr mit einem Token-Bucket ist eine der schwierigsten Operationen, wenn es um hohe Line-Rates geht. Während man bei anderen Operationen, die wir in diesem Artikel angesprochen haben, immer noch irgendwelche architekturellen Tricks benutzen kann, um Funktionen trotz zu langsamer Prozessoren oder Speicher dennoch wenigstens ungefähr zu implementieren, ist die Metering-Funktion insofern einzigartig, als dass alle Aufzeichnungen in einer einzigen Lokation sitzen müssen und durch die Natur der High Speed Umgebung Vergleiche und Updatens nicht einfach gecached werden können.

Die Abbildung 10 zeigt ein Two Rate / Three Colour Meter (TRTCM), bei dem der erste Behälter die Commit Rate und der zweite Behälter die Excess Rate repräsentiert.

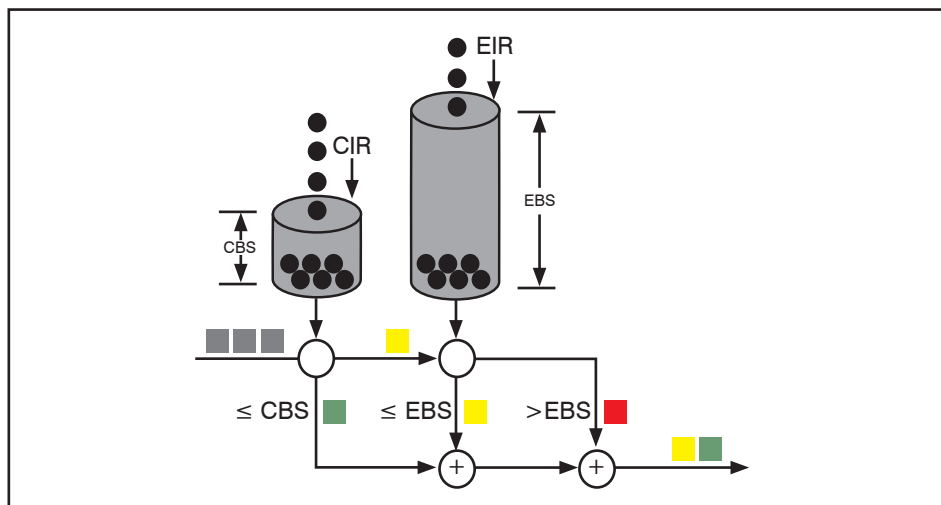


Abbildung 10: TRTCM

In die Behälter kommen die aktuellen Token für die Input Raten (CIR und EIR) und die Behälter dürfen bis zu den jeweiligen Bucket Size Limits (CBS und EBS) gefüllt werden. Jeder Flow hat ein eindeutiges Bandbreitenprofil, welches die Input Raten und Bucket Sizes beschreibt. Dieses Profil wird in den Metering-Aufzeichnungen niedergelegt, wenn der Flow etabliert wird. Jedes zu einem Flow gehörige Paket zieht zunächst Token vom Commit Bucket und wenn der leer ist, vom Excess Bucket ab. Das Paket bekommt die Farbe grün oder

gelb, je nachdem, welches Token Konto belastet wurde. Sind beide Konten leer, gibt es die rote Karte. Diese Funktion erfordert mehrere Operationen für jede Metering Instruktion in Line Rate. Nehmen wir z.B. eine 4X100G Line Card mit einer Paket-Ankunftsrate von 6,67 Nanosekunden und einer Metering-Record der Größe 128 Bits.

1. sobald ein Paket ankommt, wird die Flow-ID dazu benutzt, um die passende Metering Record aufzurufen

Kongress

ComConsult Netzwerk Forum 2014 24.03. - 26.03.14 in Königswinter

- Dienstreutralität kontra Applikations-Integration: Kampf um die Steuerung
- Netzwerk-Design in Zeiten von 10/40/100: wie sieht die Zukunft aus?
- Mobile Endgeräte im Netzwerk: teuer, unsicher, unbeherrschbar?

Das ComConsult Netzwerk Forum 2014 stellt sich den herausragenden Fragen der Entwicklung der Unternehmensnetzwerke. Nach 10 Jahren des Konzept-Stillstands ist das traditionelle Netzwerk unter Druck sowohl aus der Server- als auch aus der Endgeräte-Welt. Kombiniert man dies mit dem notwendigen Bandbreitenwechsel auf 10/40/100, dann steht das Unternehmensnetzwerk der Zukunft auf dem Prüfstand.

Die Fragen der Zukunft sind:

- Wie sieht ein Konzept zur schlüssigen Integration mobiler Endgeräte aus?
- Was bedeutet das für unsere Netzwerke?
- Wie sehen Gesamtarchitekturen von Endgerät bis zum App-Shop aus?
- Wie kann die notwendige Sicherheit wirtschaftlich umgesetzt werden?

Wir analysieren für Sie wohin die Netzwerk-Entwicklung geht, wer die Kontrahenten sind, welche Vor- und Nachteile die verschiedenen Alternativen für Ihr Unternehmen haben und wie relevant einzelne Technologien sind. Investieren Sie zukunftsicher in Ihre Netzwerke, das ComConsult Netzwerk Forum 2014 unterstützt Sie dabei.

Moderation: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr.-Ing. Behrooz Moayeri

Preise: € 2.390,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Die Zukunft privater Netze

2. das Meter berechnet die Token Credits für die Commit Rate und Input Rate Buckets basierend auf der Zeit seit dem letzten Update und der Bandbreiten-Zuordnung. Dieser Credit wird zu den aktuellen Bucket Werten hinzugerechnet
3. von jedem dieser Konten wird die Paketgröße abgezogen, wie das im Verfahren definiert wurde
4. ist die Zahl der Tokens nicht ausreichend, wird das Paket umgefärbt
5. die Meter-Funktion liefert die aktualisierte Information für Farbe und Bucket-Wert an den Marker und Dropper, die für Policing und/oder Traffic Shaping benutzt werden

Die Bandwidth Engine von MoSys hat einen Onboard-Beschleuniger für die genannten fünf atomaren Operationen. Die Tatsache, dass die Operationen alle an einer Stelle ausgeführt werden können, senkt die Anzahl der benötigten Speichergreife von sechs auf einen einzigen und befreit darüber hinaus den Host von den Rechenoperationen. Insgesamt kann die ganze Operation innerhalb von 30 Nanosekunden ausgeführt werden, das ist erheblich schneller als mit jeder anderen bekannten konstruktiven Alternative.

2.4.4 Zwischenfazit

Wenn die Netzwerk-Leistung schneller als Moore's Law wächst, benötigt man architekturelle Verbesserungen, um Schritt halten zu können. Netzwerk-Einrichtungen sind zu hochparallelen, multi-Thread System on Chip-Komplexen geworden, die einen quasi unbegrenzten Hunger auf Speicher haben. Konventionell gestaltete Speichersysteme kommen hier eben so wenig mit wie konservative Netzwerk-Prozessoren. Gefragt sind neue Konstruktionen, wie z.B. die Bandwidth Engine von MoSys, die Angebote für spezielle Problemlösungen darstellen. Aktuell kann der Chip bis hin zu Geschwindigkeiten von 384 Gbit/s. zwischenpuffern, bis zu 4,5 Milliarden Dual-Port-Zugriffe pro Sekunde verarbeiten und zusätzlich Statistik und TRTCM-Funktionen offloaden. Die Funktionen werden über die offene GiGi-Chip-Schnittstelle angeboten, das ist ein spezielles Transportprotokoll für den Informationstransfer zwischen Chips mit einer Effizienz von 90% und vollständig CRC-geschützter Kommunikation. Nur mit derartigen Innovationen kann man die nächste Generation von Netzwerk-Geräten mit erheblich gesteigerter Leistung, geringerem Stromverbrauch, geringerem Platzverbrauch und geringeren Kosten als die aktuellen Vorgängergenerationen auf den Markt bringen.

3. Konsequenzen für die Betreiber privater Netze

Wenn man es zusammenfassend betrachtet, ergibt sich folgendes Bild für die Kommunikations-Infrastruktur in einem RZ, aber in letzter Konsequenz auch für ein Campus-Netz:

- Die Implementierung von Netz-nahen Zusatzfunktionen mit Ausnahme solcher, die RegEx-Verarbeitung haben, kann mit Software auf der Basis von Virtualisierungslösungen erfolgen.
- Die Implementierung von Edge-Switches kann und wird mit Software auf der Basis von Virtualisierungslösungen erfolgen.
- Durch zunehmende Anforderungen aus Anwendungsbereichen wie Kollaboration werden die Ende-zu-Ende-Qualitätsanforderungen an Netze weiter steigen. Man kann es drehen und wenden wie man will: dadurch, dass das Netz Aufgaben übernimmt, die normalerweise ein Service-Provider hat, wird es in letzter Konsequenz zum Provider-Netz und muss entsprechend gestaltet und betrieben werden.
- Die Virtualisierungslösungen an den Netzwerk-Kanten hilft enorm bei der Steigerung von Skalierbarkeit und Zuverlässigkeit eines Netzes und unter Voraussetzung einer dynamisierten Umgebung auch unter weiteren betrieblichen Aspekten.
- Durch die Qualitätssteigerungen wachsen die Anforderungen an die nunmehr noch verbliebenen Komponenten des Kern-Netzes erheblich, nicht nur hinsichtlich der zu realisierenden Übertragungsleistung, sondern auch hinsichtlich der Möglichkeit der Verarbeitung von Qualitätsparametern durch Paket- und Flow-Manipulation bei schwindelerregenden Paket-Ankunftsdaten im einstelligen Nanosekunden-Bereich.

Das ist eine Vision für die mögliche Weiterentwicklung der Netze. Sie wird sich im Laufe der nächsten zwei Jahre materialisieren und ist nicht aus dem Kaffeesatz gelesen, sondern durch systematische Übertragung der real bereits vorliegenden Technologien für Providernetze auf privat betriebene Umgebungen. Auch wenn die Produkte für virtuelle Edge Switches heute noch so zahlreich sind wie die Top-Models, die das Wohnzimmer des Autors täglich durchqueren, werden sie schnell auf den Markt kommen, denn sie stellen auch für kleinere Anbieter eine enorme Chance dar.

Leider muss man konstatieren, dass die heute erhältliche Hardware der Vision nicht genügen kann. Auch das sieht in zwei Jahren sicher anders aus, dann wird 100 GbE in Carrier Class Quality mit hoher Wahrscheinlichkeit erschwinglich sein.

Die Vision klärt letztlich auch die Rolle von SDN. Es wird Software kommen, die Netze in ausgesprochen hohem Umfang steuert und kontrolliert. Diese Software bietet dann auch die Schnittstelle für höherwertige Funktionen bei Netzwerk-Management und Orchestrierung. Die Gründungsmitglieder der ONF sind Provider. Alle bislang vorgestellten nützlichen SDN-Lösungen kamen aus dem Provider-Umfeld oder von Universitäten. Der erfolgreiche Einsatz von SDN setzt heute noch hoch geschultes Personal in praktisch beliebiger Menge voraus, wie es Provider wie Google eben haben. Die Standardisierung ist heute nur marginal. Erst wenn sie ein operativ interessantes Volumen erreicht hat, wird das auch für Unternehmensnetze interessant. Nach Einschätzung von Infonetics wird das erst 2017 der Fall sein. Das könnte stimmen. Insgesamt ebnet die hier genannte Gesamtkonstruktion natürlich mittel- und langfristig den Weg für SDN.

Es ist immer schön, auch vollständigen Unsinn aus der zukünftigen Diskussion endgültig entfernen zu können:

- die Vorstellung, durch den Einsatz von möglichst viel Software und SDN letztlich billige Standard-Switches verwenden zu können, ist schlichter Unsinn. Die in dieser Hinsicht erfolgreichen White Box Experimente haben in einem für normale Unternehmen und Organisationen irrealen Umfeld statt gefunden und sind nicht übertragbar. Anspruchsvolle Software-Lösungen verlangen auch anspruchsvolle Hardware.
- die Vorstellung von VMware, Netze einfach mit einer Virtualisierungsschicht zu verdecken und zu einem Tunnelsystem zu machen, ist gefährlicher Unsinn. Hier werden viele Betreiber verunsichert, ohne dass es dafür eine sinnvolle Basis gäbe. Bei gesteigerter Qualität sollte das Tunnelsystem ja in der Lage sein, auf den individuellen Verbindungen QoS zu realisieren. Das gibt es ohne Zweifel, ein Tunnelsystem, was das kann, ist Carrier Ethernet. Normale Ethernet-Switches sind dazu aber nicht ohne Weiteres in der Lage, vor allem, wenn sie älter sind. Die Virtualisierungsschicht kann aber dem technischen Netz zu einem späteren Zeitpunkt keine Funktionen mehr hinzufügen, die es von sich aus nicht hat. Also um es kurz

Die Zukunft privater Netze

zu machen: VMware soll weiter Virtualisierungssoftware schreiben und Netze denjenigen überlassen, die seit drei Jahrzehnten davon etwas verstehen!

Bis alle interessanten Komponenten verfügbar sind, kann es durchaus noch zwei Jahre dauern. Was sollte der Betreiber bis dahin machen, um den kommenden Umbruch angemessen vorzubereiten? Das ist pauschal wegen der unterschiedlichen Ausgangslagen schwierig zu sagen, aber es gibt wenigstens zwei Richtlinien, die man festhalten könnte.

In der Vergangenheit hat es sich bewährt, das Netz und die Geräte, die es benutzen, sauber zu trennen. Auch für die Zukunft gibt es keinen Grund, das in übertragenem Sinne weiter zu verfolgen. Nur die Grenze verschiebt sich. Werden Rand-Switches virtualisiert, können sie physikalisch direkt zu den Servern gehören, die sie auch benutzen. Nur die Kern-Switches haben dann noch eine eigenständige physikalische Existenz. Die virtuellen Rand-Switches sollten aber in logischer Hinsicht zur „Netz“-Verwaltung gehören. Experimente mit zum Netzwerk assoziierten Funktionen werden hier für den Betrieb mehr Klarheit bringen.

Ebenfalls in der Vergangenheit bewährt hat sich die Orientierung an Standards. Es gibt wenn man vom Eigeninteresse der Hersteller einmal absieht keinen nachvollziehbaren Grund für proprietäre Lösungen. Ich meine immer, dass die Hersteller seit Jahrzehnten versuchen, die Betreiber privater Netze zu vereiern (ich meine ein anderes Wort, aber das drückt ja niemand). Das Schlimme ist, dass sie damit immer weiter Erfolg haben. So haben wir seit Langem die perverse Situation, dass der gleiche Hersteller Providern Geräte liefert, bei denen noch nicht einmal ein einzelnes Bit zuckt, ohne dass das durch einen Standard festgelegt wurde und bei privaten Betreibern glaubt, seine proprietären Steuerungsstränge ausleben zu können.

Der passende Standard für das Kern-Netz wäre Carrier Ethernet (dann hören auch endlich diese PLSB- und TRILL-Diskussionen auf). Und vielleicht erbarmt sich ja mal ein Hersteller, eine Basisversion des G.709 OAM&P-Systems, mit dem man heute nur so eine Kleinigkeit wie das Internet erfolgreich betreiben kann, auch für Zwecke privater Netze zuzuschneiden. Sonst gibt es nämlich schon wieder statt Lösungen endlose Diskussionen über Management, Automatisierung und Orchestrierung. Eigentlich wünsche ich mir das ja schon seit vielen Jahren vergeblich, aber ich sehe, dass der Wind sich dreht. Denn

die Betreiber neigen endlich zunehmend dazu, die Hersteller auf allzu komplexen und proprietären Strukturen schlicht sitzen zu lassen. So entsteht im Zusam-

menhang mit den beschriebenen Neuentwicklungen Druck, den es vorher so für keinen Hersteller gab. Und das ist wirklich spannend!

Kongress

ComConsult Netzwerk Forum 2014 24.03. - 26.03.14 in Königswinter

- Dienstneutralität kontra Applikations-Integration: Kampf um die Steuerung
- Netzwerk-Design in Zeiten von 10/40/100: wie sieht die Zukunft aus?
- Mobile Endgeräte im Netzwerk: teuer, unsicher, unbeherrschbar?

Traditionell ist die Dienstneutralität das Hauptmerkmal aller Netzwerke. Alle bisherigen Versuche, eine Applikations-spezifische Ende-zu-Ende Kontrolle und Garantie einzuführen sind in den letzten 20 Jahren gescheitert. Auch Prioritäten und Bandbreitenreservierungen für Verkehrsklassen sind trotz diverser Standards umstritten. Nun haben sich Netzwerke auf der Server-Seite in den letzten 5 Jahren deutlich verändert. Hier dominieren inzwischen die Software-Ports in den virtuellen Switches gegenüber den Hardware-Ports. Aus Sicht der Applikationen und virtuellen Maschinen spielt sich das Netzwerk mehr und mehr in Software ab. Verbindet man das mit dem zunehmenden Bedarf nach „One-Touch-Installationen“ und der schnellen Inbetriebnahme auch komplexer Applikationen innerhalb von Minuten, dann hat sich in der Server-Welt der Netzwerk-Schwerpunkt von der Hardware in die Software verschoben.

Dies wird eine Reihe von Fragen auf:

- Liegt die Zukunft der Server-Vernetzung in der Software?
- Hat das dienstneutrale Netzwerk damit ausgedient?
- Wie kann eine automatische Netzwerk-Konfiguration bei der Installation von Servern erreicht werden?
- Wie kann der Bedarf dynamischer virtueller Umgebungen am besten erfüllt werden?
- Wer wird die Zukunft bestimmen: Anbieter wie VMware mit NSX oder die traditionellen Netzwerker mit Application Aware Networks?

Der Endgerätebereich wird sich in den meisten Unternehmen in den nächsten 5 Jahren deutlich verändern. Mobile Endgeräte und neue Endgerätevarianten werden die Zukunft bestimmen. Das führt nicht nur zur Frage des besten physikalischen Anschlusses, sondern vor allem zur Frage eines schlüssigen Gesamtkonzepts von der Physik über die Architektur bis zur Sicherheit.

Die Fragen der Zukunft sind:

- Wie sieht ein Konzept zur schlüssigen Integration mobiler Endgeräte aus?
- Was bedeutet das für unsere Netzwerke?
- Wie sehen Gesamtarchitekturen von Endgerät bis zum App-Shop aus?
- Wie kann die notwendige Sicherheit wirtschaftlich umgesetzt werden?

Wir analysieren für Sie wohin die Netzwerk-Entwicklung geht, wer die Kontrahenten sind, welche Vor- und Nachteile die verschiedenen Alternativen für Ihr Unternehmen haben und wie relevant einzelne Technologien sind. Investieren Sie zukunftsicher in Ihre Netzwerke, das ComConsult Netzwerk Forum 2014 unterstützt Sie dabei.

Moderation: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr.-Ing. Behrooz Moayeri
Preise: € 2.390,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Das Wissensportal

Das Wissensportal

Auf dem Wissensportal finden Sie eine bunte Mischung aus aktuellen Informationen, persönlichen Meinungen und ausführlichen Grundlagen-Artikeln über die gesamte Themenpalette der IT- und Netzwerkwelt. Die Artikel des ComConsult Wissensportals geben Ihnen die Möglichkeit der Stellungnahme, des Kommentars oder der Diskussion mit anderen Lesern. Nutzen Sie diese Gelegenheit, die Sichtweise anderer Spezialisten zu erfahren.

Verdrängt WebRTC das klassische Telefon?

10. Januar 2014 von Markus Geller



Wenn man sich mit Sprach- oder Videodiensten auf Basis von IP beschäftigt, kennt man Begriffe wie Endpunkt oder Agenten. Und so war es bis vor ein paar Jahren üblich, dass mit diesen Begriffen auch automatisch ein physikalisches Endgerät verbunden war, mit dem der Zugriff auf die Kommunikationsinfrastruktur erfolgte...

[Kompletten Artikel lesen unter www.comconsult-research.de](http://www.comconsult-research.de)

AVB ein unterschätzter Standard?

10. Dezember 2013 von Markus Schaub



Als ich den Titel „Audio/Video Bridging“ das erste Mal von Extreme hörte, war mein spontaner Gedanke: wollen die jetzt MCUs bauen? Das, denke ich, ist auch

eines der Hauptprobleme des IEEE Standards: ein missverständlicher Name. Jeder, der nur kurz überfliegt, was es Neues auf dem Ethernet-Markt gibt, wird bei diesem Titel sofort in die nächste Zeile springen. Und jeder, der sich für MCUs interessiert, wird den „Zurück“-Button drücken, wenn er die ersten Zeilen der Charter der Arbeitsgruppe gelesen hat, da er merkt, dass es eigentlich gar nicht um Audio und Video, sondern um irgendwas mit Netzwerken geht...

[Kompletten Artikel lesen unter www.comconsult-research.de](http://www.comconsult-research.de)

IPv6 – welche Interface-Adresse wird genutzt?

8. Januar 2014 von Markus Schaub



Bei IPv4 war die Adressauswahl noch einfach: es gab eine Ziel-Adresse und eine Absende-Adresse. Nur in Ausnahmen gab es mal mehr. Bei IPv6 sind mehrere Adressen pro Interface nicht nur üblich, sondern Pflicht. Somit muss bei jeder Übertragung das Paar von Adressen gefunden werden, wo Send- und Empfangsadresse am besten zusammen passen. Im Folgenden wird dieses Verfahren vorgestellt....

[Kompletten Artikel lesen unter www.comconsult-research.de](http://www.comconsult-research.de)

Auf dem Weg zum mobilen Unternehmen

9. Dezember 2013 von Leonie Herden



Es ist gerade einmal sechs Jahre her, dass das iPhone sich anschickte, das Unternehmen zu erobern. Während der Markt für mobile Endgeräte vorher in eher konservativen Bahnen verlief, explodierte er nach dieser Initialzündung regelrecht. Und schon bald mussten sich die IT-Verantwortlichen mit der Integration dieser Endgeräte in die Unternehmens-IT herumpfen. Im Gegensatz zum vergangenen Vorstands-Hype BlackBerry kamen diese neuen Endgeräte nicht mit einer fertigen Unternehmens-Lösung daher – im Gegenteil: der Mobilmarkt wird fast ausschließlich aus dem Konsumentenmarkt getrieben. Unternehmen wie Apple und Google scheren sich schlicht nicht um den Enterprise-Markt....

[Kompletten Artikel lesen unter www.comconsult-research.de](http://www.comconsult-research.de)

Kommunizieren Sie noch oder verschlüsseln Sie schon?

27. November 2013 von Cornelius Höchel-Winter



Die jüngsten Enthüllungen über das Gebaren amerikanischer und britischer Geheimdienste hat mal wieder das Thema Datenverschlüsselung und insbesondere E-Mail-Verschlüsselung auf der Tagesordnung ganz nach oben gesetzt. Laut einer aktuellen Umfrage des GfK Vereins sorgen sich knapp 70 % der Deutschen um die Sicherheit ihrer persönlichen Daten, gleichzeitig verschlüsseln aber nur 5 % ihren E-Mail-Verkehr, weltweit wird geschätzt, dass lediglich ca. 3 % des E-Mail-Verkehrs verschlüsselt übertragen werden...

[Kompletten Artikel lesen unter www.comconsult-research.de](http://www.comconsult-research.de)

Aktuelles Seminar

Intensiv-Tag Video: Markttrends und Technologien

31.03.14 in Köln

Die ComConsult Akademie veranstaltet am 31.03.2014 ihren "Intensiv-Tag Video: Markttrends und Technologien" in Köln.

Auf dieser eintägigen Sonder-Veranstaltung unter Moderation von Petra Borowka-Gatzweiler wird der aktuelle Markt für Video und Videokonferenzen sowie deren Integration in UC-Lösungen analysiert. Es erfolgt eine Einführung in den aktuellen Technologiestand und Neuentwicklungen, insbesondere neue Standards für die Videotechnologie. Führende Hersteller erläutern ihre Neuentwicklungen, Virtualisierungs-Ansätze und stellen ihre Roadmaps vor. Ziel ist, dabei auch die unterschiedlichen Sichtweisen der Hersteller zu diesem Thema transparent zu machen.

Der Intensiv-Tag behandelt im Einzelnen folgende Themen:

- Technologien und Markttrends für Video-Lösungen
 - Wo liegen die Mehrwerte und Vorteile beim Einsatz von Video?
 - Architekturen und Einsatz-Szenarien für Video-Lösungen
 - Wofür sind funktionale Kernparameter wie Auflösung, fps, Bandbreite, Delay und Verlust-Toleranz wichtig?



- Wie werden Endgeräte mit unterschiedlicher Funktionalität integriert?
- Videofunktionalität von Mobilgeräten und ihre Integration in eine Gesamtlösung
 - Welche Rolle spielt Video für UC?
 - Der Markt für Konferenzsysteme
 - Der Markt für Video as a Service (VaaS)
- Offene Standards als Grundlage für Video-Lösungen
 - LAN Standards: IEEE AVB (Audio/Video Bridging); IEEE 802.1AS, IEEE

- 802.1Q-at, IEEE 802.1Q-av)
- H.264 AVC
- H.264 SVC
- H.265 / HEVC
- HTML 5 und RTC Web
- VP9
- VP9 vs. H265
- Anhand vorgegebener Rahmenfragen: Vorstellung der Neuentwicklungen und Roadmap für Video- und Videokonferenzlösungen von drei bis vier marktführenden Herstellern; z.B.
 - Avaya
 - Cisco
 - Lifesize
 - Polycom
 - Vidyo
- Abschließende Diskussionsrunde
 - Wo liegen die Unterschiede?
 - Welcher Ansatz ist der beste?
 - Wo ist das Optimum?

Die neuen Trends sind absolut ernst zu nehmen und schaffen eine völlig neue technische Basis für die Videokonferenz der Zukunft. Wenn Sie hier mehr erfahren möchten, so besuchen Sie diesen Intensivtag zum Thema Video-Technologien, auf dem wir diese Trends, sowie die Hersteller und ihre Positionierung in den einzelnen Märkten detailliert beleuchten und diskutieren werden.

Fax-Antwort an ComConsult 02408/955-399

Anmeldung

Intensiv-Tag Video: Markttrends und Technologien

Ich buche das Seminar
Intensiv-Tag Video: Markttrends und Technologien

☐ am 31.03.14 in Köln
zum Preis € 990,- netto

Bitte reservieren Sie mir ein Zimmer

vom _____ bis _____ 14

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

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift

Software Updates, die Geißel der IT

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.

Das haben Sie bestimmt schon einmal erlebt: Ein Software Upgrade ist, aus welchen Gründen auch immer, gewünscht oder erforderlich. Sie führen dieses Upgrade exakt nach Anweisung des Herstellers durch und danach geht gar nichts mehr. Nein, ich beziehe mich nicht nur auf Netzkomponenten, so etwas kann Ihnen am PC genauso passieren wie auf einem Smartphone. Nur – bei Netzkomponenten hat ein solcher Fehler meist weitreichende Auswirkungen.

Moment mal, Netze verfügen doch über Redundanzen! Sollte tatsächlich eine Komponente auf Grund eines Software Upgrade ausfallen, sollte doch die Redundanz deren Funktion übernehmen und die Anwender werden nicht beeinträchtigt sein, oder?

So etwas hat einer unserer Kunden neu probiert. Es sollte ein Switch Cluster am Übergang zum Weitverkehrsnetz auf eine neue Software-Version gehoben werden. Der Hersteller gibt vor, wie ein solcher Upgrade durchzuführen ist. Zunächst der eine, dann der andere Switch. Erst wenn auf beiden Switches die neue Software läuft, ist die Redundanz wiederhergestellt. Während der Zwischenschritte ist immer mindestens einer der Switches betriebsbereit. Dennoch lautete das Ergebnis der Aktion: Auf Grund eines Software-Fehlers fiel am Ende der gesamte Switch Cluster aus. Die Weitverkehrsverbindung war für Stunden unterbrochen – ein Desaster!

Switch Cluster sind eine tolle Erfindung. Mehrere Switches teilen sich eine Forwarding Database (MAC-Adresstabelle), bei einigen Herstellern sogar die gesamte Konfiguration. Damit wird es möglich, eine Link Aggregation auf mehrere Switches zu verteilen (Multi-Chassis Link Aggregation, MC-LAG). Das gefürchtete Spanning Tree Protocol (STP) kann dank MC-LAG ausgemerzt werden. Leider nur fordern viele Hersteller, dass alle Switches eines Cluster dieselbe Software-Version nutzen. Fehler in der Software wirken sich somit auf alle



Cluster-Mitglieder gleichzeitig aus, die Redundanz wird untergraben!

Ein anderer Kunde baut eine neue Werkhalle und benötigt dafür neue WLAN Access Points. Der Hersteller fordert für die aktuell lieferbaren Access Points eine bestimmte Software Version, die auf den WLAN Controllern einzusetzen ist. Leider ist diese Software Version nicht kompatibel zu einer Vielzahl vorhandener Access Points in einer bestehenden Werkhalle. Was tun? Zusätzliche WLAN Controller mit neuer Software beschaffen? Oder gleich Komponenten eines anderen Herstellers einsetzen in der trügerischen Hoffnung, dass es derlei Inkompatibilitäten bei dem nicht geben wird? So oder so, die Sache wird teuer, und die Zeche zahlt (natürlich) der Kunde.

Das bittere Fazit ist mal wieder: Redundanz schützt vor Ausfall nicht. Und: Sie hängen in jedem Fall „am Fliegenfänger“ des Herstellers. Welche Empfehlungen ergeben sich daraus?

Als erstes sollten Sie genau überlegen, wo Sie Netzkomponenten zu Clustern zusammenfassen und sich also auf Herstellerspezifische Mechanismen verlassen wollen. Im Zweifel verlassen Sie sich besser auf die altbekannten standardisierten Redundanzverfahren à la OSPF oder sogar STP. Diese Verfahren funktionieren sogar zwischen Komponenten unterschiedlicher Hersteller, Software-Versionsunterschiede sind hier meist kein Problem.

Darüber hinaus setzt natürlich jeder Software Upgrade ein genaues Studium der entsprechenden Release Notes voraus. Und danach hat immer ein Labortest zu erfolgen. Im Zweifel bleiben Sie besser bei der funktionierenden Software. Nicht jedes neue Feature wird benötigt und nicht jede Sicherheitslücke ist in Ihrer Umgebung relevant (was zu begründen wäre): Auch hier gilt der bekannte Spruch „never change a running system“.

Zuletzt gilt es, Druck auf die Hersteller auszuüben. Es kann einfach nicht sein, dass der Software Support bereits nach wenigen Jahren aufgekündigt wird und Sie am Ende funktionsfähige Komponenten wegwerfen müssen. Das ist weder wirtschaftlich noch nachhaltig. Vielleicht ist ein Anbieterwechsel letztlich doch hilfreich, um den Herstellern die Augen zu öffnen.

Kongress

ComConsult Netzwerk Forum 2014 24.03. - 26.03.14 in Königswinter

- Dienstneutralität kontra Applikations-Integration: Kampf um die Steuerung
- Netzwerk-Design in Zeiten von 10/40/100: wie sieht die Zukunft aus?
- Mobile Endgeräte im Netzwerk: teuer, unsicher, unbeherrschbar?

Das ComConsult Netzwerk Forum 2014 stellt sich den herausragenden Fragen der Entwicklung der Unternehmensnetzwerke. Nach 10 Jahren des Konzept-Stillstands ist das traditionelle Netzwerk unter Druck sowohl aus der Server- als auch aus der Endgeräte-Welt. Kombiniert man dies mit dem notwendigen Bandbreitenwechsel auf 10/40/100, dann steht das Unternehmensnetzwerk der Zukunft auf dem Prüfstand.

Moderation: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr.-Ing. Behrooz Moayeri
Preise: € 2.390,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Zweitthema

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

Fortsetzung von Seite 1

Eine strategische Technologie kann eine bestehende Technologie sein, die gereift und/oder für einen weiteren Bereich von Anwendungen brauchbar ist. Es kann aber auch eine erst im Entstehen begriffene Technologie sein, die Gelegenheiten für strategische geschäftliche Vorteile für Unternehmen, die hier früh einsteigen, nach sich ziehen kann oder das Potential hat, den Markt innerhalb der nächsten fünf Jahre massiv zu verändern. Derartige Technologien haben einen weit reichenden Einfluss auf langfristige Programme, Planungen und Initiativen von Unternehmen und Organisationen.

Es ist natürlich keineswegs so, dass auf ein Unternehmen alle genannten Top Trends gleichzeitig und in gleichem Maße zukommen. Wesentlich ist es hier vor allem, aufmerksam zu bleiben.

Und das sind die Top 10 Technologie-Trends:

1. Mobile Device Diversity and Management

In den nächsten Jahren (betrachtet wurde ein Zeitraum bis 2018) machen die wachsende Vielfalt von Geräten, unterschiedliche Methoden, wie man mit der Informationsverarbeitung umgeht, unterschiedliche Benutzer-Kontexte und Interaktions-Paradigmen universelle Strategien, die überall gleichartig erfolgreich eingesetzt werden können, faktisch unerreichbar. Eine unerwartete Konsequenz von BYOD-Programmen ist die Verdopplung oder sogar Verdreifachung der mobilen Geräteflotte. Dies erzeugt sowohl in der IT als auch hinsichtlich der Kosten einen erheblichen Druck. Regelwerke hinsichtlich der Benutzung eigener Hardware durch die Mitarbeiter eines Unternehmens müssen komplett überarbeitet und bei Bedarf aufgefrischt und



Dr. Franz-Joachim Kauffels ist Technologie- und Industrie-Analyst und Autor. Seit über 30 Jahren unabhängiger, kritischer und oft unbequemer Bestandteil der Netzwerkszene. Verfasser von über 20 Büchern in über 70 Ausgaben sowie über 2000 Artikeln, Videos und Reports.

erweitert werden. Die meisten Unternehmen haben heute nur Regelwerke für den Umgang der Mitarbeiter mit Geräten, die vom Unternehmen bereitgestellt werden. Für BYOD-Strategien entsteht ein Spannungsfeld aus Erwartungen der Mitarbeiter, was sie denn dürfen und was nicht. Dabei ist eine Balance zwischen Flexibilität, Vertraulichkeit und Schutz der privaten Daten anzustreben.

Dieses Thema wird für die Unternehmen ein Dauerbrenner, denn die Entwicklung steht nicht still. Ganz im Gegenteil ist für die nächsten Jahre mit einem wesentlichen Anstieg der Vielfalt zu rechnen. Während die Meisten z.B. bei Smartphones vornehmlich zunächst an Hersteller wie Apple, Samsung oder Nokia denken, gibt es schon große Märkte, in denen z.B. Apple statistisch nur unter „andere“ gelis-

tet wird. Blicken wir z.B. auf Indien, das ist einer der größten denkbaren Märkte mit erheblichem Wachstumspotential und teilweise schwierigen Bedingungen, alleine durch Größe und infrastrukturelle Probleme des Landes. Hier sieht es so aus, dass direkt hinter Samsung die indischen Hersteller Micromax und Karbonn kommen, die jeweils die gesamte Palette möglicher moderner Endgeräte abdecken, dabei aber auch weitere neue Android-Varianten schaffen, die zur Unterscheidung die Namen von Süßspeisen haben. (siehe Abbildung 1)

Apple hat zwar seinen Umsatz in Indien durch Preis-Senkungen, Rückkäufe und weitere Aktionen steigern können, genutzt hat es aber nichts, denn gleichzeitig sanken die Margen dramatisch. Der primäre Fehler in diesem Fall ist, dass Apple kein

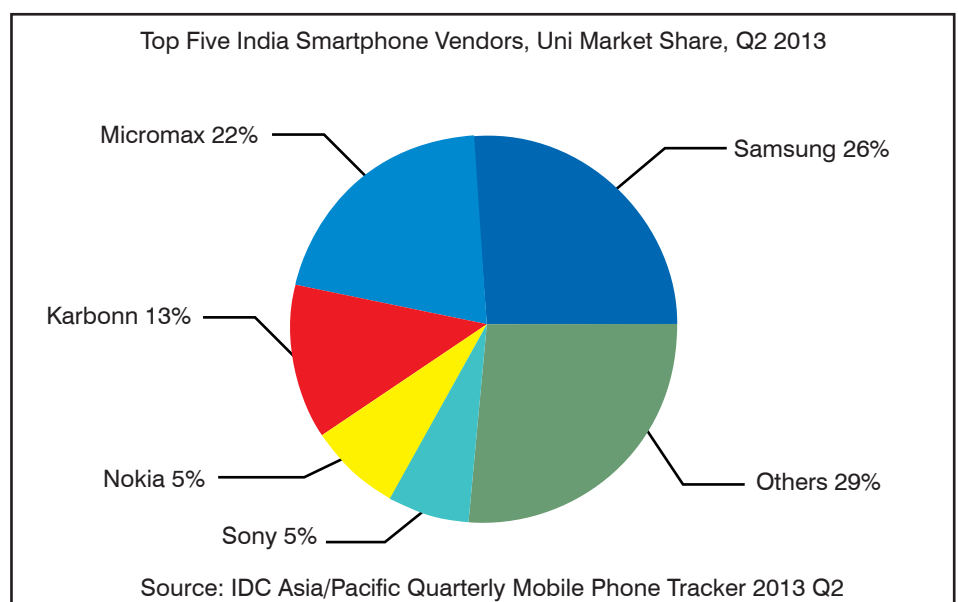


Abbildung 1: Neue Angreifer mischen auf

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

wirklich billiges Handy im Programm hat, was in diesem Zielmarkt erwünscht gewesen wäre. (siehe Abbildung 2)

Bei den neuen Deals mit China Mobile hat Tim Cook sicherlich auch einige Kröten süß-sauer schlucken müssen. Das Problem für alle Hersteller ist einfach, dass es keine technische Großtat mehr ist, eines der neuen Endgeräte zu bauen, weil auch der Integrationsgrad von der Chipseite her immer weiter steigt. Theoretisch wäre es schon heute durchaus möglich, ein „Ein-Chip-Handy“ auf den Markt zu bringen. Was immer problematischer wird, ist neue Funktionen hinzuzufügen, die den Verbraucher anlocken und die der Wettbewerb nicht hat. Die Steigerung einer Retina-Auflösung könnte man nur nutzen, wenn man sich Adlernaugen einpflanzen ließe, was sicherlich merkwürdig aussieht. In den nächsten Jahren wird noch hinzukommen, dass langsam aber sicher Patente auslaufen. Sie haben üblicherweise eine Laufzeit zwischen 17 und 20 Jahren, das US Patentamt hat jetzt einen Laufzeit-rechner online gestellt, unter www.uspto.gov. Verbunden mit der Möglichkeit, über das Internet weltweit zu verkaufen, wird es deutlich mehr Anbieter geben als heute. Schließlich gibt es permanent Neuentwicklungen. Google hat vor Kurzem seine Brille vorgestellt, ein ideales Endgerät für alle, die während der Kommunikation beide Hände frei haben müssen. Ab jetzt wird es für ein Unternehmen oder eine Organisation weder übersichtlicher noch einfacher. Neue Hersteller, neue Geräte-Varianten, neue Geräte-Typen und neue Betriebssystem-Möglichkeiten werden die Vielfalt unumkehrbar vergrößern.

2. Mobile Apps und Anwendungen

Gartner sagt vorher, dass im Laufe von 2014 eine verbesserte JavaScript Leistung HTML5 und den Browser als primäre Umgebung für die Anwendungsentwicklung in/für Unternehmen vorantreiben wird. Gartner empfiehlt den Entwicklern, Modelle für erweiterte Benutzer-Schnittstellen zu entwickeln, die reichere Sprach- und Video-Elemente enthalten und auf diese Weise Menschen auf neuen Wegen miteinander verbinden können. Apps werden weiterhin wachsen und konventionelle Anwendungen werden beginnen zu schrumpfen. Apps sind kleiner und mehr zielgerichtet, während größere Anwendungen umfassender sind. Entwickler müssen nach Wegen suchen, wie man Apps zusammenstecken und so zu größeren Anwendungen machen kann. Der Aufbau von Benutzerschnittstellen, die eine Vielfalt von Geräten umfassen, setzen ein Verständnis fragmentierter Bausteine und eine anpassungsfähige Programmierstruktur voraus,

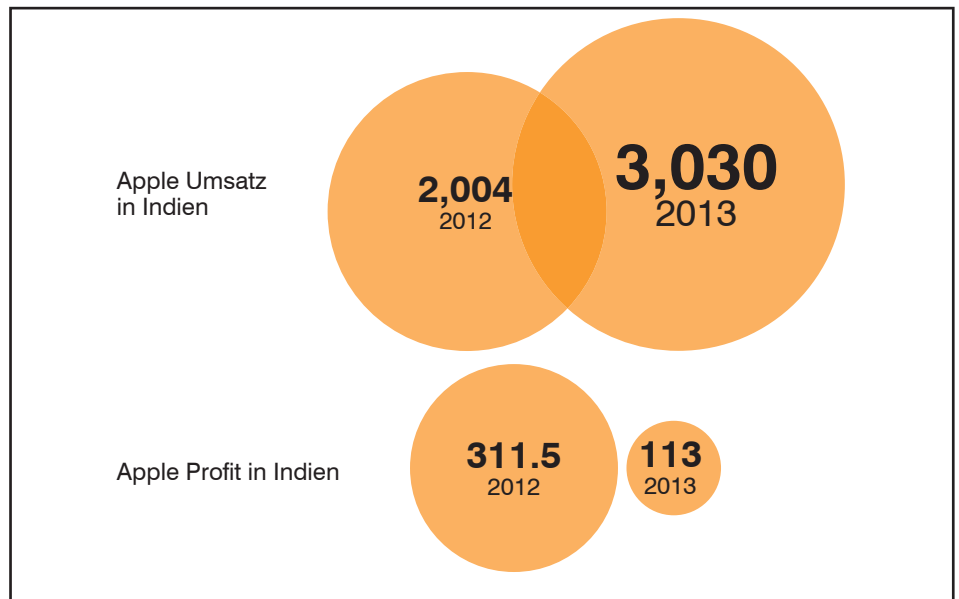


Abbildung 2: Der Umsatz steigt, die Marge sinkt

Quelle: Economic Times

die sie in die Lage versetzt, für jedes Gerät optimierten Kontext zu liefern. Der Markt für Tools zur Schaffung von Apps für Verbraucher und Unternehmen ist mit über 100 möglichen Tool-Anbietern komplex. In den nächsten Jahren wird es kein singuläres Tool geben, welches für alle Arten mobiler Anwendungen optimal ist. Man muss also erwarten, mehrere verschiedene einzusetzen. Die nächste Evolution hinsichtlich der Benutzer-Erfahrung wird die Möglichkeit sein, vorgeplante Funktionen einfließen zu lassen, die von Emotionen und Aktionen gesteuert werden, um Änderungen im Verhalten der End-Benutzer herbeizuführen.

Die zukünftige Verwendung von HTML5 ist kaum strittig. Angesichts der geschilderten Situation bei der Diversifizierung der Endgeräte muss es eine Schnittstelle geben, die überall läuft. Rückwirkend betrachtet ist genau die Kombination aus universeller Schnittstelle und einem übersichtlichen Programm (Browser) das wesentliche Erfolgsmodell für das Internet. Erst mit diesen Hilfsmitteln konnten die Massen erschlossen werden, das E-Business in Gang gebracht und letztlich die Kommunikation erheblich verbessert werden. Die Basis-Idee einer Web-Architektur ist die Kommunikation zwischen Browser und Web-Server auf der Grundlage von http-Requests und -Responses. Funktionen wie Präsentation, Rendering, Caching sowie Authentifizierung und Cookies liegen beim Browser, der Web-Server beinhaltet in einer einfachen Struktur statische Seiten und bedient die Anforderungen des Browsers. Mit der Zeit wurden die Webseiten aber dynamisiert, das Ergebnis haben Sie jeden Tag vor Augen.

Wenn Sie eine Webseite aufrufen, ist diese nicht statisch, sondern wird im Moment des Aufrufs erst zusammengesetzt. Das hat unterschiedliche Motivationen, eine sehr wesentliche Motivation ist aber neben der Einbindung artfremder Formate vor allem die Aktualität. Eine Webseite mit Börsenkursen würde Ihnen ohne aktuelle Kurse nichts nützen. Also besteht diese Webseite aus einem Rahmen und dynamischen Elementen, die in diesen Rahmen eingebaut werden. Eine weitere Abwandlung gegenüber einer statischen Webseite ist die Möglichkeit der Einbindung von Anwendungen, z.B. solchen, die aktuelle Daten holen, Authentifizierungen und Konversionen durchführen usw. Das ist alles an sich nichts wirklich Neues und führt zu einer dreilagigen Architektur. In dieser Architektur werden Darstellung, Anwendungen und Daten getrennt. Bleiben wir bei der Börsenseite. Der Rahmen, den der Benutzer am Ende sieht, ist die Darstellung, die Anwendungen holen die aktuellen Daten und passen sie an die gewünschte Darstellung an und die Daten selbst liegen auf einem Transaktionsserver z.B. beim Betreiber der Börse.

Ein eindrucksvolles Beispiel für die aktuellen Möglichkeiten ist der vollständig in HTML5 geschriebene Video-Player von Vimeo, der als bester Multiplattform-Player des Marktes gilt. (siehe Abbildung 3)

Ein weiteres eindrucksvolles Beispiel ist Keynote von Apple. Es ist für MAC- und iOS-Geräte verfügbar. Dateien können aber auch in einem Power Point-kompatiblen Format abgelegt werden. Via iCloud können Präsentationen jeweils beliebig geteilt werden. Sie holen auf al-

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

len Endgeräten das jeweilige Maximum an Qualität heraus. Sehr praktisch ist aber auch die Synchronisation über die iCloud, die alle für den Benutzer definierten Endgeräte automatisch auf dem neuesten Stand hält, an welchem Gerät man auch Änderungen durchführt. (siehe Abbildung 4)

Seit Aufkommen der ersten Browser hat sich die in den betreffenden Endgeräten verfügbare Leistung nach Moore's Law bis heute theoretisch grob um den Faktor 10.000 – 50.000 erhöht. Viele Geräte kommen aber mit deutlich weniger aus. Eine moderne Schnittstelle darf daher schon deutlich anspruchsvoller sein als ihre Vorgänger. Auf den nächsten Generationen auch mobiler Prozessoren ist durchaus auch eine Virtualisierung denkbar, die letztlich zu einer eleganteren Funktionsverteilung genutzt werden könnte. Hinsichtlich des „Zusammensteckens“ von Apps gibt es im Markt durchaus Tendenzen, Methoden der Objekt-orientierten Programmierung aufzunehmen, die einerseits in viel höherem Maße zu wiederverwendbarem Code führen, andererseits aber auch die Definition innerer Schnittstellen erleichtern, die für das „Zusammenstecken“ der Apps genutzt werden kann, wie Gartner das so schön ausdrückt. Theoretische Überlegungen, Techniken und Definitionen hierfür gibt es schon seit über 10 Jahren, sie nutzen aber nur wenig, wenn die überwiegende Zahl der Endgeräte zu wenig Leistung hat. In /1/ findet der Interessent eine gut verständliche allgemeine Einführung und in /2/ eine weitere Einführung mit praktischen Beispielen auf Basis von Object Oriented JAVA. Auf allen Plattformen gibt es schon seit sehr langer Zeit Entwicklungsbibliotheken mit entsprechenden Objekten. In Open Source Bereich ist das ohnehin Standard, aber auch alle Entwicklungen von Microsoft basieren darauf. Elemente wie Menue-Führung, Fensterverwaltung, Button Aktionen usw. müssen ja nicht immer neu programmiert werden. Das Endgerät wird dadurch nicht mehr oder weniger belastet. Eine spannende Frage aktuell ist, wie viel Javascript-Funktionalität in den Browser verlegt wird, siehe Web RTC. Solche Funktionen brauchen danach nicht mehr programmiert, sondern nur noch aufgerufen und parametrisiert werden. Man kann es eigentlich gar nicht anders beschreiben als ein zum Bersten gefülltes Fass: viel bessere Programmiertechniken als die aktuellen warten sehnsüchtig auf Erlösung durch hinreichende Prozessorleistung. Persönlich bin ich davon überzeugt, dass der Wechsel eher explosionsartig und weniger kontinuierlich erfolgt. Sobald das geschehen ist, werden wir die heute noch teilweise sehr primitiven Apps nicht mehr wieder-

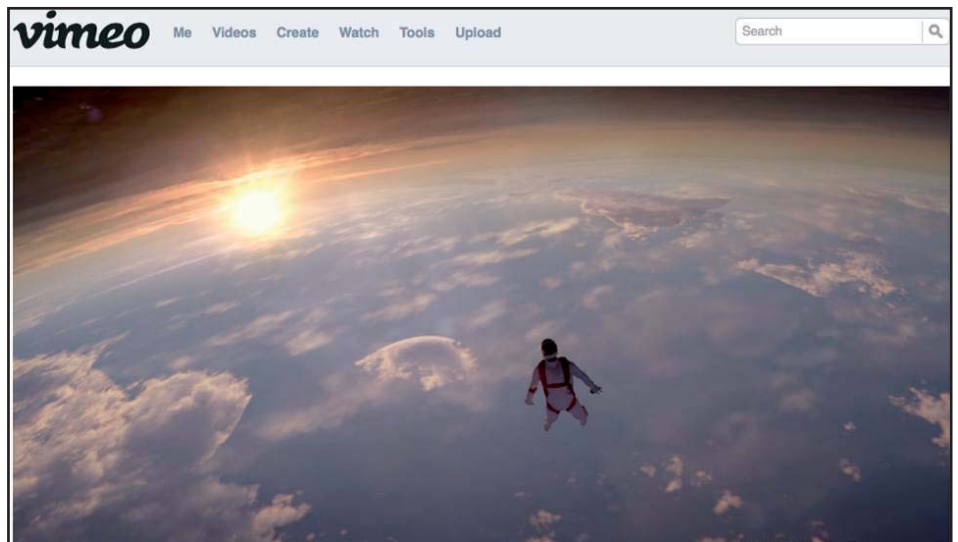


Abbildung 3: Vimeo Video Player



Abbildung 4: Apple Keynote

Quelle: Apple

erkennen. Und in diesem Zuge sind eben auch erheblich erweiterte Modelle für die Benutzerschnittstelle möglich. Die Änderung des Benutzerverhaltens durch vordefinierte Funktionen ist an sich nichts Neues: auch wenn man, wie der Autor, die alte Rechtschreibung noch „drin“ hat und intuitiv immer noch benutzt, streicht Word dadurch entstehende Fehler immer wieder an. Irgendwann gibt man auf und schreibt das Wort, wie Word es möchte.

3. Internet of Everything

Das Internet dehnt sich auch jenseits von PCs und mobilen Endgeräten auf Einrichtungen von Unternehmen, wie z.B. Feldgeräte, und technische Gegenstände bei Benutzern, wie Autos oder Fernseher, aus. Das Problem ist, dass die meisten Unternehmen und Anbieter von Technologie die Möglichkeiten eines in diesem Sinne ausgedehnten Internets erst einmal er-

kunden müssen und weder operativ noch von ihrer Organisation her darauf eingerichtet sind. Betrachten wir die Digitalisierung der meisten wichtigen Produkte, Dienstleistungen und Aktiva. Die Kombination von Datenströmen und Dienstleistungen, die dadurch entstehen, dass man alles digitalisiert, führt zu vier grundsätzlichen Nutzungsmodellen: Verwaltung, Kapitalisierung, Betrieb und Erweiterung. Diese vier Modelle können auf jedes der vier „Internets“ angewendet werden (Personen, Dinge, Information und Orte). Unternehmen sollten sich nicht auf die Denkweise beschränken, dass nur das „Internet of Things“ (z.B. mit Maschinen und Vermögenswerten) das Potential zur Entwicklung der vier Nutzungsmodelle hat. Unternehmen aus allen Industriezweigen können diese vier Modelle nutzen.

Das IoT oder IoE ist momentan ein unglaublicher Hype. Die Aktien von Chip-

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

Herstellern wie NXP-Semiconductors (NASDAQ: NXPI) oder Micron Technologies (NASDAQ: MU) haben jeweils deutliche Schübe erhalten, wenn die Hersteller nur ankündigten, etwas für das IoT herzustellen. Genauso hat der Markt die Äußerungen von Cisco dazu komplett ignoriert. Das gehört zu einer solchen Hype-Frühphase. Abgesehen von teilweise noch fehlenden technischen Möglichkeiten zur Umsetzung, müssen Unternehmen in weiten Bereichen überhaupt erst einmal eruieren, welche Potentiale für sie bestehen. Natürlich gibt es hier Selbstläufer wie Zählerfernablesung oder intelligente Steuerungen für Wohnungen, aber das ist nur der Beginn.

Wie sehr die Entwicklung voranschreitet, kann man sicherlich sehr eindrucksvoll mit dem von Intel auf der Consumer Electronics Show 2014 in Las Vegas vorgestellten „Edison“ verdeutlichen. Intel hat sein System-on-a-Chip-System „Quark“ auf die Größe einer SD-Karte geschrumpft. Dieser komplette PC der Pentium-Klasse wurde in 22 nm-Technologie gebaut unterstützt das vollständige Linux-Betriebssystem und kann mit WiFi und Bluetooth 4.0 mit dem Rest der Welt kommunizieren. Ein wichtiger Anwendungsbereich ist „Wearable Computing“. Intel demonstrierte nicht nur Baby-Kleidung, die die wichtigsten Vital-Signale eines Kleinkindes automatisch überwacht, mit dem Edison, sondern hat auch eine Initiative „Make it Wearable“ ins Leben gerufen, bei der die besten Ideen aus der kreativen möglichen Kundschaft aufgenommen werden sollen. Immerhin setzt Intel 1,3 Mio US\$ als Preise für die besten Ideen zu Edison-Anwendungen aus. Es wird einen App-Store für Edison-Anwendungen geben. Nach Meinung von Analysten kostet der Edison in größerer Stückzahl nicht mehr als 30 US\$. Die Überwachung eines Kleinkindes ist sicher keine Schlüsselanwendung und die meisten Kleinkinder werden auch sicher wie bisher ohne solche Funktionen groß. Insgesamt ist „Wearable Computing“ ein sehr großer Markt, der bis 2018 nach Prognosen von Juniper Research auf rund 19 Mrd US\$ wachsen wird. Alleine Apple sollte für eine bis zu 350 US\$ teuren „Uhr“ in Ergänzung des iPhones bereits im ersten Jahr durchaus 10 Mio. Abnehmer finden können. Die aktuellen Modelle für Smartwatches z.B. von Samsung sind bisher nicht gut angenommen worden, einfach deshalb, weil sie funktional überfrachtet und schlicht und ergreifend so hässlich wie die ersten Digitaluhren der 70er Jahre sind. Das freut vielleicht einen Nerd, aber die stilbewusste Kundschaft möchte hier mehr. Immerhin werden sich bedeutende Gelgenheiten für App-Programmierer auf den Bereichen Gesundheit, Fit-

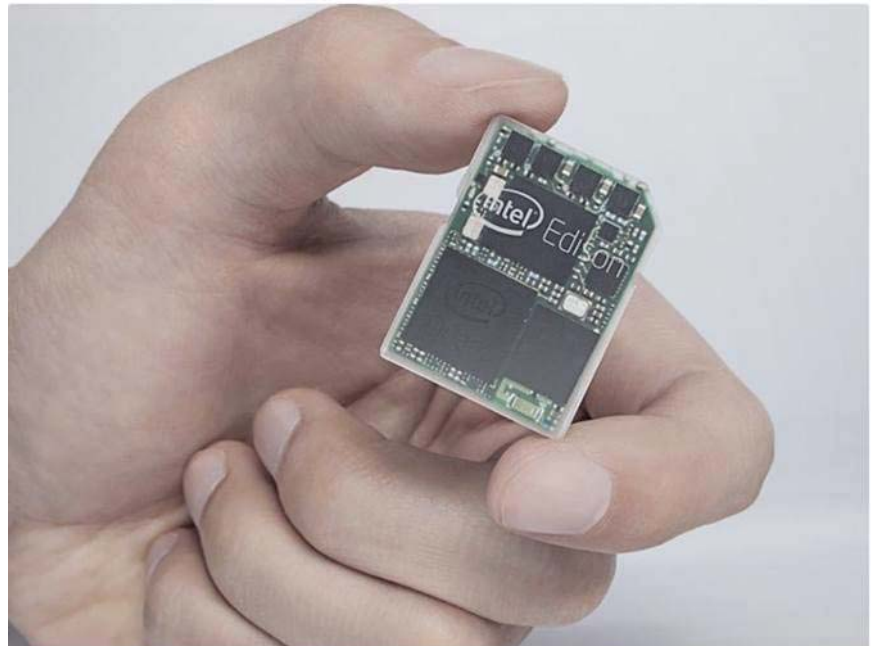


Abbildung 5: Intel Edison

Quelle: Intel

ness, Sport, Lifestyle und Kommunikation ergeben. Es gibt Unmengen Designstudien und Gerüchte über das Aussehen der iWatch, deshalb habe ich einfach Abbildung 6 genommen, die ich persönlich für Damen am elegantesten finde.

das Thema hier mit einem Hinweis darauf abschließen.

4. Hybrid Cloud und IT als Service Broker

Es ist unbedingt erforderlich, Personal Clouds und externe Private Cloud Services zusammenzuführen. Unternehmen sollten private Cloud Services mit einer hybriden Zukunft im Hinterkopf entwickeln und si-



Abbildung 6: iWatch Designstudie Isamu Sanada

Quelle: Apple

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

herstellen, dass eine zukünftige Integration und Interoperabilität möglich ist. Hybride Cloud Services können auf viele Arten aufgebaut werden und konstruktiv von recht statisch zu sehr dynamisch variieren. Die Verwaltung einer derartigen Komposition wird es vielfach erforderlich machen, die Rolle des Cloud Service Brokers mit Leben zu füllen, der Aggregation, Integration und Kundenanpassung von Diensten vornimmt. Unternehmen, die sich von der privaten Cloud in eine Hybrid Cloud Umgebung weiterentwickeln, kommen automatisch in diese Problematik. Begriffe wie „Overdrafting“ und „Cloudbursting“ werden oft herangezogen, um zu beschreiben, was Hybrid Cloud Computing möglich machen wird. Die überwiegende Mehrzahl von Hybrid Cloud Services wird aber zu Beginn wesentlich weniger dynamisch sein. Frühe Hybrid Cloud Services werden sehr statische konstruierte Kompositionen sein (wie z.B. die Integration zwischen einer internen privaten Cloud und einer öffentlichen Cloud für eine bestimmte Funktionalität oder Daten). Erst mit der Verfügbarkeit von Cloud Service Brokern werden sich ausgeprägtere Strukturen entwickeln wie z.B. private IaaS-Angebote, die die Leistungen externer Service-Anbieter auf dem Wege von Regelwerken und Nutzungs-Szenarien einbeziehen können.

Hybrid Clouds für BYOD-Support oder generell Anbindung mobiler Geräte gehören sicherlich zu den ersten Anwendungsfällen eines solchen Szenarios. Einem Unternehmen wird mittelfristig nichts anderes übrig bleiben, als die neuartigen Geräte mit dem Gleichen zu unterstützen, was ihnen auch sonst erst das Überleben trotz Speichermangels erlaubt: einer Cloud. Anscheinend bleibt zunächst nichts anderes übrig, als den Cloud Service als private Cloud selbst aufzubauen. Dabei muss es gar kein Cloud Service im engen Sinne sein, sondern es wird vielfach ausreichend sein, Rechen- und Speicherleistung mit modernen Geräten so zur Verfügung zu stellen, dass die Tablets eben genau so versorgt werden, wie es die auf ihnen laufenden Unternehmensanwendungen brauchen. Doch mittelfristig ist es damit nicht getan.

Solange sich ein Benutzer in den Gebäuden des Unternehmens befindet, kann man deren Versorgung aus dem eigenen RZ via WLAN-Infrastruktur vornehmen. Hier ist die Frage spannend, wie man die Versorgung noch sicherstellen kann, wenn sich der Mitarbeiter aus dem Gebäude herausbewegt. Grundsätzlich wird ein Endgerät beim Verlassen des lokalen Versorgungsbereiches, also z.B. eines Gebäudes mit flächendeckender

WLAN-Versorgung im Rahmen einer LTE-Versorgungsstruktur letztlich in einen öffentlichen Cloud-Service eines Providers eingebunden. Von seiner Struktur her ist der Provider in der Lage, die einzelnen Benutzer und ihre Verkehrsströme völlig voneinander zu isolieren. Das ist seine Basis für ein weitergehendes funktionales Angebot, die Abrechnung gegenüber den Teilnehmern und eine grundsätzliche Sicherheitsmaßnahme. Wie der Provider das alles genau macht, bleibt dem Kunden verborgen und er hat heute auch keinerlei weiteren Zugriff auf die Infrastruktur des Leistungsanbieters. Für Corporate Customers gibt es weiter gehende Angebote z.B. hinsichtlich der Durchreichung von Qualitätsparametern. Heute beschränken sich diese Angebote aber vorwiegend auf Aspekte der logischen Verbindung. Das muss aber nicht so bleiben. Es kann durchaus sinnvoll sein, dass ein Provider auch andere Elemente seiner Infrastruktur als IaaS Infrastrukturservice verfügbar macht. Wir können z.B. davon ausgehen, dass der Provider eine umfassende virtualisierte Speicherlösung betreibt. Aus dieser ergeben sich zusätzliche Angebote wie z.B. eMail und Web-Zugriff sowie die Möglichkeit, in einem weiter unspezifizierten, aber in der Regel hersteller-neutralen Cloud-Angebot Daten abzulegen, wie das heute schon z.B. mit der Telekom-Cloud möglich ist. Letztlich wird es Aufgabe von Unternehmen und Providern sein, diese so gewachsenen Infrastrukturen so aufeinander abzustimmen, dass ein Benutzer unabhängig von seinem Standort innerhalb oder außerhalb des Unternehmens durchgängig und ohne sichtbaren Bruch auf die für seine Arbeit relevanten Anwendungsdienste und Daten in sicherer und performanter Form zugreifen kann.

Jeder Besitzer eines „neuartigen Endgerätes“ lebt schon heute in einer Hybrid Cloud bestehend aus einem privaten Teil mit dessen Möglichkeiten, die ja auch durchaus eine unabhängige Speicherlösung umfassen können, und einem öffentlichen Teil, der durch den Provider implementiert ist. Je nach Versorgungslage mit Funkabdeckung kann er von privaten in den öffentlichen Teil wechseln und vice versa. Die einzigen Unterschiede ergeben sich im Hinblick auf die Reaktionsfähigkeit durch die unterschiedliche Geschwindigkeit der Netzanbindungen. Und genau diese Grundqualität erwartet er auch hinsichtlich der Unternehmensanwendungen, die er für seine Arbeit benötigt! Die Hybrid Cloud hat darüber hinaus noch einige wichtige weitere Vorzüge hinsichtlich der Redundanz, die man je nach Anwendungsfall auch gut brauchen kann. Beim Ausfall des RZs des Unternehmens kann

der Mitarbeiter auf dem öffentlichen Teil der Hybrid Cloud weiterarbeiten. Für die nähere Zukunft werden Sonderfunktionen, wie sie sich im Rahmen der Virtualisierung ergeben, eine wichtige Rolle spielen. Darauf kommen wir gleich nochmal.

5. Cloud/Client-Architektur

Cloud/Client-Computing-Modelle befinden sich in einem Umbruch. In einer Cloud/Client-Architektur ist der Client eine reichhaltige Anwendung, die auf einem mit dem Internet verbundenen Gerät läuft und der Server besteht aus einer Menge von Anwendungs-Services, die auf einer in zunehmendem Maße skalierbaren elastischen Cloud Computing Plattform beherbergt sind. Die Cloud ist Kontrollstelle und System und Anwendungen können multiple Client-Geräte umfassen. Die Client-Umgebung kann eine native Anwendung oder Browser-basiert sein. Die zunehmende Leistung des Browsers ist für viele stationäre und mobile End-Geräte verfügbar. Robuste funktionale Möglichkeiten in vielen mobilen Endgeräten, die wachsenden Anforderungen an Netze, die Kosten von Netzen und die Notwendigkeit der Verwaltung der Nutzung der Bandbreite schaffen Anreize, in manchen Fällen den Bedarf für Rechenkapazität und Speicher einer Cloud-Anwendung zu minimieren und im Gegenzug die Rechen- und Speicherkapazität im Client-Gerät zu nutzen. Die zunehmend komplexer werdenden Anforderungen mobiler Benutzer werden aber in Zukunft dazu führen, dass Apps immer mehr Rechen- und Speicher-Kapazität auf der Server-Seite benötigen.

Seit es Rechner und Netze gibt, gibt es immer wieder aufflammende Diskussionen über Funktionsverteilung. Eine frühe Variante war, ob die für einen Ethernet-Zugriff nötigen Firmware-Funktionen nun von einem Ko-Prozessor auf der Karte oder doch besser vermöge eines Prozesses im Rechner realisiert werden sollen. Aktuell haben wir viele Neuauflagen dieser Thematik, die sich um Hardware-Offload komplexerer Kommunikations- und Speicherfunktionen drehen. Bei der VM-Migration sieht man, dass ein optimiertes Verfahren einen Prozessor nur zu wenigen Prozent auslastet, ein schlechtes zu 70 bis 80 %. Es gibt dabei nie einen dauerhaften Gewinner, sondern immer nur eine ganz gute architekturelle Alternative, die so lange benutzt wird, bis sich Komponenten wieder dramatisch weiterentwickelt haben. Das findet nicht nur auf einer Kleinteile-Ebene statt, sondern wir haben ja auch Diskussionen darüber, ob wir nicht ganze Funktionseinheiten in Netzen angesichts steigender Leistung der virtualisierten Systeme schlicht und ergrei-

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

fend komplett in Software implementieren können. Wie schnell Betrachtungsweisen kippen können, muss selbst Gartner mit dieser These erfahren. Apps werden nur dann wirklich mehr Rechen- und Speicherkapazität auf der Serverseite benötigen, wenn die Endgeräte, auf denen sie laufen, relativ schlapp sind. Die These hätte nur dann längeren Bestand, wenn die relative Schwäche von Endgeräten eine Zeit lang anhalten würde. Aber genau das Gegenteil ist der Fall: mobile Endgeräte haben sich in 2013 fulminant weiterentwickelt, bei manchen gab es sogar einen regelrechten Sprung. Und das ist erst der Anfang. Intel und Konsorten schieben auch auf der Ebene der Mobilprozessoren immer weiter fleißig Leistung nach. Das könnte sogar noch mehr sein, wenn man nicht immer wieder auf den Stromverbrauch achten müsste. Aber auch hier ist durch die Einführung der 22- und 14nm-Technik für den Bau der „Atom“-Prozessoren ein wesentlicher Durchbruch gelungen, weil in dieser Technik ein nunmehr dreidimensionaler Tri-Gate Feldeffekt-Transistor erheblich weniger Leistung benötigt, um umzuschalten als in den Vorgängergenerationen. (siehe Abbildung 7)

Die ersten Modelle, die einer wirklich neuen Leistungsklasse zugeordnet werden können, sind das iPad 5 von Apple und das Surface 2 PRO. Mit Windows 8 steht auf dem Surface PRO standardmäßig Virtualisierung bereit, weil Windows 8 MS Hyper-V unterstützt, in Abbildung 8 sehen wir ein Szenario für Windows 7.

Und genau das ist für meine Begriffe der Anfang einer völligen Neubewertung der Rollen von Cloud-Servern und Endgeräten: wenn die Endgeräte Virtualisierung unterstützen, können viele der Vorteile, die die Virtualisierung auch schon in RZs erfolgreich gemacht haben, übernommen werden. Zunächst wird es um besseren Schutz, Versionenparallelität und weitere lokale Organisationsaufgaben gehen. Aber die Virtualisierung bietet ja weitere Konzepte, besonders hinsichtlich der Ausfallsicherheit und für den unterbrechungsfreien Betrieb. Also wird es in absehbarer Zeit möglich sein, dass VMs zwischen einer Cloud und den Endgeräten hin- und herwandern. Die so geschaffenen dynamischen Umgebungen führen auch zu völlig neuen Möglichkeiten hinsichtlich der BYOD-Thematik. Aber auch heute kann man mit schlauneren Endgeräten sehr viel machen. VMware hat ja vor einiger Zeit den Hersteller Desktope gekauft. Dessen Hauptprodukt ist eine Software für Desktop-Virtualisierung aus der Cloud. Ein Cloud-Service-Provider kann das Produkt einsetzen, um Unternehmen einen Service anzubieten, der sie nach Anga-

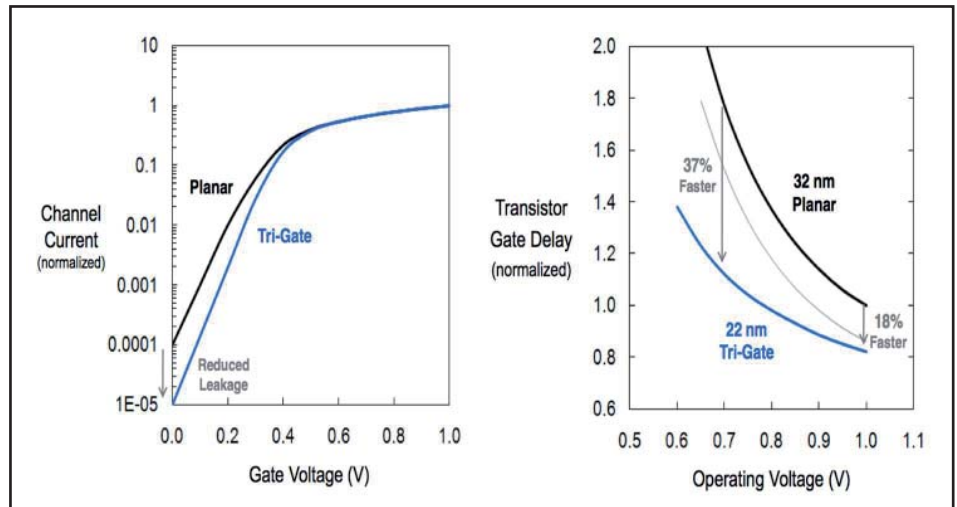


Abbildung 7: Planar-Transistor vs. Tri-Gate: Leistungsmerkmale

Quelle: Intel

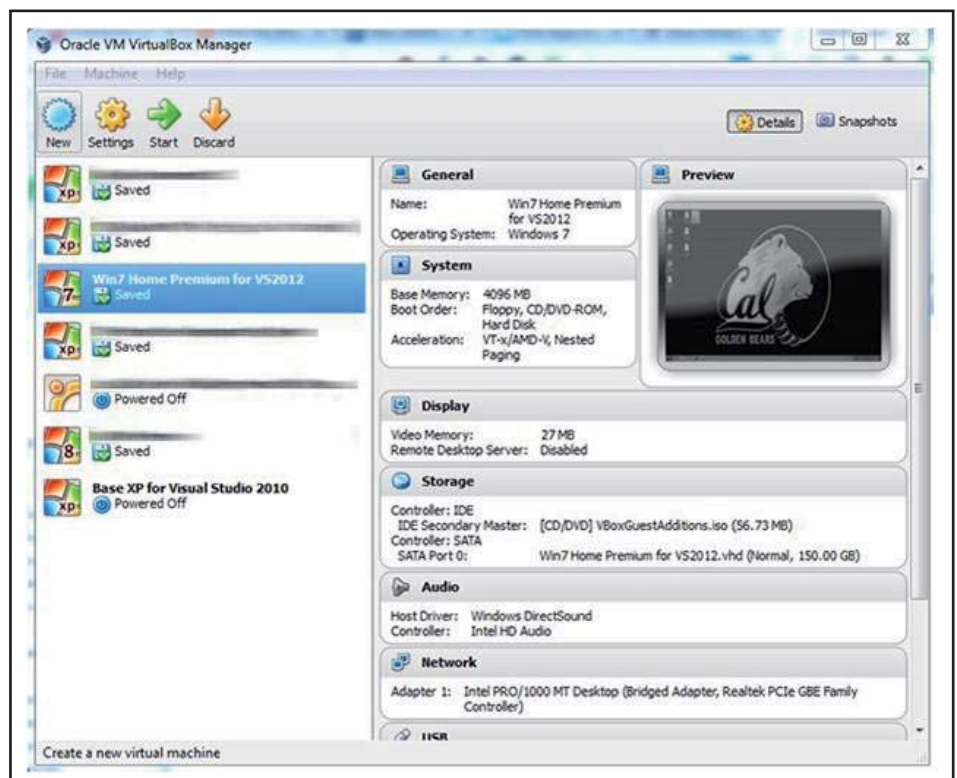


Abbildung 8: Virtualisierung auf Tablets

Quelle: Gewirtz

ben des Herstellers von vielen Mühen der BYOD-Problematik befreit. Ergebnis ist ein virtueller Windows Desktop. Das Produkt wird neuerdings auch von Cisco vertrieben. Betrachtet man aber Endgeräte wie das iPad 5 oder das Surface PRO, ist es völliger Hohn, diese leistungsfähigen Geräte durch Aufspielen eines virtuellen Desktops zu kastrieren. Der funktioniert zwar gut, bleibt in seinen Möglichkeiten aber weit hinter dem zurück, was z.B. Office 365 bietet.

Die Tendenz geht ganz klar in Richtung

leistungsfähigerer Geräte. Sie sind auch vom Preis her ausgesprochen attraktiv und es gibt hier wirklich keinen Anlass, hier aus wirtschaftlichen Erwägungen tiefer zu zielen. Apple hat aktuell übrigens etwas Mühe mit der Integration. Ein Surface PRO ist ein Windows-Gerät und alles, was in einem Unternehmen mit Windows in der Vergangenheit mit Windows gemacht wurde, kann ohne Weiteres auch auf das Surface übertragen werden. Apple hat an vielen Stellen eigene Entwicklungen, die vielleicht für die Geräte als solche genial sind, aber die Integration erschwe-

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

ren. In Deutschland müssen verschiedene Steuererklärungen ausnahmslos elektronisch abgegeben werden, auch von winzigen Selbständigen wie dem Autor. Die Finanzbehörden unterstützen aber keine ELSTER-Version für MAC-Betriebssysteme. Die Windows-Emulation auf dem Mac knirscht mit ELSTER aber auch und selbst das Portal ELSTER-Online sagt ganz deutlich, dass man damit rechnen muss, dass es nicht richtig funktioniert, wenn man einen Mac verwendet. Das ist nur ein konkretes Beispiel für den Nachholbedarf von Apple, einen umfangreichen Vergleich findet man in /5/. Es könnte sich die Tendenz entwickeln, dass interessierte Hersteller „BYOD-fähige“ Endgeräte herstellen, die man zwar unter den Weihnachtsbaum legen und privat benutzen kann, die aber direkt umfangreich für eine Integration in übliche IT-Strukturen von Unternehmen ausgerüstet sind. Wegen der vielfältigen Sicherheitsprobleme, die durch die vielen verschiedenen Android-Versionen entstehen, wie es eine aktuelle Studie von IBM X-Force® gezeigt hat, könnte es durchaus sein, dass Windows-basierte Tablets hier in eine echte Marktlücke stoßen. Es wird nicht ausbleiben, dass sich der Tablet Markt in gewisser Weise in einfache Geräte, die nicht oder kaum für die Einbindung in die Unternehmens-IT geeignet sind und bessere Modelle, die genügend Leistung und die Möglichkeit der Virtualisierung haben, bekannte Unternehmens-Standards unterstützen und sich auch sinnvoll in ein Sicherheitskonzept einbinden lassen.

Zusammenfassend kann ich mich dieser These von Gartner nicht wirklich anschließen. Seit dem Untergang der grünen SNA-Terminals hat es immer wieder Diskussionen um „Thin Clients“ gegeben, die so aussahen, dass der Client eben recht simpel sein sollte und die Leistung aus dem Netz erbracht wird. Alle diese Initiativen mit „Netz-Computern“ und ähnlichen Geräten haben nicht gefruchtet. Das liegt einfach daran, dass PCs oder sonstige „Fat Clients“ den Gesetzen des allgemeinen Massenmarktes unterlegen haben und bei immer höherem Leistungsumfang immer preiswerter wurden. Die Kastration der Endgeräte und die damit notwendige Aufrüstung der Server hat dazu geführt, dass schon bei der Vorbereitung einer Installation der angebliche Kostenvorteil gekippt ist, weil Server eben immer vergleichsweise teuer sind. Und es ist überhaupt nicht einzusehen, wieso das dieses Mal anders verlaufen sollte. Komplexer werdende Anforderungen der Benutzer stellen im Zusammenhang mit einer erhöhten Leistungsfähigkeit der Endgeräte vielleicht höhere Anforderungen an die Datenspeicherung in einer Cloud, aber

nicht notwendigerweise an die eigentliche Server-Rechenleistung. Man darf auch nie vergessen, dass die ja eindeutig gewünschte Mobilität zur Notwendigkeit drahtloser Verbindungen führt. Auch mit LTE werden diese zu bezahlbaren Kosten immer wesentlich weniger Leistung haben als vom Anwendungsumfeld her vergleichbare fest verdrahtete Anschlüsse. Also ist ein wesentliches Optimierungsziel auch die Minimierung der für die sinnvolle Zusammenarbeit von Cloud-Servern und leistungsfähigen mobilen Endgeräten benötigten Bandbreite.

6. Die Ära der „Personal Cloud“

Die Personal Cloud markiert das Umdenken weg von Geräten hin zu Diensten. In dieser neuen Welt werden die spezifischen Eigenschaften von Geräten für ein Unternehmen immer weniger wichtig, obwohl die Geräte natürlich nach wie vor nötig sein werden. Benutzer verwenden eine ganze Sammlung von Geräten, der PC ist nur eine der Möglichkeiten. Kein Gerät wird zur primären Schaltstelle. Diese Rolle wird eher der Personal Cloud zufallen. Anstatt dass man sich auf ein Gerät konzentriert, werden der Zugriff zur Cloud und der Inhalt, der dort gespeichert oder aus der Cloud geteilt wird, verwaltet und gesichert.

Was den privaten Bereich betrifft, ist diese Prognose schon längst von der Realität eingeholt worden. Ein tägliches praktisches Beispiel ist der Umgang mit TV-Inhalten. Kunden von Entertain oder Maxdome können nicht nur Inhalte praktisch beliebig auf grade passende End-

geräte vom Smartphone bis zum Internet-fähigen Großbild-TV mit vierfacher HD-Auflösung verteilen, sondern auch Sonderfunktionen, z.B. bei Entertain bis hin zur Sprachsteuerung, nutzen. Wenn mir irgendwann irgendwo irgendwie einfällt, dass ich noch eine Sendung aufnehmen möchte, kann ich sie über das Smartphone programmieren. Noch simpler sind die Angebote von Inhaltsanbietern wie ran.de für den Sport. Auch hier kann man, Netzabdeckung vorausgesetzt, überall mit einem Gerät, das man grade bei sich hat, fernsehen. Für die meisten Smartphone-Nutzer fallen auch eMail und Internet-Zugang in die Kategorie der überall unabhängig vom Endgerät leicht nutzbaren Dienste. Extrem praktisch sind dabei die Synchronisationsfunktionen. Basis für all dies sind Cloud-Angebote von Providern (z.B. Telekom-Cloud) oder Herstellern (z.B. iCloud). Die als Beispiel genannten Fälle bieten auch Funktionen für viele andere Arten von Daten eines individuellen Benutzers, z.B. Fotos oder Videos. Die technischen Möglichkeiten sind also offensichtlich alle vorhanden, das Problem, was Unternehmen eigentlich eher haben, ist, dass sie nicht wissen, wie sie damit richtig umgehen sollen. Auf der Provider-Seite fehlen noch Dinge wie Haftpflicht-Versicherungen für Schäden durch verloren gegangene Daten, aber eigentlich sind die Probleme mehr organisatorisch als technisch.

7. Software Defined Anything

Software Defined Anything SDx ist ein kollektiver Begriff, der die wachsende Marktbedeutung für verbesserte Standards zur

Kongress

ComConsult Netzwerk Forum 2014 24.03. - 26.03.14 in Königswinter

Traditionell ist die Dienstneutralität das Hauptmerkmal aller Netzwerke. Alle bisherigen Versuche, eine Applikations-spezifische Ende-zu-Ende Kontrolle und Garantie einzuführen sind in den letzten 20 Jahren gescheitert. Auch Prioritäten und Bandbreitenreservierungen für Verkehrsklassen sind trotz diverser Standards umstritten. Nun haben sich Netzwerke auf der Server-Seite in den letzten 5 Jahren deutlich verändert. Hier dominieren inzwischen die Software-Ports in den virtuellen Switches gegenüber den Hardware-Ports. Aus Sicht der Applikationen und virtuellen Maschinen spielt sich das Netzwerk mehr und mehr in Software ab. Verbindet man das mit dem zunehmenden Bedarf nach „One-Touch-Installationen“ und der schnellen Inbetriebnahme auch komplexer Applikationen innerhalb von Minuten, dann hat sich in der Server-Welt der Netzwerk-Schwerpunkt von der Hardware in die Software verschoben.

Moderation: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr.-Ing. Behrooz Moayeri
Preise: € 2.390,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

Programmierbarkeit der Infrastruktur und der Interoperabilität von Rechenzentren umfasst, die von dem Wunsch nach Automation, die kennzeichnend für Cloud Computing ist, und schneller Provisionierung von Infrastruktur genährt wird. Als Kollektiv umfasst SDx verschiedene Initiativen wie OpenStack, OpenFlow, das Open Compute Project und Open Rack, die vergleichbare Visionen haben. In dem Maße, wie sich individuelle SDx-Technologie-Silos entwickeln und Konsortien gründen, muss man nach kommenden Standards Ausschau halten, die Brücken zwischen den Möglichkeiten bauen und gleichzeitig individuelle Technologie-Anbieter dazu auffordern, ihr Commitment zu echten Interoperabilitäts-Standards innerhalb ihrer spezifischen technologischen Domänen zu demonstrieren. Während Offenheit immer ein Ziel sein wird, was man von den Anbietern einfordert, können unterschiedliche Interpretationen von SDx Definitionen alles andere als offen sein. Anbieter von SDN- (Netz), SDDC- (Rechenzentrum), SDS- (Speicher) und SDI- (Infrastruktur) Technologien versuchen alle, in ihren jeweiligen Bereichen die Führung zu halten, während die Entwicklung von SDx-Initiativen eigentlich Vielfalt und Ausweitung von Marktchancen bedeutet. Daher werden Anbieter, die einen Infrastrukturbereich dominieren, sehr zurückhaltend mit der Implementierung von Standards sein, die das Potential haben, die Margen zu senken und dem Wettbewerb Türen zu öffnen, auch wenn der Verbraucher durch Vereinfachung, Kostensenkung und Effizienz bei der Konsolidierung profitieren würde.

Beschränken wir die Diskussion auf SDN. Es gibt sicherlich keinen Netzwerk-Profi, der geschafft hat, der Legion von Darstellungen zu diesem Thema zu entgehen und nicht wenigstens die wesentlichen technischen Aussagen mit zu bekommen. Selten oder gar nicht wird allerdings die Frage gestellt, für wen SDN als Konzept überhaupt gemacht wurde. Betrachten wir die Fakten. Die grundsätzlichen Konzepte kommen von Wissenschaftlern wie Nick McKeown, denen bei einer Technik vollkommen wurscht ist, wer davon konkret einen Nutzen haben könnte. Ein möglicher Nutzen wird wenn überhaupt theoretisch abstrahiert. Dann kam die ONF und erstaunlicher Weise sind die wesentlichen Gründungsmitglieder Provider. Diese hatten sich längst mit dem Problem befasst, wie man in Zukunft die immer weiter wachsenden Datenströme angemessen beherrschen könnte. Dieses Problem ist ja nicht auf LANs beschränkt. Und so kam es dazu, dass die ersten Berichte über erfolgreiches SDN von den Providern kamen, allen voran Google. Provider, ob nun Netz- oder Service-Provider oder

beides, machen sich schon seit Jahren immer weiter Gedanken über eigene Referenzarchitekturen. Diese Architekturen sind einerseits eine Vorgabe für Hersteller, die bei dem zu erwartenden Auftragsumfang ruhig einmal nach Spezifikation bauen können, andererseits erhofft man sich natürlich, dass die eigene Architektur besser ist als die der Mitbewerber. Manche Referenzarchitekturen sind geheim, andere nicht, am bekanntesten ist sicherlich die von Facebook. Eine Provider-Architektur bezieht sich auf alles, was man so benötigt, also primär Server, Speicher und Netzwerk-Komponenten. Übergreifender Sinn ist immer die Kostensenkung aus verschiedenen Richtungen. Ein großer Provider hat natürlich auch das notwendige Personal, um umfangreiche Programmieraufgaben durchführen zu können. Da macht es auch nichts, wenn man mal ein paar Tausend Flows von Hand definieren muss, wenn nicht grade schnell jemand dafür wiederum ein Programm schreibt. Das Netz ist derart umringt von hoch ausgebildetem Personal, dass ihm gar nichts anderes übrig bleibt, als ordentlich zu arbeiten. Das sind synthetische Umgebungen, die mit der richtigen Welt und den praktischen Problemen dort rein nichts gemein haben. Was ist aber jetzt in den letzten zwei Jahren passiert? Die offenen Spezifikationen von der ONF sind nicht wirklich weiter gekommen, viele Dinge sind einfach überhaupt noch nicht ausgereift oder noch nicht einmal angepackt, die Hersteller überschlagen sich dennoch mit SDN-Ankündigungen über die auf SDN basierenden Anwendungen, die ja letztlich angeblich so vorteilhaft sein sollen, sind noch nicht einmal am Horizont als kleine Schatten zu sehen. Wie ist die Diskussion um VMware NSX hier einzuordnen? Ganz einfach, es ist natürlich der Traum von VMware, dass Provider ausschließlich deren Virtualisierungssoftware nutzen. Technisch ist klar, dass NSX nur dann sinnvoll ist, wenn es eine entsprechende Hardware-Unterstützung gibt. In meiner Reihe „Chip, Chip, Hurra“ habe ich ja entsprechende Switch-ASICs vorgestellt, die die Tunnel direkt in Hardware implementieren, sonst wird das Ganze zu langsam. VMware könnte hier Glück bei Service Providern haben, reine Netzwerk-Provider werden sich nicht von einer einzelnen Virtualisierungssoftware abhängig machen. Tatsache ist, dass VMware die Idee der Tunnel keineswegs erfunden hat. Jedes Provider-Netz arbeitet im Grunde mit Tunneln und es gibt enorm viele Standards. Die Ansicht, dass das bei Cisco oder anderen Netzwerkern wirklich Probleme erzeugen könnte, halte ich für völlig absurd. So viel kann eine Community überhaupt nicht vergessen haben. Was man aber aktuell bei Cisco sieht,

könnte einem als Größenwahn erscheinen. Der Nexus 9000 ist ein reinrassiger 40G-Switch, man kann ihn zwar auch als 10G-Switch einsetzen, aber das ist hinausgeworfenes Geld. Die Anwendungszentrische Infrastruktur ACI ist irgendwie eine Mischung aus Dingen, die man schon gehört hat und anderen, die sehr komplex erscheinen. Ob es nun am Ende 850 oder 1734 APIs gibt, kratzt den Betreiber des RZs eines durchschnittlichen mittelständischen Unternehmens oder einer Behörde reichlich wenig, weil er schlicht und ergreifend nichts Sinnvolles damit anfangen kann.

Dinge, für die man angeblich SDN benötigt, sind Automatisierung und Dynamisierung komplexerer virtueller Umgebungen. Wer aber hat solche Umgebungen und wer kann tatsächlich daraus einen Nutzen ziehen? Doch nur wirklich große Betreiber. Es gibt aktuell nichts in den SDN-Definitionen, was auf direktem Wege zu Vorteilen für einen „normalen“ Betreiber führt. Das war auch nie der Sinn von SDN. Aber ist SDN damit schon am Ende, wenn Mega-RZs das Konzept einsetzen? Der Master-Plan ist offensichtlich völlig anders. Man geht bei interessierten Kreisen davon aus, dass kleinere und mittlere Rechenzentren, wie sie heute noch in großen Mengen vorhanden sind, immer unwirtschaftlicher werden und relativ schnell einen Punkt erreichen, an dem sie die immer weiter steigenden Anforderungen mit eigener Kraft und eigenem Personal nicht mehr bewältigen können. Unternehmen, die in diesen Sektor fallen, sollen sich dann die benötigte Leistung bei Cloud-Service-Providern holen und für diese Leistung nach Umfang oder Verbrauch bezahlen. Das ist ein Markt von fast unvorstellbarer Größenordnung, denn er umfasst letztlich Zehntausende kleinerer und mittlerer RZs.

Um gegenüber dem privaten Betrieb eine Marge erzeugen zu können, die es erlaubt, die Leistung zu wettbewerbsfähigen Kosten anzubieten, müssen die CSPs natürlich hochgradig optimiert und automatisiert sein. Und in diesem Zusammenhang machen SDN-Konzepte, die neuen Angebote bestimmter Hersteller wie Cisco aber auch die Welt von VMware endlich wirklich Sinn.

Nun kann man ja sagen, dass das an und für sich keine wirklich fundamental neuen Gedanken sind, denn Begriffe wie „Outsourcing“ gibt es schon seit Jahrzehnten. Allerdings sind die Anforderungen an Rechenzentren in den letzten 5 Jahren stärker gestiegen als in den drei Jahrzehnten zuvor. Fast täglich prasseln auf einen Betreiber neue Vorschriften, Techniken, Verfahren, Leistungsanforderungen

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

gen und weitere Problemstellungen ein, die er möglichst ohne zusätzlichen finanziellen oder personellen Aufwand „gestern“ gelöst haben soll. Natürlich ohne merkliche Unterbrechung des Betriebs. Wann können Angebote von CSPs besonders erfolgreich werden? Naja, genau dann, wenn sie in einer verfahrenen Situation schnell dringende Probleme lösen. Je schneller und besser das funktioniert, desto großzügiger wird auf die Kosten geblickt.

Ich hatte die Diskussion hier auf SDN beschränkt, für die anderen Bereiche des „SDx“ gelten analoge Überlegungen. Mit dem nächsten Top-Trend von Gartner können wir die Diskussion unmittelbar weiterführen.

8. Web-Scale IT

Web-Scale IT ist ein Denkmuster für globales Computing, welches die Möglichkeiten großer Cloud Service Provider innerhalb einer unternehmenseigenen IT-Umgebung durch die Neubewertung von Positionen über verschiedene Bereiche hinweg einbringt. Große Cloud Service Provider wie Amazon, Google, Facebook usw. erfinden den Weg, in dem IT-Services erbracht werden können, neu. Ihre Möglichkeiten gehen weit über ein Maß hinweg, welches durch die schiere Größe charakterisiert wird, sondern umfassen auch Maßstäbe in Geschwindigkeit und Agilität. Möchten Unternehmen und Organisationen als private Betreiber von IT-Infrastrukturen mithalten, müssen sie Architekturen, Prozesse und Vorgehensweisen der exemplarischen Cloud Provider emulieren. Gartner nennt die Kombination all dieser Elemente „Web-scale IT“. Web-scale IT versucht, die IT-Wertschöpfungskette in systematischer Art und Weise zu ändern. Rechenzentren werden aus einer industriellen Ingenieurs-Perspektive entworfen, die jede Möglichkeit sucht, Kosten und Verschwendung zu minimieren. Das geht weit darüber hinaus, Rechenzentren lediglich Energie-effizient zu machen, sondern umfasst auch das In-house-Design von Schlüsselkomponenten wie Server, Speicher und Netze. Web-orientierte Architekturen erlauben den Entwicklern, sehr flexible und widerstandsfähige Systeme zu bauen, die sich auch nach Fehlern sehr schnell erholen.

Die Probleme, vor denen unternehmenseigene Rechenzentren heute stehen, hatten Provider schon vor Jahren. Dabei war ihre Ausgangslage wesentlich ungünstiger, weil ja vielfach noch gar keine Erfahrungen mit dem Einsatz neuer Technologien in großem Maßstab vorlagen. Aber wie wir wissen, haben die Provider und Betreiber wirklich großer Plattformen wie

z.B. Amazon, eBay & Co. die anstehenden Aufgaben erfolgreich gelöst. Rechenzentren, die dazu benutzt werden, Cloud-Services zu schaffen, stellen zunächst einmal ein signifikantes Investment dar, sowohl hinsichtlich der grundsätzlichen Anschaffungskosten als auch im Hinblick auf die Folgekosten. Daher ist die wichtigste Anforderung: Optimierung der tatsächlich erledigten Arbeit pro investiertem Dollar/Euro. Leider werden die Ressourcen in einem RZ aber oft nur zu einem geringen Grad ausgenutzt, bei Servern primär wegen schlichter Arbeitslosigkeit, bei Speichern eher wegen zu hoher Fragmentierung. Um dieses Problem anzupacken, kann man die Agilität zugrunde liegender RZ-Netze verbessern und Anreize für die optimierte Nutzung von Ressourcen schaffen. Anbieter von Cloud-Services neigen dazu, geografisch verteilte Netze von Rechenzentren aufzubauen. Diese geografische Verteilung verringert die Latenz auf dem Weg zu den Benutzern und erhöht die Zuverlässigkeit in dem Sinne, dass der Ausfall einer einzelnen Site nicht dazu führen muss, dass der Service ganz wegfällt. Ohne geeignetes Design und Management können diese räumlich verteilten Rechenzentren die Kosten für die Bereitstellung eines Services aber noch beträchtlich in die Höhe treiben. Außerdem müssen räumlich verteilt bereitgestellte Services auch dafür designet sein, wenn man einen Nutzen von ihnen haben möchte. Dazu müssen Netz- und RZ-Ressourcen gemeinsam optimiert werden. Das geht mit der Einführung neuer Systeme und Mechanismen für die Verteiltheit einher.

In reinen Cloud-DCs ist die Automatisierung eine zwingende Anforderung für die Skalierbarkeit und daher ein fundamentales Design-Prinzip. In einem gut geführten Cloud-DC ist das typische Verhältnis von Mitgliedern der Betriebsmannschaft zu Servern 1:1000. Automatisierte, auf Recovery angelegte Rechentechniken können mit der überwiegenden Zahl möglicherweise auftretender Probleme gut umgehen.

Aber: es gibt weitere gravierende Unterschiede zwischen privaten RZs in Unternehmen und Organisationen und Cloud-DCs. Die schiere Größe von Cloud-DCs, die auch durchaus eine sechsstellige Anzahl von Servern umfassen kann, ermöglicht es z.B. trotz erheblich größerer Anlaufkosten, für bestimmte Dienste eine Wirtschaftlichkeit herbeizuführen, die von Corporate-DCs niemals erzielt werden könnte. Economy of scale funktioniert im normalen Wirtschaftsleben meist prächtig und ist eine Triebfeder für die in den letzten Jahren immer stärker werdenden Konzentrationsprozesse auf praktisch allen Bereichen.

Wenn man es richtig anpackt, können private RZs durchaus von den Entwicklungen bei Cloud-DCs profitieren. Sehr pragmatisch ist der Ansatz, erfolgreiche Technologien und Geschäftsmodelle der Cloud-DCs zu adaptieren und sie dabei auf die Größenordnung eines privaten RZs herunter zu brechen. Nach aktuellem Stand ist das aber mit sehr individueller Planung verbunden und es gibt nur sehr

Kongress

ComConsult Netzwerk Forum 2014 24.03. - 26.03.14 in Königswinter

- Dienstneutralität kontra Applikations-Integration: Kampf um die Steuerung
- Netzwerk-Design in Zeiten von 10/40/100: wie sieht die Zukunft aus?
- Mobile Endgeräte im Netzwerk: teuer, unsicher, unbeherrschbar?

Das ComConsult Netzwerk Forum 2014 stellt sich den herausragenden Fragen der Entwicklung der Unternehmensnetzwerke. Nach 10 Jahren des Konzept-Stillstands ist das traditionelle Netzwerk unter Druck sowohl aus der Server- als auch aus der Endgeräte-Welt. Kombiniert man dies mit dem notwendigen Bandbreitenwechsel auf 10/40/100, dann steht das Unternehmensnetzwerk der Zukunft auf dem Prüfstand.

Wir analysieren für Sie wohin die Netzwerk-Entwicklung geht, wer die Kontrahenten sind, welche Vor- und Nachteile die verschiedenen Alternativen für Ihr Unternehmen haben und wie relevant einzelne Technologien sind. Investieren Sie zukunftsicher in Ihre Netzwerke, das ComConsult Netzwerk Forum 2014 unterstützt Sie dabei.

Moderation: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr.-Ing. Behrooz Moayeri
Preise: € 2.390,- netto



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

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

wenige allgemeingültige technologische Wahrheiten. Kritisch ist hier das Inhouse-Design, das normale Betreiber nicht leisten können. Grade hier stecken aber sehr viele interessante Möglichkeiten.

Eine ausführliche Analyse der Arbeitsweise von Web-Scale-RZs, die hierbei auftretenden Probleme und Lösungswege und Ansätze zur möglichen Übertragung auf privat betriebene Rechenzentren findet der Interessent in /6/.

Die letzten zwei Top-Trends nach Gartner sollen hier nur genannt, aber nicht weiter diskutiert werden, weil sie mit Aspekten der infrastrukturellen Entwicklung der IT in Unternehmen und Organisationen zu nächst nur wenig zu tun haben.

9. Smart Machines

Bis 2020 wird die Ära der „Smart Machines“ mit einer erheblichen Verbreitung von kontextsensitiven, intelligenten persönlichen Assistenten, schlaun Ratgebern, fortgeschrittenen globalen Industriesystemen und der öffentlichen Verfügbarkeit früher Beispiele autonomer Fahrzeuge aufblühen. Die Smart Machine Ära wird die durchschlagendsten Änderungen in der Geschichte der IT haben. Neue Systeme werden einige der frühesten Visionen dessen, was Informations-Technologie leisten kann, erfüllen – indem sie Dinge tun, von denen wir bisher dachten, dass nur Menschen sie tun könnten und keinesfalls Maschinen. Gartner erwartet, dass Personen solche Geräte kaufen, kontrollieren und einsetzen werden, um erfolgreicher zu sein. Aus dem gleichen Grund werden auch Unternehmen in Smart Machines investieren. Spannungen zwischen Kundenanpassung und zentraler Kontrolle werden im Rahmen der Umwälzungen durch die Smart Machines kaum stattfinden. Wenn überhaupt werden Smart Machines die Kundenanpassung stärken, wenn die erste Welle der Käufe durch Unternehmen vorüber ist.

10. 3-D Drucker

Man erwartet, dass die weltweiten Auslieferungen von 3D-Druckern in 2014 um 75% wachsen und sich dann in 2015 nochmals verdoppeln werden. Während es ja schon seit 20 Jahren sehr teure Geräte in der Produktion gibt, sprechen wir jetzt über einen Preisrahmen zwischen 500 und 50.000 US\$. Der Hype im Consumer Markt hat Unternehmen darauf aufmerksam gemacht, dass 3D-Drucken real und verfügbar sowie ein gangbarer und kosten-effektiver Weg ist, Kosten durch verbesserte Designs, verschlanktes Prototyping und Herstellung kleiner Mengen ist.

Konsequenzen für private IT-Infrastrukturen

Sehr auffällig ist, dass die ersten acht Top Trends inhaltlich zusammenhängen. Das hat es in diesem Umfang bei den Prognosen von Gartner noch nicht gegeben. Jeder Trend hat das Potential, die IT erheblich zu beeinflussen. Alle zusammen haben daher das Potential, die bisherigen Vorstellungen von unternehmensweiter IT weitest gehend umzukrempeln. Natürlich trifft nicht jeder Trend alle Unternehmen in gleichem Maße, aber angesichts der massiven thematischen Ballung kann man durchaus den Gedanken vertreten, dass die nächsten Jahre für alle privaten RZs und die mit ihnen betriebene Informationsverarbeitung erhebliche technische und organisatorische Änderungen mit sich bringen. Sieht man genau hin, sind alle acht primären Top-Trends auslegungsfähig. Es gibt vielfach mehrere unterschiedliche technologische Methoden zur Unterlegung des Trends. Das macht es sehr schwierig, zu allgemein gültigen Aussagen zu kommen.

Dennoch gibt es einige wenige Dinge, die recht verlässlich eintreten werden oder bereits eingetreten sind, wie z.B. die fulminante Weiterentwicklung bei den mobilen Endgeräten oder die private Cloud.

Für mich persönlich haben die Trend-Aussagen aber eine Reihe von Konsequenzen für die Bewertung aktueller Diskussionen:

- **Dienst-neutrales oder Anwendungsbewusstes Netz.** Unter verschiedenen Begriffen diskutieren wir schon seit rund 15 Jahren über eine derartige Thematik. Die Diskussion lebt immer dann auf, wenn eigentlich eine neue Leistungsstufe bei der technischen Übertragung wünschenswert wäre, aber entweder noch zu teuer ist oder schlicht noch nicht in Serie produziert werden kann. Beim Übergang von 0,1 auf 1 GbE, von 1 GbE auf 10 GbE gab es Diskussionen um Priorisierungen und differenzierte Priorisierung. Jetzt stehen wir an der Schwelle des Übergangs von 10 GbE auf 40 oder 100 GbE und mit der Zuverlässigkeit des Eintretens von Weinhachten kommt jetzt die Diskussion um SOA (gab es schon) und ACI (ist jetzt neu von Cisco) wieder auf. Jede Art von Anwendungsbewusstem Netz macht letztlich nichts anderes als mehr oder minder fein detaillierte Priorisierung. Das erhöht aber die Komplexität. Es gibt ein abschreckendes Beispiel, nämlich Carrier Ethernet. Der Unterschied zwischen CE 1.0 und dem im letzten Jahr eingeführten CE 2.0 liegt primär in der Möglichkeit, differenzierte Qualitätsparameter über alle in Reihe und/oder

verschachtelt hintereinander genutzten Teilnetze durchzusetzen. So können Provider Kunden entsprechende Angebote machen. Der Aufwand ist indes fürchterlich. Aus dem eigentlich einfachen Carrier Ethernet wird ein komplexes Monstrum. Wer möchte, kann sich das auf der Website des Metro Ethernet Forums in vielen Präsentationen gerne selbst ansehen. Cisco hat natürlich als Hersteller ein hohes Interesse, den Kunden in diese Richtung zu führen, denn wenn er darauf eingeht, gibt es kein Zurück mehr und der Kunde ist ewig auf proprietäre Steuerungsmechanismen angewiesen.

- Die Diskussion über die **zukünftige Kontrolle im Netz** wurde ja dadurch angestoßen, dass VMware letztlich mit Nicira NVP eine Software anbietet, die ein Netz völlig verdeckt, es zu einem System von Tunneln macht und diese Tunnel für die Kommunikation zur Verfügung stellt. Damit möchte VMware die Kontrolle über das Netz erlangen. Das erinnert mich an die Ziegen. Seit sie ein Eichhörnchen gesehen haben, wollen sie einen buschigen Schwanz. Die technische Implementierung dieser Tunnel-systeme ist durchaus möglich, es gibt ja sogar neue Switch-ASICs, die das direkt in Hardware unterstützen. Man kann durchaus darüber diskutieren, aus einem Netz ein Tunnelsystem zu machen, das ist die grundsätzliche Vorgehensweise von Providern, weil sie auf diese Weise Traffic Engineering, Multi-Mandantenfähigkeit, technologische Flexibilität und Sicherheitsfunktionen durchsetzen. Aber bitte dann doch auf Basis internationaler Standards, die es ja für die Provider gibt. Angesichts der vielen beweglichen Ziele im Rahmen der Weiterentwicklung der IT von Unternehmen ist es wirklich eine völlige Schnapsidee, die Kontrolle über das Herzstück der Kommunikation, ohne die nichts funktioniert, an einen Hersteller von Virtualisierungssoftware zu übergeben, der in den vergangenen Jahren nicht gerade gezeigt hat, dass er die „Innereien“ von Netzwerken wirklich versteht, ich erinnere nur an die unseligen Diskussionen hinsichtlich der Implementierung von VM-Kommunikation und -Migration über Netze. VMware mag heute in verschiedenen Bereichen Marktführer sein, das kann sich aber schnell ändern, z.B. wenn Virtualisierungs- und Hypervisorfunktionen in breiterem Maße Bestandteil der ohnehin verwendeten Standard-Betriebssysteme werden, wie es bei verschiedenen LINUX-Versionen aber auch bei Windows 8 heute bereits der Fall ist. Jede Art von Entwicklung des Netzes muss sich auch weiterhin am

Die Top 10 Technologie Trends 2014 nach Gartner und ihre Bedeutung für private Informations-Infrastrukturen

bisher über 30 Jahre erfolgreichen Modell der Nutzung von Standards an allen Stellen, wo dies möglich ist, orientieren.

- **Auslagerung von Funktionen.** Jedes private RZ wird seine Leistungsfähigkeit auf den Prüfstand stellen lassen müssen. Dabei kann es durchaus vorkommen, dass die Kosten für bestimmte oder alle Leistungen trotz technischer Optimierungen zu hoch sind. In diesem Fall müssen Funktionen ganz oder teilweise an einen Provider ausgelagert werden. Die Erfahrung mit Outsourcing-Tendenzen in den letzten Jahrzehnten zeigt jedoch immer wieder, dass am Ende der Rotstift nicht das ausschlaggebende Instrument ist, sondern andere unternehmenspolitische Entscheidungen wichtiger sind. Das ist aber ein sehr individuelles Thema.

- **Neubewertung von BYOD.** Lässt man unbedachte Benutzer mit ungeeigneten Endgeräten auf die Unternehmens-IT los, wird BYOD zum multidimensionalen apokalyptischen Horror-Szenario in den Bereichen Funktionalität, Organisation, Management, Integration mit bestehenden Anwendungen, Datenschutz, Sicherheit und Kosten. Geht man mal von den Darstellungen interessierter Anbieter in ihrer Rolle als Jubel-Perser weg, sieht man, dass bestehende BYOD-Angebote heute vielfach nur dazu genutzt werden, eMails zu checken. Ein wirklich hoher Aufwand für das bisschen Nutzen! Ich habe keine Hoffnung, dass die Nutzer besser werden. Wesentlich verbessern wird sich aber die Situation bei den Endgeräten. Alle Geräte werden mehr Leistung haben, es wird Geräte geben, die systemisch für den Einsatz im Umfeld von Unternehmen gut geeignet sind, vielfach wird es Unterstützung für Virtualisierung geben und mit HTML5 auch mindestens eine Geräte-unabhängige Schnittstelle. All dies wird dazu führen, dass nicht nur die oftmals versprochenen Produktivitäts-Vorteile von BYOD-Konzepten eintreten, sondern darüber hinaus auch weitere Optimierungsziele in der IT realisiert werden können, wie z.B. eine wesentlich flexiblere Nutzung und Auslastung von Servern, Speichern und anderen Hardware-Ressourcen. Objekt-orientierte App-Kombinationen werden Standard-Software ersetzen, konventionelle Versorgungsbereiche mit festmontierten Dosen werden praktisch entfallen und unternehmensweit können flexiblere Arbeitsmodelle eingeführt werden. Auch wenn diese Entwicklungen vielleicht erst in den nächsten 12 bis 36 Monaten durchschlagen, muss man bereits heute damit beginnen, vernünftige skalierbare

Betriebskonzepte für Mobile Device Management und Sicherheit einzuführen und auf einer kleineren Anzahl von Geräten ausprobieren, damit sie dann zur Verfügung stehen, wenn sie wirklich benötigt werden. Parallel dazu muss man sich natürlich überlegen, wie man die bestehenden Anwendungen mit der nunmehr neu gestalteten Außenwelt sinnvoll verbindet. Desktop-Virtualisierung kann in diesem Zusammenhang nur eine Übergangs-Krücke sein.

Kurzum: es bleibt spannend!

Literatur

- /1/ Gellersen, Gaedke „Object Oriented Web Application Development“, als pdf unter

- www.srtechgroup.com/oo_web.pdf
- /2/ Dradjadina, Ray „Create Advanced Web Applications with Object-Oriented Techniques“, www.msdn.microsoft.com/en-us/magazine/cc163419.aspx
- /3/ Hoff S. „Das Internet of Things“, Netzwerk Insider Juni 13
- /4/ Höchel-Winter, C. „Wireless Tings“ Wissensportal www.comconsult-research.de
- /5/ Gewirtz, David „Could I use a 64-bit iPad 5 or a Surface 2 as my main work computer?“ auf www.zdnet.com
- /6/ Kauffels, F.-J. Report: „RZ-Netzwerk-Infrastruktur Redesign“, 6. Auflage 2013 www.comconsult-research.de

Kongress

ComConsult Netzwerk Forum 2014 24.03. - 26.03.14 in Königswinter

- **Dienstneutralität kontra Applikations-Integration: Kampf um die Steuerung**
- **Netzwerk-Design in Zeiten von 10/40/100: wie sieht die Zukunft aus?**
- **Mobile Endgeräte im Netzwerk: teuer, unsicher, unbeherrschbar?**

Traditionell ist die Dienstneutralität das Hauptmerkmal aller Netzwerke. Alle bisherigen Versuche, eine Applikations-spezifische Ende-zu-Ende Kontrolle und Garantie einzuführen sind in den letzten 20 Jahren gescheitert. Auch Prioritäten und Bandbreitenreservierungen für Verkehrsklassen sind trotz diverser Standards umstritten. Nun haben sich Netzwerke auf der Server-Seite in den letzten 5 Jahren deutlich verändert. Hier dominieren inzwischen die Software-Ports in den virtuellen Switches gegenüber den Hardware-Ports. Aus Sicht der Applikationen und virtuellen Maschinen spielt sich das Netzwerk mehr und mehr in Software ab. Verbindet man das mit dem zunehmenden Bedarf nach „One-Touch-Installationen“ und der schnellen Inbetriebnahme auch komplexer Applikationen innerhalb von Minuten, dann hat sich in der Server-Welt der Netzwerk-Schwerpunkt von der Hardware in die Software verschoben.

Dies wird eine Reihe von Fragen auf:

- Liegt die Zukunft der Server-Vernetzung in der Software?
- Hat das dienstneutrale Netzwerk damit ausgedient?
- Wie kann eine automatische Netzwerk-Konfiguration bei der Installation von Servern erreicht werden?
- Wie kann der Bedarf dynamischer virtueller Umgebungen am besten erfüllt werden?
- Wer wird die Zukunft bestimmen: Anbieter wie VMware mit NSX oder die traditionellen Netzwerker mit Application Aware Networks?

Das ComConsult Netzwerk Forum 2014 stellt sich den herausragenden Fragen der Entwicklung der Unternehmensnetzwerke. Nach 10 Jahren des Konzept-Stillstands ist das traditionelle Netzwerk unter Druck sowohl aus der Server- als auch aus der Endgeräte-Welt. Kombiniert man dies mit dem notwendigen Bandbreitenwechsel auf 10/40/100, dann steht das Unternehmensnetzwerk der Zukunft auf dem Prüfstand.

Moderation: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr.-Ing. Behrooz Moayeri
Preise: € 2.390,- netto



Buchen Sie über unsere Web-Seite

www.comconsult-akademie.de

Aktuelle Veranstaltungen

Lokale Netze für Einsteiger, 27.01. - 31.01.14 in Aachen

Dieses Seminar vermittelt kompakt und intensiv innerhalb von 5 Tagen die Grundprinzipien des Aufbaus und der Arbeitsweise Lokaler Netzwerke. Dabei werden sowohl die notwendigen theoretischen Hintergrundkenntnisse vermittelt als auch der praktische Aufbau und der Betrieb eines LANs erläutert. Ausgehend von einer Darstellung von Themen der Verkabelung und der grundlegenden Übertragungsprotokolle werden die wichtigen Zusammenhänge zwischen der Arbeitsweise von Switch-Systemen, den darauf aufsetzenden Verfahren und der Anbindung von PCs und Servern systematisch erklärt.

Preis: € 2.490,-- netto

IP-Wissen für TK-Mitarbeiter, 10.02. - 11.02.14 in Bonn

Dieses Seminar vermittelt kompakt und effizient das IP-Wissen, das TK-Mitarbeiter ohne Vorkenntnisse zur Planung und zum Betrieb von IP-basierten Telefonie-Lösungen benötigen. Alle Seminarinhalte werden von einem Referenten mit hoher Praxiserfahrung betreut. Ziel ist dabei bewusst, statt einer umfassenden Theorieschulung gezielt die Aspekte vorzustellen und unter Praxis-relevanten Gesichtspunkten zu beleuchten, die erfahrungsgemäß aus Sicht einer IP-basierten Telefonielösung wichtig sind.

Preis: € 1.590,-- netto

Rechenzentrumsdesign - Technologien neuester Stand, 10.02. - 12.02.14 in Bonn

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

Wireless LAN professionell, 17.02. - 19.02.14 in Bonn

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, Markt- und Produktsituation werden bewertet, und Empfehlungen für eine optimale Auswahl werden gegeben.

Preis: € 1.890,-- netto

IPv6 Grundlagen - SeminarPlus, 24.02. - 25.02.14 in Düsseldorf

Diese Schulung bietet allen Planern, Betreibern, Administratoren und Softwareentwicklern, die von einer Migration zu IPv6 betroffen sind, ein tiefes Verständnis der Basis von IPv6. Das neue Konzept „SeminarPlus“ sorgt mit vielen Videos zur individuellen Vor- und Nachbereitung der Präsenz-Schulung für einen optimalen Lernerfolg.

Preis: € 1.790,-- netto

IP-Telefonie und Unified Communications erfolgreich planen und umsetzen, 24.02. - 26.02.14 in Düsseldorf

Dieses Seminar behandelt die Projektschritte, Einsatz- und Migrations-Szenarien, einsetzbare Basis-Technologien, Komponenten und erweiterte TK-Anwendungen, Bewertungskriterien für eine TK-Lösung und gibt eine Übersicht über den bestehenden TK-Markt etablierter Hersteller wie Alcatel-Lucent, Avaya, Cisco und Siemens aber auch des Newcomers Microsoft.

Preis: € 1.890,-- netto

SIP (Session Initiation Protocol) - Basis-Technologie der IP-Telefonie, 10.03. - 12.03.14 in Stuttgart

SIP ist der Schlüssel für offene, leistungsfähige und kostenoptimale Kommunikationslösungen. Es umfasst Sprache, Video, Daten und Präsenz. Lernen Sie, was SIP leistet, worin sich wesentliche Herstellerlösungen unterscheiden und wie Sie das Beste aus beiden Welten zukunftsorientiert nach Ihrem Bedarf optimieren.

Preis: € 1.890,-- netto

Virtualisierungstechnologien in der Analyse, 10.03. - 12.03.14 in Stuttgart

Dieses 3-tägige Seminar liefert einen umfassenden und zugleich detaillierten Einblick in die aktuellen Virtualisierungstechnologien der marktführenden Anbieter. Vom Server über das Netzwerk bis zum Speicher und schließlich auch zum Client werden die Möglichkeiten und Grenzen der Virtualisierungslösungen analysiert. Dabei bleiben auch Sicherheitsaspekte nicht unberücksichtigt. Basis hierfür bilden neben den technischen Grundlagen und Hintergründe die Erfahrungen aus dem Projektalltag sowie die Diskussion mit den Teilnehmern.

Preis: € 1.890,-- netto

Sicherheitsmanagement mit BSI-Grundschriftmethodik/ ISO 27001, 10.03. - 12.03.14 in Stuttgart

Angemessene Sicherheit mit optimalem Aufwand: geht das? Die Antwort liegt in der Nutzung bewährter Architekturen und Lösungen bei gleichzeitiger Erfüllung von Compliance Richtlinien. Anders formuliert: das Rad muss nicht von jedem Unternehmen neu erfunden werden. Dieses Seminar stellt die BSI-Grundschriftmethodik vor und zeigt auf, wie auf der Basis dieser etwas abstrakten Methodik eine Praxis-gerechte Sicherheitslösung mit optimalem Aufwand erreicht werden kann.

Preis: € 1.890,-- netto

TCP/IP-Netze erfolgreich betreiben, 10.03. - 12.03.14 in Aachen

IP ist die Grundlage jeden Netzes. Die Protokolle TCP und UDP bilden die Basis jeder Anwendungskommunikation. Für den erfolgreichen Betrieb werden zudem Kenntnisse über Routingprotokolle, DHCP, DNS benötigt. Dieses Seminar vermittelt praxisnah das notwendige Wissen.

Preis: € 1.890,-- netto

Zertifizierungen

ComConsult Certified Network Engineer

Lokale Netze

27.01. - 31.01.14 in Aachen
05.05. - 09.05.14 in Aachen
08.09. - 12.09.14 in Aachen
17.11. - 21.11.14 in Aachen

TCP/IP-Netze erfolgreich betreiben

10.03. - 12.03.14 in Aachen
02.06. - 04.06.14 in Aachen
20.10. - 22.10.14 in Aachen

Internetworking

07.04. - 11.04.14 in Aachen
23.06. - 27.06.14 in Aachen
03.11. - 07.11.14 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

25.02. - 28.02.14 in Aachen
20.05. - 23.05.14 in Aachen
28.10. - 31.10.14 in Aachen

Trouble Shooting für Netzwerk-Anwendungen

18.03. - 21.03.14 in Aachen
24.06. - 27.06.14 in Aachen
18.11. - 21.11.14 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

24.02. - 26.02.14 in Düsseldorf
05.05. - 07.05.14 in Bonn
29.09. - 01.10.14 in Stuttgart
17.11. - 19.11.14 in Bonn

Session Initiation Protocol Basis-Technologie der IP-Telefonie

10.03. - 12.03.14 in Stuttgart
02.06. - 04.06.14 in Düsseldorf
20.10. - 22.10.14 in Berlin

Umfassende Absicherung von Voice over IP und Unified Communications

17.03. - 18.03.14 in Bonn
30.06. - 01.07.14 in Bonn
03.11. - 04.11.14 in Bonn

Optionales Einsteiger-Seminar: IP-Wissen für TK-Mitarbeiter

10.02. - 11.02.14 in Bonn
08.04. - 09.04.14 in Aachen
15.09. - 16.09.14 in Düsseldorf

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: € 4.840,-- netto statt € 5.370,-- 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