

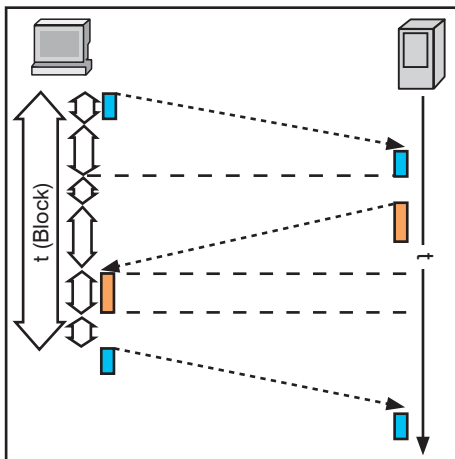
Schwerpunktthema

Wie viel Delay ist tolerierbar?

von Dr. Behrooz Moayeri

Seit einigen Monaten macht im Zusammenhang mit der Bewertung von LAN-Komponenten ein neuer Ausdruck die Runden: Ultra-Low Latency (ULL). Switch-Hersteller stufen ihre neuen Produkte gerne als ULL-Komponenten ein. Von Sub-Microsecond Latency ist die Rede. Nicht nur der einzelne Switch, sondern auch das Netzdesign soll auf ULL ausgerichtet werden.

Daher werben die Hersteller für vermaschte Konstrukte als Ablösung der klassischen Hierarchien im RZ. Auf der anderen Seite gibt es Veröffentlichungen zum Locator/ID Separation Protocol, das u. a. helfen soll, die Latenz beim Zugriff



auf Ressourcen in einer auf verschiedene Data-Center-Lokationen verteilten Cloud zu minimieren¹. Wie sind diese neuen Verfahren und Konzepte einzustufen? Wie viel Delay ist tolerierbar? Der vorliegende Beitrag geht diesen Fragen nach. Zunächst wird auf die Delay-Minimierung innerhalb der Rechenzentren eingegangen. Im zweiten Teil folgt die Betrachtung der Delay-Minimierung außerhalb der Rechenzentren, d. h. beim Zugriff der Clients auf Server.

weiter auf Seite 8

Zweitthema

Cloud Storage in der Analyse

von Dr. Jürgen Suppan

Wir alle haben unsere Wunschliste mit Produkten, die wir dringend oder weniger dringend brauchen. Cloud Storage stand vermutlich auf keiner dieser Listen und kam etwas überraschend.

Nachdem nun die ersten Kinderkrankheiten beseitigt sind und ein erster Reifegrad erreicht ist, analysieren wir für Sie: wo steht Cloud Storage im Verhältnis zu unserem Bedarf und wie schneidet es gegenüber unseren Top-Vier Kriterien ab!

Die zukünftige Handhabung von Speicher gehört mit zu den wichtigsten Fragen in der zukünftigen Gestaltung von IT-Architekturen.

weiter auf Seite 20

Neues Seminar

RZ-RZ-Kopplung - alles nur eine Frage der Bandbreite?

auf Seite 6

Geleit

Wie viel UC braucht die TK?

auf Seite 2

Aktueller Kongress

Rechenzentrum Infrastruktur-Redesign Forum 2011

ab Seite 4

Standpunkt

Auf die Prozesse kommt es an

auf Seite 19

Zum Geleit

Wie viel UC braucht die TK?

Wir kommen jetzt in die entscheidende Marktphase, in der viele der alten TK-Anlagen nun nicht länger wirtschaftlich betrieben werden können. Entsprechend explodiert die Zahl der Projekte. Und es gibt ein zentrales Problem, um das sich nahezu jedes Projekt dreht. Dies ist die Frage: brauchen wir Unified Communications und wenn ja, wie viel davon?

Diese Frage mag überraschen, diskutieren wir doch seit Jahren, dass die Zukunft des Marktes in der UC liegt. Entsprechend sahen wir lange Zeit Hersteller, die wichtige UC-Positionen nicht abdecken können, mit einem sehr kritischen Blick. Nun haben die Hersteller natürlich dazu gelernt. Die typischen „alten“ TK-Produkte wurden um UC-Funktionen angereichert, selbst vor der Virtualisierung der Lösungen schreckte man nicht zurück. Es gibt kein besseres Beispiel für diese Entwicklung als die HiPath 4000. Gefördert wird diese Entwicklung durch einen Markt mit vielen TK-lastigen Kleinberatern, die UC nicht zwingend fördern und einem wenig flexiblen Vertrieb der großen TK-Anbieter. So hat sich der Markt noch weiter in seine Glaubens-Lager zerlegt. Auf der einen Seite das Lager der klaren UC-Befürworter, dem diametral gegenüber die TK-Traditionalisten gegenüber stehen und dazwischen jede Menge Schattierungen.

Im Endeffekt fällt der Kunde die Entscheidung. Und da können wir nicht verkennen, dass dieser nach wie vor mehrere ganz schwerwiegende Probleme mit UC hat:

- Die Zahl der Arbeitsplätze, die UC-Funktionalität wirklich wirtschaftlich sinnvoll nutzen können, schwankt in den Unternehmen zwischen 10% und 90%. Vermutlich liegt der Mittelwert momentan deutlich unter 50%. Auch wenn sich das in Zukunft kontinuierlich in Richtung UC verschieben wird, existieren im jetzigen Zustand im Mittel noch viele reine TK-Arbeitsplätze.
- Viele der angebotenen UC-Lösungen sind schlicht so komplex, dass das angestrebte Ziel einer neuen Form von Kommunikation weder auf der Seite der Betreiber noch auf der Seite der Anwender wirtschaftlich erreicht werden kann. Auf der Betreiberseite ist der Betrieb zu komplex, auf der Anwenderseite werden viele UC-Funktionen nicht genutzt, weil sie nicht intuitiv genug in der Bedienung sind.



- Die Glaubwürdigkeit einiger Anbieter hat in den letzten Jahren deutlich gelitten. Zwar wird immer von UC und der Zukunft geredet, doch bleibt die tatsächliche Weiterentwicklung der Produkte hinter den Versprechungen zurück.
- Mit Irritation sehen die Kunden, dass die Hersteller wichtige Standards, die eigentlich untrennbar mit UC verbunden sind, blockieren. Dies gilt vor allem für SVC auf der Videoseite. Die heute angebotenen Video-Lösungen skalieren schlichtweg nicht mit großen Teilnehmer-Zahlen. Nimmt man die Vision von UC aber ernst, dann müssen sie das.

Wie glaubwürdig ist ein Hersteller, dem der Vertrieb seiner MCUs wichtiger ist als die Umsetzung einer skalierbaren UC-Plattform?

- Die Cloud wirft neue und entscheidende Fragezeichen auf. Zwar war der externe Betrieb der TK-Lösung immer mal wieder ein Thema, aber er konnte sich nie wirklich durchsetzen. Die Cloud, was immer wir darunter verstehen wollen, schreibt die Spielregeln neu aus. Damit stellt sich automatisch die Frage, ob in Zukunft ein Teil der Anschlüsse (oder alle?) besser durch einen externen Betreiber erbracht werden sollten. In der Vergangenheit war das zum Beispiel aufgrund der engen Vermaschung mit internen Schnittstellen nicht sinnvoll. Das Argument verliert aber an Bedeutung in genau dem Maße, in dem diese Schnittstellen auch in der Cloud verfügbar sind. Ich verweise auf die ComConsult Analyse zu „Public und Private Clouds“.

In dieser Lage suchen die Kunden nach einer einfachen Lösung, bei der man nichts falsch machen kann. Die Lösung soll auf der einen Seite alle Türen offen halten, auf der anderen Seite aber keine Bindungen an eigentlich ungeeignete Produkte generieren.

Damit sind wir wieder an einer Stelle, an der wir vor 2 Jahren schon einmal ge-



Abbildung 1: Video „Seminar: Public und Private Clouds in der Analyse“ analysiert die Vor- und Nachteile von Public und Private Clouds und leitet daraus Handlungsempfehlungen ab, Dr. Suppan auf ComConsult-Study.tv

Wie viel UC braucht die TK?

standen haben. Man nehme die traditionelle TK-Lösung und ergänze sie um eine angeflanschte UC-Lösung. Da ist es sehr praktisch, dass sich die alten TK-Lösungen in diesen 2 Jahren ja auch weiter entwickelt haben und als reine Software-Lösungen nun deutlich besser in den Rahmen eines gemeinschaftlichen Betriebs passen.

Damit ist aber auch klar: der Gewinner ist Microsoft. Lync-Server passt in dieses Muster wie die Faust aufs Auge. Da damit auch weitreichende Office- und Kollaborations-Anforderungen abgedeckt werden, hat man alle Türen für die Zukunft offen. Und als Krönung des Ganzen gibt es jetzt noch Office 365 mit der Option eines Cloud-Services mit dem identischen Funktionsumfang plus der Möglichkeit eines hybriden Betriebs (auch die anderen Hersteller gehen stark in diese Richtung, aber keiner belegt die Rolle des Ergänzungsprodukts so intensiv wie Microsoft).

Die Frage nach der Rolle des Lync-Servers in unserem aktuellen TK und UC-Markt war noch nie so drängend wie jetzt.

Finde ich das gut? Natürlich nicht. Nicht

weil ich etwas grundsätzliches gegen den Lync-Server hätte (im Gegenteil, das User Interface setzt Maßstäbe für die ganze Branche). Aber die Strategie der offenen Türen, so dass man nichts falsch machen kann und jederzeit so viel UC hat wie man braucht, ist schlichtweg sehr riskant. Nehmen wir einmal an, die Prognose, dass UC die Zukunft gehört, ist richtig. Was bedeutet das dann für diesen Ergänzungsbetrieb? Ganz einfach: der Kunde hat sich startend mit vielleicht 10% der Arbeitsplätze im Prinzip dazu bekannt, auf Dauer Lync-Server an alle Arbeitsplätze zu bringen. Ganz klar formuliert: die Entscheidung für einen Lync-Ergänzungsbetrieb ist eine langfristige Entscheidung für Microsoft. Das Problem ist nur, dass der Kunde das gar nicht weiß und so ggf. auch nicht will. Aber wo soll die Entwicklung des stufenweisen Ausbaus des UC-Bereichs denn anders enden?

Hinzu kommt, dass der parallele Betrieb zweier Produkte weder die Betriebskosten senkt noch die Nutzung für den Endanwender vereinfacht. Auf welcher Seite liegt denn zum Beispiel die Integration mobiler Endgeräte? Wie werden Rufweitschaltungen umgesetzt, wenn Anwen-

der in beiden Welten arbeiten? Und wie sieht die Video-Lösung auf Dauer aus?

Das Problem in dieser Situation ist, dass wir in keiner heilen Welt leben. Auf der einen Seite die Vielzahl der UC-Produkte, die auch nach Jahren der Entwicklung dem UC-Anspruch immer noch nicht gerecht wird. Auf der anderen Seite der Wunsch nach einem Risiko-armen Einstieg. Da gibt es keine klaren Gewinner.

Es ist an der Zeit, noch einmal die Strategien der wesentlichen Marktteilnehmer zu analysieren. Die Kunden, die jetzt aus ihrer alten TK-Anlage aussteigen wollen, brauchen klare und Risiko-freie Perspektiven.

Dies ist das große Ziel, das wir uns mit dem diesjährigen ComConsult UC-, Video und Kollaborations-Forum gesetzt haben. Wir werden sie über die aktuelle Entwicklung dieser Arbeiten auf dem Laufenden halten. Die Ergebnisse gibt es wie immer exklusiv für die Teilnehmer des Forums im November.

Ihr
Dr. Jürgen Suppan

Frühbucherphase noch bis zum 30.09.2011

Kongress

ComConsult Voice-, Video- und Kollaborations-Forum 2011 21.11. - 24.11.11 in Königswinter

Unified Communications integriert begrifflich viele verschiedene Formen von Kommunikation. Es beginnt mit der reinen Sprachkommunikation und endet bei der Kollaboration im IT-Sinne. Entsprechend schwer ist die Abgrenzung zu etablierten Produkten aus dieser Spannweite, sei es eine traditionelle TK-Anlage oder sei es ein Portalserver für Projekt-Kollaboration. Das ComConsult Voice-, Video und Kollaborations-Forum 2011 analysiert und diskutiert genau diese Fragen. Es basiert auf hochaktuellen Analysen von ComConsult-Research und wird kombiniert mit aktuellen Erfahrungen aus laufenden Projekten. Führende Hersteller werden nach ausgewählten Kriterien eingeladen, ihre Technologie- und Strategie-Sichtweise zu präsentieren und sich der Diskussion zu stellen.

Moderatoren: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr. Frank Imhoff
Preis: € 2.190,-* zzgl. MwSt. bzw. € 1.790,-* zzgl. MwSt. - *gültig bis zum 30.09.11



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

Kongress

Rechenzentrum Infrastruktur-Redesign Forum 2011

07. - 10.11.11 in Königswinter

Die ComConsult Akademie veranstaltet vom 07.11. - 10.11.11 ihren neuen Kongress „Rechenzentrum Infrastruktur-Redesign Forum 2011“ in Königswinter.

Das ComConsult Rechenzentrum Infrastruktur-Redesign Forum 2011 ist die zentrale Veranstaltung des Jahres, in der hochqualifizierte Referenten die neuen Einflussfaktoren, Möglichkeiten und Strategien vorstellen und mit den Herstellern diskutieren, um Anregungen und praktische Hinweise für die Gestaltung des Rechenzentrums der jetzt beginnenden Zukunft zu erarbeiten.

Noch nie hat es eine Zeit gegeben, in der so viele neue Anforderungen gleichzeitig auf die RZ-Netze zukommen. Web-Architekturen, Virtualisierung, I/O-Konsolidierung und Speicher-Konvergenz sind hier die wichtigsten Schlagworte. Durch sie wird das RZ-Netz zum Systembus. Aber was heißt das für Bandbreite, Latenz, Reaktionsfähigkeit, Sicherheit, Struktur und Betrieb? Es gibt viele neue Standards und Produkte, die alle in irgendeiner Weise zur Lösung beitragen. Aber wie genau? Welche Kombinationen sind sinnvoll, was ist eher überflüssig? Und schließlich: welche Systeme unterstützen diese ganzen Neuheiten? Bestehende Systeme stoßen hier schnell an Leistungsgrenzen und zwar nicht nur hinsichtlich der puren Bandbreite. Alle diese Anforderungen und möglichen Lösungen können nicht losgelöst voneinander betrachtet werden.

Die Ausgangssituation ist durch folgende Stichworte zu kennzeichnen:

- Leistungsexplosion Virtueller Server
- I/O-Konvergenz und Anschlussproblematik
- Anforderungen Virtueller Gesamtlösungen: das Netz als Systembus
- Integration neuer Storage-Technologien

Das ist die Perspektive eines Planers, der das Problem hat, neue, schnelle Rechner sinnvoll mit der Infrastruktur zu verbinden. Die Frage ist aber, ob das alleine zum Verständnis der Gesamt-Problematik ausreicht und zu befriedigenden Lösungen führt.

Aus der Perspektive der Systemarchitektur kann man vereinfachend sagen, dass die



Kommunikation im RZ von drei neuen Verkehrsströmen geprägt wird:

- Kommunikation zwischen virtuellen Maschinen als Teil von verteilten (Web-) Architekturen
- Systemkommunikation aus dem Umfeld der Virtualisierung wie z.B. das Wandern Virtueller Maschinen, High Availability und Fault Tolerance
- Verlagerung von Plattenspeicher aus dem Direct Attached Bereich hin zu Storage Area Networks

Einzig die erste Anforderung steht in unmittelbarem Zusammenhang zu Anwendungen und ihrer Bearbeitung. Die beiden anderen Anforderungen ergeben sich unmittelbar aus dem Konzentrationsprozess, der in praktisch allen RZs bereits läuft.

Kurz charakterisiert ist das Ziel des Konzentrationsprozesses die Ablösung bestehender Strukturen aus vielen singulären älteren Servern durch ein modernes virtuelles System mit wenigen Servern hohen Konzentrationsgrades. Auch wenn wir heute im Rahmen der Virtualisierung nur 10-20 „alte“ Server auf Virtuelle Maschinen in einem neuen Server abbilden, wird diese Zahl mit der Zeit getrieben durch die Entwicklung bei Prozessoren und Virtualisierungssystemen deutlich steigen.

Unabhängig davon, wie viele ältere Server nun in einem oder wenigen neuen Servern konzentriert werden, müssen natürlich sämtliche Funktionen und Betriebsmittel der alten Server nachgebildet werden, wenn die Migration auch für die Anwen-

dungen erfolgreich sein soll.

Die dadurch entstehenden Verkehrsströme sind neu - es gab sie vorher nicht.

Das zementiert eine unangenehme Tatsache:

In einem RZ entstehen durch neuartige Kommunikationsformen aus der Anwendungsebene und die durch die Virtualisierung bedingte Systemkommunikation neue Verkehrsströme erheblichen Umfangs, die durch eine Extrapolation bisheriger Verkehrsströme nicht vorherbestimmt werden können.

Das bedeutet, dass die herkömmliche Methode der Erweiterung eines RZ-Netzes auf der Grundlage dessen, was über die bisherigen Anforderungen für die Kommunikation der älteren singulären Server im Rahmen der Anforderungen durch die Anwendungen bekannt war, nicht mehr zielführend ist.

Was passiert in den nächsten Jahren? Knapp zusammengefasst:

- Web-Architekturen werden dominanter
- Virtualisierung erreicht ungeahnte Konzentrationen
- Die Virtualisierung wird sich auch auf Speicher erstrecken
- Es gibt neue reaktionsschnelle Speicher (SSDs), die integriert werden müssen
- 10 GbE wird Mindeststandard im RZ und im Campus
- 40/100 GbE ist die nächste Stufe
- Die heute diskutierten Standards für die Strukturierung werden umgesetzt
- Es entstehen völlig neue, latenzarme und hochskalierbare Architekturen.

Besonders spannend ist die Strukturproblematik. Möchte man ein RZ-Netz gleichermaßen skalierbar und latenzarm haben, kommt man unter dem Begriff „Matrix der kürzesten Wege“ zu völlig neuen Architekturen. Mittlerweile haben die Hersteller mindestens vier unterschiedliche Ansätze dafür vorgestellt, die alle einer intensiven Diskussion bedürfen, weil sie mit einer Ausnahme zwar zu extrem leistungsfähigen Netzen führen, die aber weitestgehend auf proprietären Techniken basieren:

- Fat Tree-Strukturen für Unified Fabrics
- Virtual Chassis

Rechenzentrum Infrastruktur-Redesign Forum 2011

- ScaleOut
- Single-Hop-Netze

Als wäre dies nicht schon genug Diskussionsstoff, wird das Jahr 2011 auch davon gekennzeichnet sein, dass die Chip-Hersteller völlig neue Switch-ASICs auf den Markt bringen, die unter Benutzung der 40 nm-Technologie die Vielfalt integrierter Funktionen (FCoE, DCB, TCP-Offload, iSCSI-Offload, ...) und die mögliche Portdichte und damit den Konzentrationsgrad erheblich erhöhen und dabei den Energieverbrauch und die allgemeinen Infrastrukturkosten wesentlich senken helfen werden. Diese neuen Chips werden noch in diesem Jahr zu Produkten führen, von denen wir vor zwei Jahren nicht zu träumen gewagt haben.

Man muss aber auch konstatieren, dass diese neuen Chipentwicklungen im Access-Bereich zu einer weitgehenden Austauschbarkeit der Komponenten der einzelnen Hersteller führen, was man spätestens dann sieht, wenn fast alle die gleichen Switch-ASIC Chipsätze benutzen. Um sich überhaupt noch abgrenzen zu können,

bauen die Hersteller unisono den Kern der neuen Netze mit proprietären Virtual Chassis auf und betreiben die Implementierung der an dieser Stelle geeigneten Standards nur zögerlich.

Das führt wiederum zu einer momentan sehr zögerlichen Akzeptanz der neuen Strategien durch den Markt. Nicht nur bei Netzen, sondern auch an anderen Stellen haben Betreiber lernen müssen, dass es ein gefährliches und teures Vergnügen sein kann, sich in die Hände eines Herstellers zu begeben und sich hinsichtlich sämtlicher Weiterentwicklungen von dessen selbstgebastelten Verfahren abhängig zu machen.

Eine professionell verantwortliche Planung für die Zukunft kann also nur in der durchgängigen Implementierung von Standards liegen. Provider machen das seit langem vor, sie akzeptieren nichts anderes. Es ist also offensichtlich möglich, vollends standardisierte große leistungsfähige Netze zu bauen und erfolgreich zu betreiben. Mit welcher Motivation sollte der Betreiber eines Corporate RZs weniger strenge Maß-

stäbe anlegen?

Also kann man sich eigentlich auf wenige Kernfragen konzentrieren:

1. Was benötigt das Corporate RZ in den nächsten Jahren wirklich?
2. Wie ist der Stand der Standardisierung für die benötigten Funktionen?
3. Inwieweit tragen die neuen Herstellerstrategien zu verantwortlichen, durchgängig standardisierten Lösungen für die wirklich benötigten Funktionen bei?

Besonders die erste Frage bedarf einer sehr differenzierten Betrachtung. In der allgemeinen Betrachtung scheint z.B. die Frage nach der Beherrschbarkeit etwas unter zu gehen. Mit immer stärkerem Konzentrationsgrad bei den Servern müssen wir auch mit immer mehr VMs rechnen. Was heute noch sehr übersichtlich sein kann, könnte in nur wenigen Jahren zu einer für den sicheren Betrieb gefährlichen Komplexität führen. Haben sich die Hersteller eigentlich dazu schon etwas überlegt?

Frühbucherphase bis zum 31.08.2011

Fax-Antwort an ComConsult 02408/955-399

Anmeldung Rechenzentrum Infrastruktur- Redesign Forum 2011

Ich buche den Kongress
**Rechenzentrum Infrastruktur-Redesign
Forum 2011**

mit Workshop am letzten Tag

☐ vom 07.11. - 10.11.11 in Königswinter
zum Preis € 2.190,--* zzgl. MwSt.

inklusive Report "RZ Netzwerk- Infrastruktur Redesign"

☐ zum Sonderpreis von nur € 338,--
zzgl. MwSt.

*gültig bis 31.08.2011

☐ Bitte reservieren Sie mir ein Zimmer

vom _____ bis _____ 11

ohne Workshop am letzten Tag

☐ vom 07.11. - 09.11.11 in Königswinter
zum Preis € 1.790,--* zzgl. MwSt.

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift



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

Neues Seminar

RZ-RZ-Kopplung - alles nur eine Frage der Bandbreite?

Die ComConsult Akademie veranstaltet am 26.09.11 erstmalig ihr Seminar „RZ-RZ-Kopplung - alles nur eine Frage der Bandbreite?“ in Köln.

Diese 1-tägige Veranstaltung konzentriert sich auf aktuelle Fragestellungen, die bei der Kopplung von modernen Rechenzentren entstehen. Neben der Motivationslage für eine solche Kopplung werden die physikalischen Rahmenbedingungen für ein solches Vorhaben aufgezeigt, um die natürlichen Grenzen hinsichtlich Latenz und Übertragungsrate greifbar zu machen. Anschließend werden aus den Bereichen Server, Storage und Netzwerk die aktuellen Cluster-Mechanismen sowie Techniken für Spiegelungs- und Replikationsverfahren im Kontext entfernter RZ-Standorte besprochen. Die Kombination aus strategischen Überlegungen und technischen Maßnahmen bilden abschließend eine Leitlinie, welche Systeme in welcher Form bei einer RZ-RZ-Kopplung zu berücksichtigen sind.

Das Seminar geht unter anderem auf folgende Fragen ein:

- BSI / Basel III: wo kommen die Anforderungen zur RZ-RZ-Kopplung her und wie sehen sie aus?
- Dark Fiber, CWDM, DWDM: welche Kopplungsmöglichkeiten gibt es auf der physikalischen Ebene? Wie funktionieren diese Technologien? Welche Distanzen können hiermit überbrückt werden? Mit



welchen Komponenten kann dies heute realisiert werden?

- Asynchrone / Synchrone Spiegelung: worin besteht der Unterschied? Welche Rolle spielt die physikalische Begrenzung durch die Lichtgeschwindigkeit? Welchen Einfluss hat das eingesetzte Übertragungsprotokoll?
- Wie sehen die in der Praxis ermittelten Datenübertragungsraten aus?
- Was bedeuten „Hochverfügbarkeit“ und „Fehlertoleranz“ auf Ebene der Server-Dienste? Welche technischen Anforderungen ergeben sich daraus?
- Welche Rolle spielt die Verschlüsselung auf den Weitverkehrsstrecken? Welche Techniken stehen hierfür zur Verfügung?
- Wie sind die aktuellen Netzwerk-Techniken Shortest Path Bridging (SPBm),

Transparent Interconnection of Lots of Links (TRILL) sowie proprietäre Entwicklungen in diesem Bereich der Layer-2-Ausdehnung einzuschätzen (z.B. Cisco FabricPath, Juniper QFabric)?

- Welche weiteren Herausforderungen ergeben sich aus der Präsenz gleicher Subnetze an unterschiedlichen geografischen Standorten? Wie sind aktuelle Entwicklungen zu beurteilen, die diese Schwierigkeiten adressieren (z.B. Cisco Location ID Separator Protocol, LISP)?
- Wie geeignet sind die Speicherübertragungsprotokolle Fibre Channel (FC), iSCSI, Fibre Channel over Ethernet (FCoE), Network File System (NFS) für die Replikation von Speicherinhalten?
- Welche Anforderungen bringen Server- und Datenbank-Cluster wie der Oracle Real Application Cluster?
- Welche Möglichkeiten zur Lastverteilung gibt es über mehrere RZ-Standorte (z.B. Global Server Load Balancing)?
- Wie können diese technischen Maßnahmen geeignet in einer Leitlinie zum Betrieb redundanter RZ-Standorte dargestellt werden? Für welche Systeme sollte eine synchrone Spiegelung vorgesehen werden? Wo reicht Asynchronität aus? Was muss überhaupt nicht gespiegelt werden?

Durch das Seminar führt Dipl.-Inform. Matthias Egerland von der ComConsult Beratung & Planung GmbH.

Fax-Antwort an ComConsult 02408/955-399

Anmeldung

RZ-RZ-Kopplung - alles nur eine Frage der Bandbreite?

Ich buche das Seminar

RZ-RZ-Kopplung - alles nur eine Frage der Bandbreite?

☐ 26.09.11 in Köln

zum Preis von € 890,- zzgl. MwSt.

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift



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

Aktuelle Neuerscheinungen bei ComConsult-Study.tv

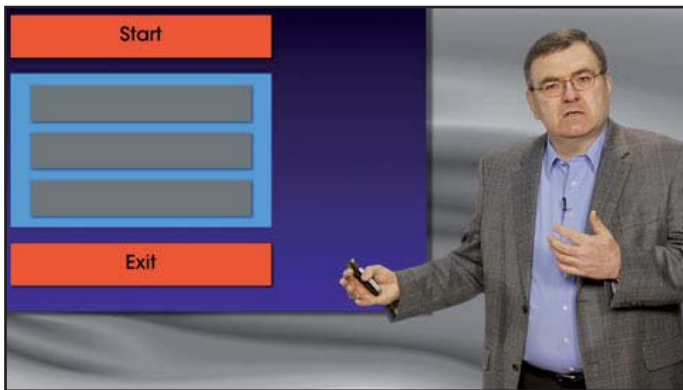
Themenbereich: Betrieb und Architekturen

Seminar: Besser Präsentieren

Referent: **Dr. Jürgen Suppan**

Zeit: 00:33:38 gesamt

Preis: Kostenlos



Erfolgreich präsentieren zu können, ist ein wichtiger Baustein der beruflichen Entwicklung. In dieser Video-Serie erklärt Dr. Suppan, wie jeder von uns seine Präsentations-Fähigkeiten deutlich verbessern kann.

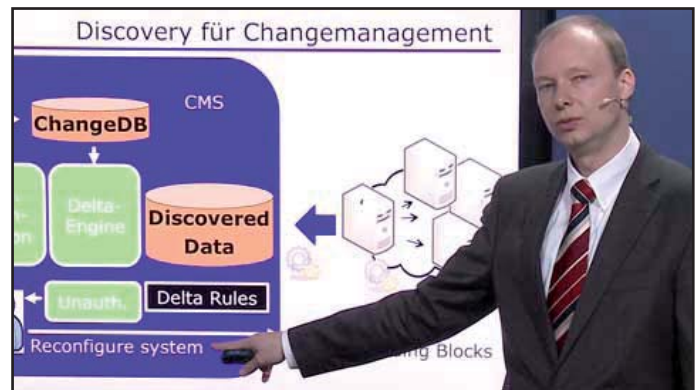
Themenbereich: Betrieb und Architekturen

Discovery Architektur

Referent: **Dipl.-Ing. Ingo Voßkötter**

Zeit: 00:30:52

Preis: Kostenlos mit Abo



Discovery stimmt immer! Wir haben diesen Traum der vollautomatischen Bestandspflege als Basis für Change und Incident Management. Ingo Voßkötter analysiert, ob dieser Traum umsetzbar ist. Er stellt die Lösungs-Ansätze vor, die sich in der Praxis bewährt haben.

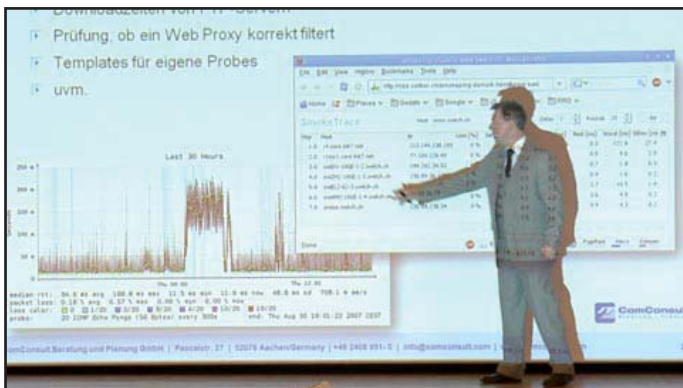
Themenbereich: Netzwerke

Die Zukunft der Netzanalyse

Referent: **Dr.-Ing. Jochen Wetzlar**

Zeit: 00:40:42

Preis: Kostenlos mit Abo



Als weiteres Highlight für die Abonnenten von ComConsult-Study.tv einen der best bewerteten Vorträge des ComConsult Redesign-Forums 2011.

Themenbereich: Analyse und Strategie

Cloud Storage in der Analyse

Referent: **Dr. Jürgen Suppan**

Zeit: 00:56:13

Preis: Kostenlos mit Abo



Dr. Suppan stellt eine hochaktuelle Analyse von ComConsult Research zur Bedeutung von Cloud Storage für Behörden und Unternehmen vor.

Schwerpunkthema

Wie viel Delay ist tolerierbar?

Fortsetzung von Seite 1



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

Was ist ULL?

Ultra-Low Latency ist ein Begriff, der ursprünglich aus dem Bereich Optical Networking kommt. In diesem Zusammenhang geht es darum, dass optische Komponenten wie Wellenlängenmultiplexer möglichst wenig Delay in den Signalpfad einfügen. In letzter Zeit wurde der Begriff von der Ethernet-Industrie übernommen. Dabei ist von „Sub-Microsecond Latency“ die Rede, das heißt die Switches sollen weniger als eine Mikrosekunde Delay bei der Übertragung jedes Paketes verursachen.

Unklar ist, welche Latency-Definition dabei gilt. Es gibt nämlich vier denkbare, die in der Abbildung 1 dargestellt sind. Je nach Wahl des Messverfahrens ergeben sich unterschiedliche Werte, die miteinander nicht vergleichbar sind. Für die Anwendungen ist die LIFO-Latenz entscheidend, denn damit wird gemessen, um welche Zeit das letzte Bit des Paketes (das letztlich für den vollständigen Empfang des Paketes entscheidend ist) durch den Switch verzögert wird. Die LIFO-Latenz ist nur gleich der FIFO-Latenz, wenn die Bitraten von Input und Output gleich

sind. Für den Vergleich zwischen zwei Switches, die im Modus Store & Forward arbeiten, sollte der LIFO-Wert herangezogen werden. Nicht sinnvoll ist die Betrachtung der FIFO-Latenz.

Was können Switch-Hersteller tun, um Delays zu minimieren? Sie können Switch-Architekturen entwerfen, in denen die Pakete im Switch minimal verzögert werden. Wo entstehen im Switch die längsten Latenzzeiten? Sie entstehen nicht hauptsächlich auf den wenige Zentimeter oder Dezimeter langen Pfaden im Switch, sondern vor allem in den Zwischenspeichern (Puffern). Das Ziel sollte also sein, dass die Pakete nur minimale Zeit in den Puffern verweilen, auch wenn der Puffer groß dimensioniert sein sollte.

Aber es gibt noch eine Einflussgröße, und das ist der Switching-Modus.

Déjà-vu-Erlebnis

Diejenigen, die länger als 15 Jahre mit Lokalen Netzen zu tun haben, bekommen beim Verfolgen der ULL-Diskussion ein Déjà-vu-Erlebnis. In der ersten Hälfte des 1990er Jahrzehnts, also den An-

fängen der LAN-Switch-Technik, war eine Diskussion um das schnellste Switching-Verfahren ausgebrochen. Kalpana, einer der ersten LAN-Switch-Hersteller, der später von Cisco Systems übernommen wurde, brachte Switches auf den Markt, die eine Weiterleitung im Cut-Through-Modus unterstützten. Bei diesem Modus beginnt der Switch mit der Übertragung am Ausgangsport, ohne auf das letzte am Eingangsport empfangene Bit eines Frames zu warten. Im Gegensatz dazu wartet ein im Modus Store & Forward arbeitender Switch immer auf das letzte am Eingangsport empfangene Bit eines Paketes, bevor er das Paket am Ausgangsport weiter leitet. Der Geschwindigkeitsvorteil des Cut-Through-Modus ist einleuchtend (siehe Abbildung 2). Damals arbeiteten die Switches an allen Ports mit 10 Mbit/s. Bei einem Paket der Länge von ca. 1.500 Bytes lag der Geschwindigkeitsvorteil des Cut-Through-Modus in der Größenordnung einer Millisekunde. Eine Millisekunde entsprach und entspricht ungefähr einer zusätzlichen Kabellänge von 200 km. Bei fünf kaskadierten Switches (in einer typischen Hierarchie Access, Distribution, Core, Distribution, Access) käme dann der Gegenwert eines Kabelwegs von 1000 km zusammen, für die meisten Anwendungen auch auf der Ebene der Benutzerwahrnehmung relevant.

Der Cut-Through-Modus erledigte sich aber relativ schnell, nachdem 1995 der Fast-Ethernet-Standard vom Institute of Electrical and Electronic Engineers (IEEE) verabschiedet worden war und die ersten Switches mit Uplinks auf den Markt kamen, die 100 Mbit/s unterstützten. Wenn nämlich ein Frame an einem Port mit 10 Mbit/s empfangen wird, kann an einem mit 100 Mbit/s arbeitenden Zielport nicht mit der Weiterleitung begonnen werden, solange das letzte Bit des Pakets nicht vom Switch empfangen worden ist, denn

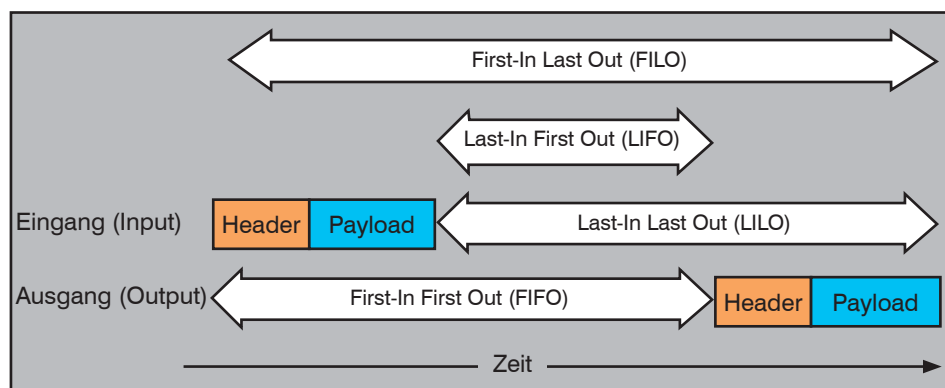


Abbildung 1: Verschiedene Verfahren der Latenzmessung

¹ Beispiel: ComConsult Study.tv, LISP-Seminar, <http://www.comconsult-study.tv/de/LISP::1553:1314.html>

Wie viel Delay ist tolerierbar?

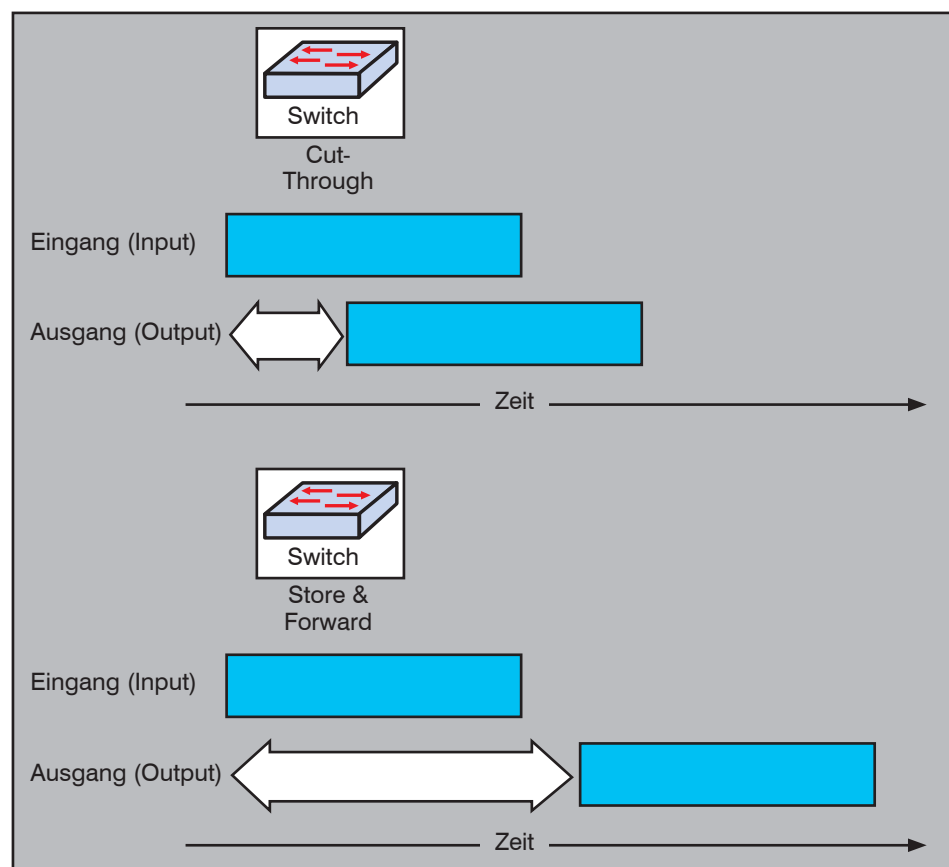


Abbildung 2: Cut-Through- versus Store&Forward-Modus

bekanntlich kostet aufgrund der unterschiedlichen Bitraten der beiden Ports das Senden des Paketes nur 10 % der Zeit, die für den Empfang erforderlich ist. Da das Paket am Stück gesendet werden muss, ist das Ende des empfangenen Paketes abzuwarten. Der Geschwindigkeitsvorteil des Cut-Through-Modus fällt dann nicht mehr so deutlich aus. (siehe Abbildung 3)

Die Lokalen Netze wurden von 1995 an in der Regel nach einem hierarchischen Konzept aufgebaut, bei dem die Uplinks normalerweise mit einer höheren Bitrate arbeiten als die Downlinks.

So kam es, dass der Cut-Through-Modus nach 1995 langsam in die Vergessenheit geriet, bis der Begriff in letzter Zeit wieder aufgetaucht ist.

ULL: wer braucht das?

Der neue Zusammenhang, in dem wieder vom Cut-Through-Modus die Rede ist, heißt Ultra-Low Latency (ULL). Der Anspruch ist, die Latenz von Switches in den Bereich von „Sub-Microsecond“ zu drücken. Geht man von einem Switch mit einer Mindestverarbeitungszeit von einer Mikrosekunde für die Weiterleitung eines Frames aus, die allein für das Nachschla-

gen in Adresstabellen und anderen Verarbeitungsschritten im Switch erforderlich ist, lassen sich die theoretischen Minima der Switch-FIFO-Latenzen bei Switches, die durchgängig mit 10Gigabit Ethernet Ports ausgestattet sind, gemäß der Tabelle 1 ausrechnen (die FIFO-Latenz wäre aufgrund der einheitlichen Bitraten der Ports auch gleich der LIFO-Latenz). Das sind auch die für die Applikationen wirksamen Latenzen.

Pro Switch sind je nach Frame-Länge bis über 6 Mikrosekunden (genauer: $7,373 \mu s$ minus $1,0 \mu s$, also $6,373 \mu s$, ungefähr $6,4 \mu s$) Unterschied zwischen den beiden Modi Cut-Through und Store & Forward anzunehmen. Eine Mikrosekunde entspricht einem Kabelweg von 200 m. Bei fünf kaskadierten Switches berechnet sich daher der maximale latenzbezogene Nachteil von Store & Forward wie folgt:

Maximum_Delta (Cut-Through; Store & Forward) = $5 \times 6,4 \mu s = 32 \mu s \triangleq 6.400 \text{ m}$

Der Latenzunterschied, der 6.400 m Kabelweg entspricht, ist ungefähr mit dem Unterschied vergleichbar, der sich ergäbe, wenn man die Systeme eines Data Centers auf zwei Lokationen an den beiden Enden eines großen Campus verteilte.

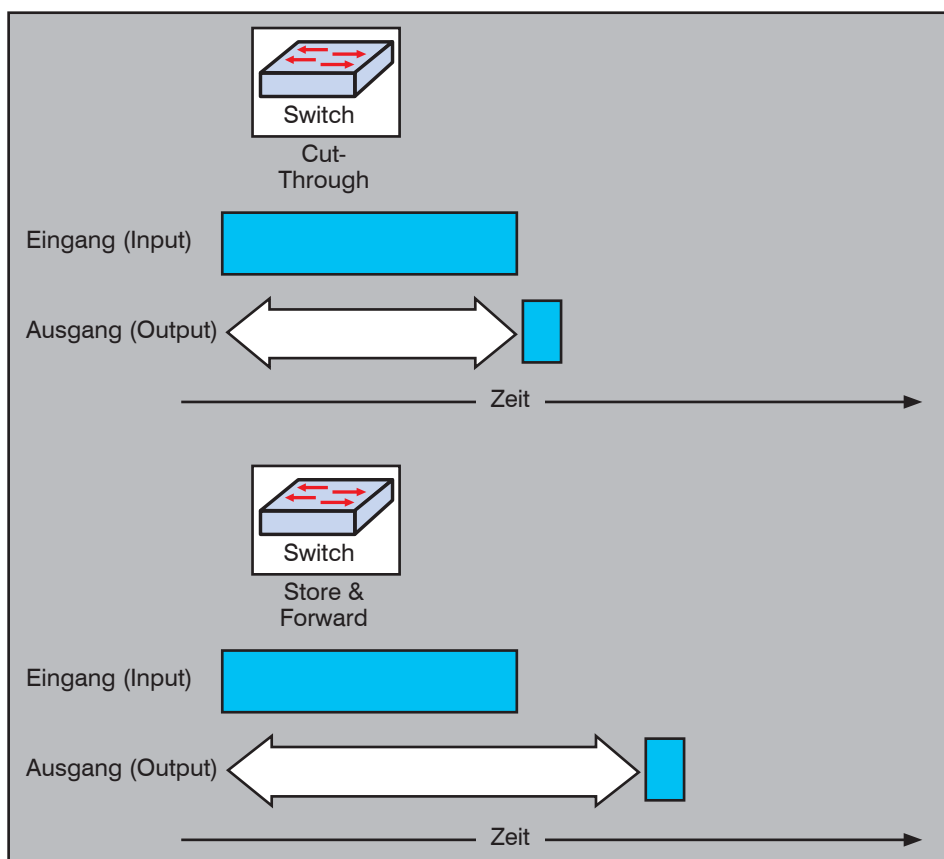


Abbildung 3: Cut-Through- versus Store&Forward-Modus bei unterschiedlichen Bitraten der Ports

Wie viel Delay ist tolerierbar?

Länge in Bytes	Länge in Bits	Mindest-FIFO/LILO-Latenz in Nanosekunden	
		Cut-Through	Store & Forward
72	576	1.000	1.000
144	1152	1.000	1.000
288	2304	1.000	1.000
576	4608	1.000	1.000
1152	9216	1.000	1.000
2304	18432	1.000	1.843
4608	36864	1.000	3.686
9216	73728	1.000	7.373

Tabelle 1: Vergleich der Mindestlatenzwerte bei Cut-Through und Store & Forward

Ist das überhaupt relevant? Müssen nicht alle rechenzentrumsinternen Datenströme Entfernungen tolerieren, die sogar weit über die Campusausdehnungen hinausgehen? Hinsichtlich Disaster Recovery würden zwei Data Center auf einem Campus nicht alle Szenarien abdecken, wenn man an solche Ereignisse wie Erdbeben, Tsunami, lang anhaltenden Stromausfall, Flugzeugabsturz, Hochwasser oder terroristische Anschläge denkt. Viele Unternehmen und Organisationen erheben den Anspruch, solche Ereignisse zu überleben und innerhalb sehr kurzer Zeit danach wieder arbeitsfähig zu sein. Die RZ-Betreiber sind in solchen Fällen gut beraten, zentrale IT-Ressourcen auf Lokationen zu verteilen, die auch bei größten anzunehmenden Unfällen (so etwa bei einem solchen Ereignis wie der März-Katastrophe in Japan) nicht gleichzeitig betroffen wären.

Nun kann man natürlich zwischen Disaster Recovery und dem normalen RZ-Betrieb unterscheiden. Im ersten Fall kann zum Beispiel der Anspruch „Business Continuity“ dahin gehend präzisiert werden, dass man nach der Katastrophe schnell wieder arbeitsfähig wird, aber nicht ganz ohne Datenverlust. In einem solchen Fall könnte es zum Beispiel reichen, dass eine Organisation mit Daten weiter arbeitet, die wenige Stunden vor dem Schadensfall gesichert worden sind, zum Beispiel mittels Snapshots von einem RZ auf eine andere Data-Center-Lokation ein paar hundert Kilometer weiter.

Dieselbe Organisation kann jedoch eine RZ-Infrastruktur konzipieren, die im Normalfall auf einen Campus oder sogar ein Gebäude konzentriert ist. In diesem Fall gäbe es natürlich einen signifikanten Latenzunterschied zwischen einem Netz auf der Basis von Cut-Through-Switches und einem Netz, das aus Store&Forward-Switches besteht. Die Größenordnung des Unterschieds liegt schlimmstenfalls im zweistelligen Mikrosekundenbereich.

Eine andere Frage ist, ob sich selbst solche großen Latenzunterschiede auf der Ebene der Anwendungen spürbar auswirken.

Wann wirken sich Mikrosekunden aus?

Das Zeitdiagramm einer typischen Client-Server-Applikation ist in der Abbildung 4 dargestellt.

Wenn man eine Transaktion auf der Anwendungsebene (zum Beispiel eine Datenbanktransaktion) als Folge von n Request-Response-Paaren modelliert, ergibt sich für die Dauer der Transaktion die Formel 1.

Dabei bedeuten:

- t (Transaktion): Dauer der Transaktion insgesamt, bestehend aus n Request-Response-Paaren
- t (Request _{i}): Dauer der Übertragung des Request-Blocks i auf dem langsamsten Übertragungssegment des Ende-zu-Ende-Pfades
- t (Latency1 _{i}): Netz-Latency des Request-Blocks i
- t (Server _{i}): Server-Latency nach dem Request-Block i
- t (Latency2 _{i}): Netz-Latency des Response-Blocks i

$$t(\text{Transaktion}) = \sum_{i=1}^{i=n} t(\text{Request_}i) + \sum_{i=1}^{i=n} t(\text{Latency1_}i) + \sum_{i=1}^{i=n} t(\text{Server_}i) + \sum_{i=1}^{i=n} t(\text{Latency2_}i) + \sum_{i=1}^{i=n} t(\text{Response_}i) + \sum_{i=1}^{i=n} t(\text{Client_}i)$$

Formel 1: Dauer der Transaktion

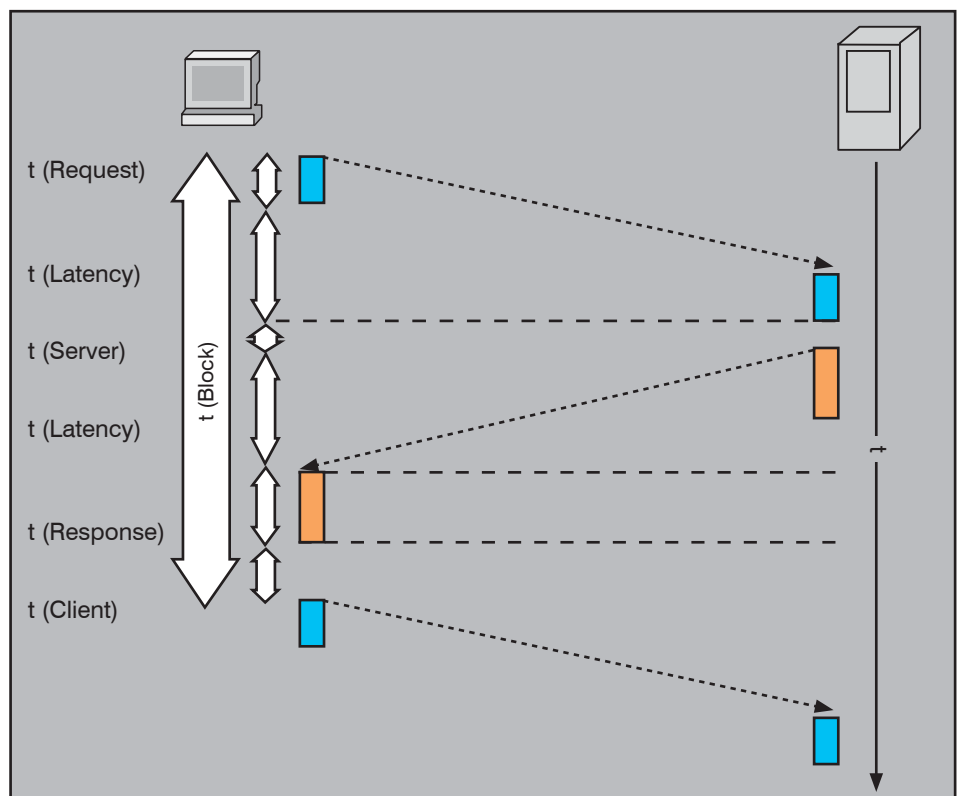


Abbildung 4: Zeit-Diagramm einer Client-Server-Applikation

Wie viel Delay ist tolerierbar?

- $t(\text{Response}_i)$: Dauer der Übertragung des Response-Blocks i
- $t(\text{Client}_i)$: Client-Latency nach vor dem Request-Block $(i+1)$

Mit entscheidend ist dabei die Blockgröße. Dies ist die Länge der Sequenz, die pro Richtung übertragen werden kann, ohne dass auf eine Bestätigung in entgegengesetzter Richtung gewartet werden muss.

Für die weitere Berechnung treffen wir folgende Vereinfachungen:

- Alle Request-Response-Paare einer Transaktion weisen gleiche Latenzwerte aller Komponenten auf.
- Die Netzlatenz ist in beiden Richtungen gleich.
- Die Blockgröße ist so klein, dass ein Block an einem Stück übertragen werden kann, ohne auf eine Bestätigung warten zu müssen (das entspricht der obigen Definition von „Block“).

Somit bekommen wir die einfachere Formel, in der RTT für Round Trip Time steht:

$$[t(\text{Transaktion}) = n \times (t(\text{Request}) + \text{RTT} + t(\text{Response}) + t(\text{Server}) + t(\text{Client}))]$$

$t(\text{Request})$ und $t(\text{Response})$ sind wie folgt zu berechnen:

$$t(\text{Request}) = \text{Länge des Request-Blocks} / \text{Bitrate}$$

$$t(\text{Response}) = \text{Länge des Response-Blocks} / \text{Bitrate}$$

Wenn man für die Summe aus $t(\text{Request})$ und $t(\text{Response})$ einfach die ausgetauschte Datenmenge einsetzt, vereinfacht sich die o. g. Formel noch weiter:

$$t(\text{Transaktion}) = n \times (\text{RTT} + t(\text{Server}) + t(\text{Client})) + \text{Datenmenge} / \text{Bitrate}$$

Somit wirkt sich die Round Trip Time im Netz in drei Fällen nicht signifikant auf die Transaktion aus:

- Wenn $\text{RTT} \ll t(\text{Server})$
- Wenn $\text{RTT} \ll t(\text{Client})$
- Wenn $n \times \text{RTT} \ll \text{Datenmenge} / \text{Bitrate}$, d. h. wenn $n \times \text{RTT} \times \text{Bitrate} \ll \text{Datenmenge}$

Die Frage ist also, für welche Anwendungen der oben berechnete Latenzunterschied von $32 \mu\text{s}$, d. h. der RTT-Unterschied von $2 \times 32 \mu\text{s} = 64 \mu\text{s}$ zwischen Cut-Through und Store&Forward relevant ist. Nicht relevant ist dieser Unterschied in

folgenden Fällen:

- wenn die Entfernung zwischen Client und Server wesentlich größer als 6,4 km sein kann, d. h. für alle Anwendungen, die für das WAN und mobile Clients tauglich sein müssen,
- wenn Server und Clients im Durchschnitt wesentlich mehr als 64 Mikrosekunden Verarbeitungszeit für die Bearbeitung eines Request-Response-Paares benötigen, oder
- wenn das Produkt aus n , RTT und Bitrate wesentlich kleiner als die gesamte Datenmenge ist, die pro Transaktion übertragen werden muss.

Der letztgenannte Fall lässt sich für die Bitrate 10 Gbit/s und $\text{RTT} = 64 \mu\text{s}$ berechnen:

$$\text{Datenmenge} \gg n \times 64 \mu\text{s} \times 10 \text{ Gbit/s} = n \times 640.000 \text{ Bits} = n \times 80.000 \text{ Bytes}$$

Bei Datenbanktransaktionen sind mittlere sechsstellige Byte-Zahlen pro Transaktion keine Seltenheit; die letztgenannte Bedingung trifft daher nicht zu. Auch ist es nicht üblich, Client und Server eines typischen Datenbankprotokolls wie Net8 von Oracle weiter entfernt als wenige km aufzustellen.

Bleibt noch die Frage, wie schnell Server und Clients auf Requests bzw. Responses reagieren und ob deren Latenz wesentlich über 64 Mikrosekunden liegt. Diese Latenz ist natürlich nicht nur von den Prozessoren der betreffenden Maschinen abhängig, sondern auch von der I/O-Leistung, die den Maschinen zur Verfügung steht. Letzteres würde bei der

Nutzung eines zentralen Speichers der einem Server durchschnittlich zur Verfügung stehenden I/O-Leistung entsprechen. Geht man vom Wert 100.000 IOPS bei einem High-End-Speichersystem aus, welches 100 physikalische Server bedient, stehen einem Server durchschnittlich 1000 IOPS zur Verfügung. Selbst eine einzige I/O-Aktion würde also durchschnittlich einer Millisekunde entsprechen. ULL lohnt sich also nur dann, wenn pro Transaktion nur wenige I/O-Aktionen anfallen. Auf keinen Fall darf dann pro Request eine I/O-Operation erforderlich sein, denn eine Millisekunde I/O-Zeit ist ja bekanntlich wesentlich mehr als $64 \mu\text{s}$.

Anders stellt sich die Situation dar, wenn alle mit der Transaktion zusammenhängenden I/O-Operationen ohne Inanspruchnahme von Festplattenleistung vorstatten gehen, also wenn zum Beispiel statt Festplatten Solid State Devices (SSD) eingesetzt werden. Diese führen I/O-Operationen wesentlich schneller aus.

Hersteller wie Cisco geben bei der Werbung für ULL-Switches an, diese seien besonders geeignet für Börsenanwendungen in einer sogenannten Co-location, in der die Trading-Rechner zusammengeschaltet sind². Es geht darum, dass Transaktionen binnen Millisekunden über die Bühne gehen. Das sind dann keine von Menschen gesteuerten Transaktionen, sondern zum Beispiel Trading-Aufträge, die automatisch und ohne Zutun von menschlichen Händlern über die Bühne gehen. Bestimmte Ereignisse lösen den Kauf oder Verkauf bestimmter Werte aus. Hier könnten theoretisch Millisekunden über Millionen Gewinn oder Verlust entscheiden.

Seminar

RZ-RZ-Kopplung - alles nur eine Frage der Bandbreite? 26.09.11 in Köln

Diese 1-tägige Veranstaltung konzentriert sich auf aktuelle Fragestellungen, die bei der Kopplung von modernen Rechenzentren entstehen. Neben der Motivationslage für eine solche Kopplung werden die physikalischen Rahmenbedingungen für ein solches Vorhaben aufgezeigt, um die natürlichen Grenzen hinsichtlich Latenz und Übertragungsrates greifbar zu machen. Anschließend werden aus den Bereichen Server, Storage und Netzwerk die aktuellen Cluster-Mechanismen sowie Techniken für Spiegelungs- und Replikationsverfahren im Kontext entfernter RZ-Standorte besprochen. Die Kombination aus strategischen Überlegungen und technischen Maßnahmen bilden abschließend eine Leitlinie, welche Systeme in welcher Form bei einer RZ-RZ-Kopplung zu berücksichtigen sind

Referent: Dipl.-Inform. Matthias Egerland

Preis: € 890,- zzgl. MwSt.



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

²<http://www.cisco.com/en/US/products/ps11541/index.html>

Wie viel Delay ist tolerierbar?

Kritisch ist hier anzumerken:

- Ereignisse, welche die Börse betreffen, sind globale Ereignisse. Kursschwankungen in Japan machen sich in Europa frühestens nach hunderten Millisekunden bemerkbar. Dann sind Mikrosekunden Switch Latency bzw. wenige Millisekunden Transaktionszeit nicht mehr so relevant.
- Spätestens nach den vielen teuren Panen, die durch automatisiertes Trading ohne Kontrollen durch Menschen passiert sind, ist die Frage nach der Sinnfälligkeit und den Risiken solchen Trdings berechtigt. Letztendlich ist High Frequency Trading nichts anderes als ein Automat, der auf externe Ereignisse mit Aktionen reagiert. Dazu ist eine Software erforderlich, die mit Heuristiken, Wahrscheinlichkeiten etc. arbeitet. Wenn diese Software nicht optimiert ist und Fehlentscheidungen trifft, nutzt ein schnelles System auch nichts. Werden jedoch Menschen für die Kontrolle von Trading und letzte Checks eingesetzt (damit Automaten keine milliarden schweren Schäden anrichten), brauchen die Aufträge keine Ultra-Low Latency.

Wie in den 1990er Jahren wird sich mit der Einführung von 40/100Gigabit Ethernet in den meisten Netzen ohnehin wieder eine hierarchische Anordnung von Links im Netz einstellen, die mit unterschiedlichen Bitraten arbeiten. Dann wird der Cut-Through-Modus wie vor 15 Jahren in Vergessenheit geraten. Was bleibt dann von ULL? Switches, die Weiterleitungsentscheidungen auch beim Ansatz Store & Forward schnell genug fällen, um allen Applikationen gerecht zu werden.

Full Mesh versus Tree

Was ist aber mit dem Netzdesign? Stimmt das, was zum Beispiel Juniper Networks über die Vorteile einer Full-Mesh-Topologie gegenüber einem hierarchischen, baumförmigen Design sagt?³

Die Quintessenz solcher Aussagen ist diese: Der „Tree“, also die baumförmige hierarchische Anordnung der Switches, ist mit zu viel Latenz verbunden. Sie muss durch Full Mesh ersetzt werden. Es geht darum, dass die Paketübertragung zwischen zwei beliebigen Punkten im RZ-Netz über möglichst wenige Switches geht, damit nicht zu viel Latenz auf dem Übertragungspfad anfällt.

Die Nachrichtentechnik und speziell die Netztechnik blickt auf jahrzehntelange Erfahrung zurück. Die hierarchischen Netzstrukturen sind nicht ohne Grund entstan-

den oder nur weil bestimmte Hersteller dazu passende Komponenten verkaufen wollten. Die Latenzproblematik stellt sich in anderen Größenordnungen auch in der weltweiten Telekommunikation. Auch da geht es um Latenzminimierung. Es wäre sicherlich vom Vorteil, wenn alle Provider-Switches dieser Welt über direkte Kabel miteinander vollvermascht wären. Dies ist aber aus Gründen der Skalierbarkeit der Switches und des hohen Verkabelungsaufwands nicht möglich. So entwickelte sich im Laufe der Jahrzehnte ein hierarchisches Netz.

Die Entstehungsgeschichte hierarchischer RZ-Netze war ähnlich. In den 1990er Jahren haben viele Unternehmen ihre wenigen Server direkt an große zentrale Switches angeschlossen. Bus- und ringförmige Lokale Netze wurden durch sogenannte Collapsed Backbones abgelöst. Auch vor 20 Jahren wusste man, dass damit Latenzen minimiert werden können.

Der Collapsed Backbone funktionierte aber nur bis zu einer niedrigen Anzahl zentraler Switches. Irgendwann waren mehr Ports für die Vollvermaschung der zentralen Switches erforderlich, als diese den Clients und Servern zur Verfügung stellen konnten. Also behalf man sich der guten alten Graphentheorie, die schon immer ein verlässlicher Freund und Helfer der Nachrichtentechnik gewesen ist.

Was ist die optimale Netztopologie?

Stellen wir uns einem grundlegenden Problem der Netztechniker: Was ist die optimale Topologie für ein Netz, das aus Access Switches mit jeweils m Ports besteht und e Endgeräte miteinander verbinden muss?

Je nach Optimierungskriterien fällt die Antwort anders aus. Wenn zum Beispiel die Minimierung der Anzahl der Switches im Vordergrund steht, sollte m so groß wie möglich sein, damit die Anzahl n der Switches möglichst klein ist:

$$n = e \bmod m$$

Was ist aber dann, wenn die Minimierung der Anzahl der Switches (am besten auf den Wert $n = 1$, aus ULL-Gesichtspunkten der beste Wert) nicht das einzige Kriterium ist, das ein Planer eines Rechenzentrums berücksichtigen muss?

Dann sind wir nämlich in der Realität des Netzdesigns angekommen. Eine Vielzahl von Kriterien ist zu berücksichtigen, darunter:

- Managebarkeit
- Hardware-Kosten
- Verfügbarkeit
- Gesamtleistung im Sinne von Durchsatz
- Latenzminimierung

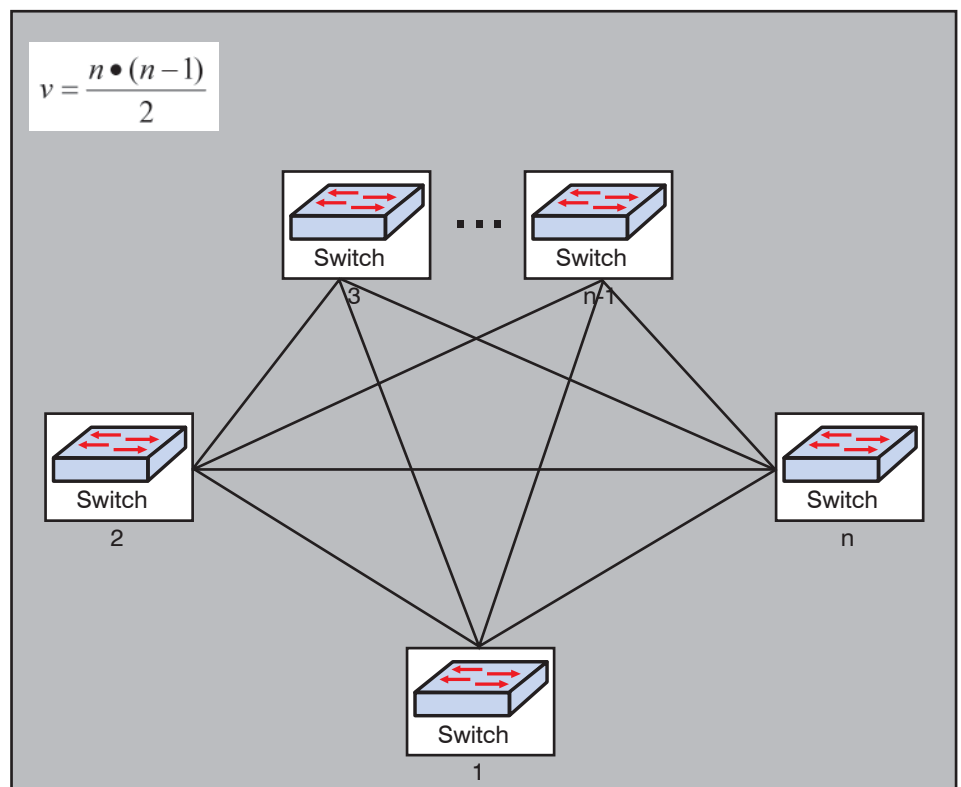


Abbildung 5: Vollvermaschung von n Switches

³http://junipernetworks.com/us/en/company/press-center/press-releases/2011/pr_2011_02_23-13_04.html

Wie viel Delay ist tolerierbar?

Einige dieser Kriterien sind mathematisch zu fassen, andere nicht. Welche Netztopologie ist in einem RZ-Netz am besten managebar? Wie wichtig ist der Aspekt Managebarkeit insgesamt? Sicher ist die Antwort der RZ-Betreiber auf diese Fragen unterschiedlich. Die Antwort ist davon abhängig, wie groß das RZ ist, wie viele Server-Ports zu bedienen sind, wie das Wartungskonzept im RZ aussieht etc. Der Autor arbeitet zum Zeitpunkt der Ausarbeitung dieses Beitrags sowohl in einem Projekt, in dem der Betreiber die Konzentration hunderter Gigabit nebst ein paar Dutzend 10Gigabit Ethernet Ports auf zwei Switches bevorzugt, als auch in anderen Projekten, in denen die RZ-Betreiber auf Top of the Rack Switches bestehen. Wer hat Recht? Keinem dieser RZ-Betreiber darf man Unwissenheit oder Festhalten an überholten Konzepten vorwerfen. Die RZ-Randbedingungen und Präferenzen der unterschiedlichen Betreiber sind unterschiedlich.

Bei einer Vielzahl von RZ-Betreibern scheint sich die Präferenz für Top of the Rack Switches durchgesetzt haben. Betrachten wir diese Gruppe etwas genauer. Wir gehen davon aus, dass n solche Switches erforderlich sind, unabhängig davon, wie groß die Zahl e der zu realisierenden Server Ports ist.

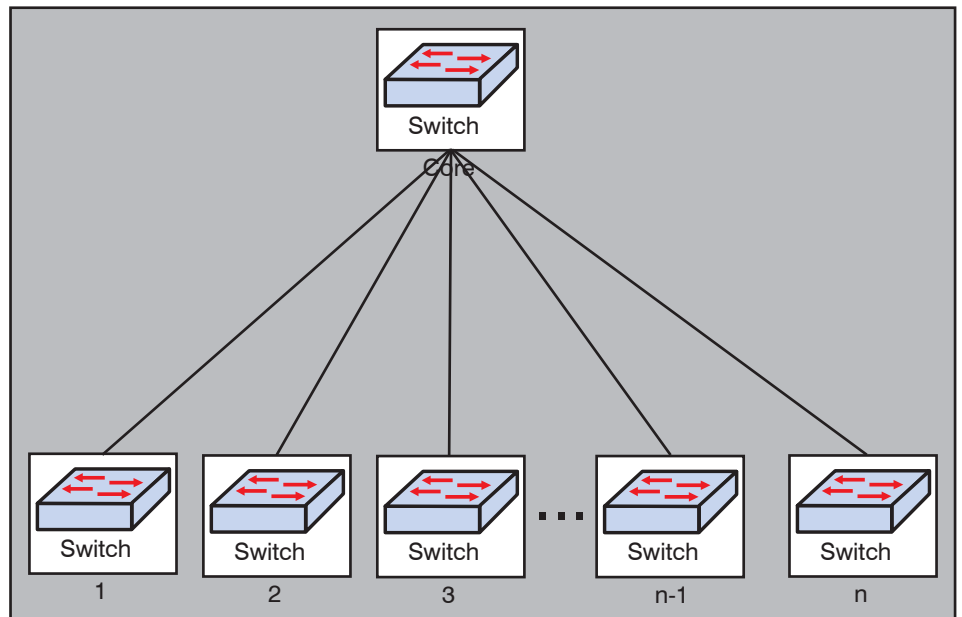


Abbildung 6: Core und Access Switches

Eine Vollvermaschung der ToR-Switches gemäß Abbildung 5 ist zwar aus Latenzsicht optimal, skaliert aber nur begrenzt, denn bekanntlich brauchen wir dafür eine hohe Anzahl von Inter-Switch-Verbindungen v .

Außerdem können bei einem solchen Design, wenn es auf Up- und Downlinkseite der Switches mit derselben Bitrate arbeitet, punktuelle Überlastsituationen dadurch entstehen, dass temporär oder permanent asymmetrische Lasten zu übertragen wären. Wenn also an den Switch x zum Bei-

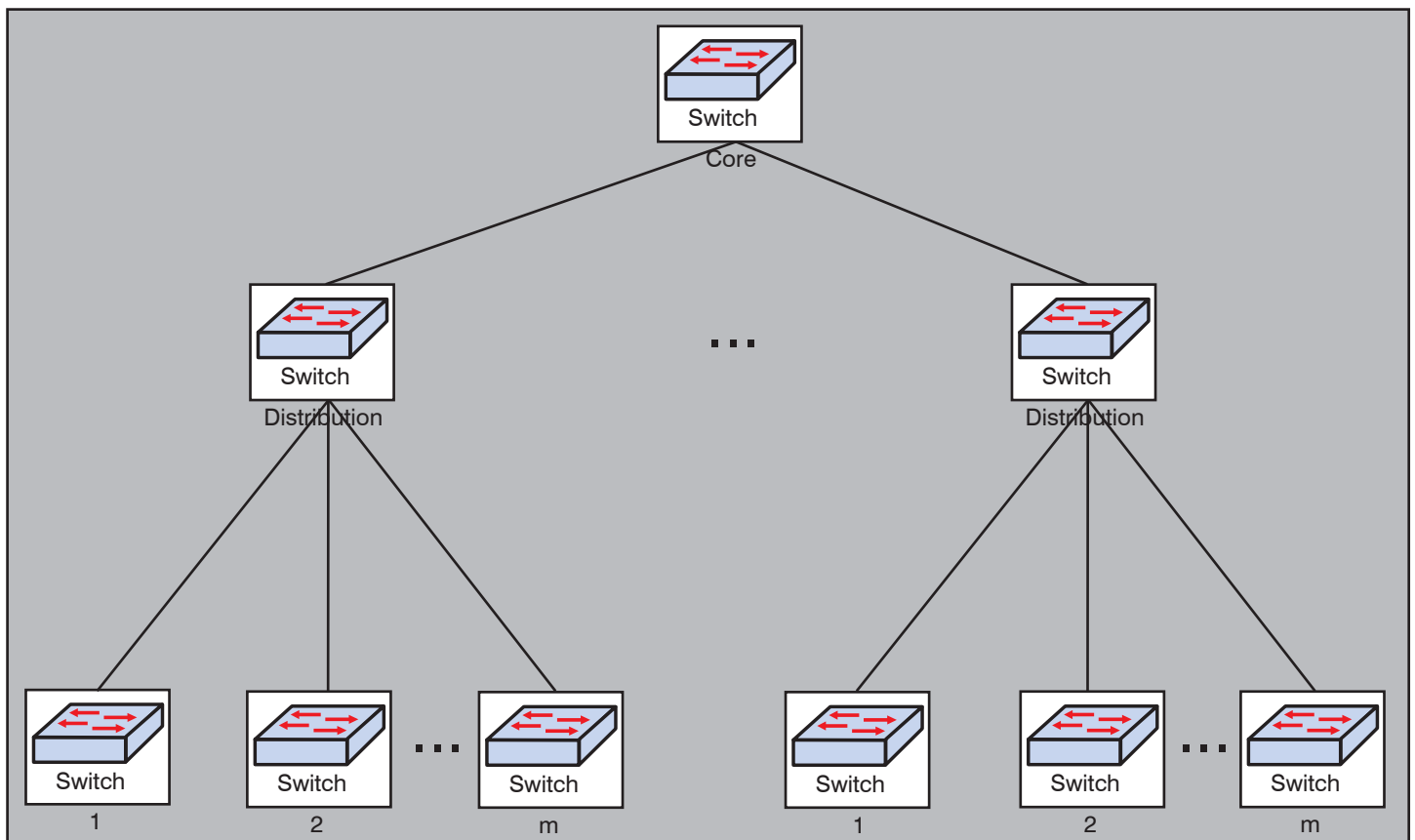


Abbildung 7: Dreistufige Netzhierarchie

Wie viel Delay ist tolerierbar?

spiel ein Datensicherungsserver angeschlossen ist, kann die Verbindung zum Switch y, eine einzige Verbindung, die als „Shortest Path“ gilt, genau im Moment der Datensicherung überlastet sein.

Und hier sind wir an einer wesentlichen Ursache hierarchischer Netzstrukturen angelangt: dem Kommunikationsprofil. Wir haben es in den RZ-Netzen häufig nicht mit dem Profil „jeder mit jedem“ zu tun, sondern mit der Situation „viele mit wenigen“. Ein Datensicherungsserver, ein zentrales Speichersystem etc. gehören zu den „wenigen“, „normale“ Server zu den „vielen“.

Die Vollvermaschung ist nur bedingt skalierbar, und es gibt wenige Maschinen, welche die Last vieler anderer Maschinen aufnehmen müssen. Der Anforderung nach Skalierbarkeit und Asymmetrie wird man in der Regel durch eine hierarchische Struktur gerecht, auch wenn sie nur aus zwei Ebenen besteht: Access Switches und Core Switches, wie in der Abbildung 6 dargestellt.

Damit wird der Hop Count im Vergleich zur Vollvermaschung von zwei auf drei erhöht, was mit einer Erhöhung der Latenz im Netz verbunden ist. Solange die Uplinks mit der gleichen Bitrate wie die Downlinks arbeiten, beträgt der Latenznachteil 50 % (von 2 auf 3 erhöht). Hat man aber Switches zur Verfügung, deren Uplinks (auch aus Lastgründen) ohnehin mit einer wesentlich höheren Bitrate als die Downlinks arbeiten, relativiert sich der Latenzunterschied wieder, denn mit erhöhter Bitrate sinkt die Zeit, die ein Frame braucht, um einen Hop zurück zu legen.

Die meisten RZ-Netzbetreiber werden mit einer zweistufigen Hierarchie auskommen. In den größten RZ-Netzen reicht aber diese Struktur nicht aus; man denke etwa an Rechenzentren mit hunderten Schränken und damit hunderten Server Access Switches, womöglich verteilt auf verschiedene Lokationen. In solchen Umgebungen muss eine dritte Hierarchiestufe dazu kommen, wie in der Abbildung 7 dargestellt.

Auch hier gilt: die Verschlechterung der maximalen Hop Counts von 3 (im Falle der zweistufigen Hierarchie) auf 5 (im Falle der dreistufigen Hierarchie) wird dadurch kompensiert, dass Uplinks eine höhere Bitrate als die Downlinks unterstützen.

Die Anzahl der Hierarchiestufen und damit Hop Counts im Netz ist damit immer von der Größe des Netzes abhängig. Daran, dass größere Netze tendenziell mehr Hierarchiestufen brauchen als kleinere, hat sich in den letzten Jahrzehnten nichts geändert.

Zwischenfazit zu ULL

Als Zwischenfazit zu Ultra-Low Latency ist festzuhalten:

- ULL ist bei Kommunikationsströmen irrelevant, welche wesentlich längere Entfernungen als wenige Kilometer zurücklegen müssen. Solche langen Entfernungen sind für Disaster Recovery geboten.
- Wenn überhaupt wirkt sich ULL nur bei sehr leistungsfähigen sonstigen Komponenten aus: Prozessoren auf beiden Seiten, Storage-Systeme.
- Mit der Einführung von 40/100Gigabit Ethernet wird der Cut-Through-Modus als wesentliche ULL-Eigenschaft in den meisten Netzen irrelevant.
- Transaktionen, die Menschen mit Reaktionszeiten im Sekundenbereich involvieren, brauchen ULL nicht, sondern nur Transaktionen mit Kommunikation zwischen Maschinen.
- ULL ändert nichts daran, dass die Anzahl der benötigten Hierarchiestufen mit der Größe eines RZ-Netzes steigt.

Switch- und Chip-Hersteller können weder mathematische Gesetze der Graphentheorie noch physikalische Gesetze aufheben. Signale sind maximal mit Lichtgeschwindigkeit zu übertragen. Dabei bleibt es.

Wenn jedoch mit ULL gemeint ist, dass neue Switches durch den Einsatz schneller Chips für die Minimierung der Latenzen im Switch sorgen, dann ist das eine Selbstverständlichkeit. Mit fortschrei-

tender Chip-Technologie steigt die Geschwindigkeit von Switches, und die Latenz in den Switches sinkt. Allein schon die steigenden Bitraten erfordern immer schnellere Chips.

Delay im WAN

LAN Switch Delays sind je nach Anwendung vielleicht innerhalb eines Campus relevant, nicht jedoch im Wide Area Network (WAN). Wenn die Übertragungswege wenige Kilometer überschreiten, dann werden die Latenzunterschiede zwischen verschiedenen Switch-Produkten irrelevant. Im WAN entstehen lästige Latenzen vor allem durch andere Effekte:

- Ungünstige Wege (zum Beispiel von Deutschland nach Indien über Nordamerika, den Pazifik und den Indischen Ozean)
- Überlast an bestimmten Punkten und auf bestimmten Segmenten des Übertragungsweges

Hier muss man zwischen unvermeidbaren und vermeidbaren Latenzeffekten unterscheiden. Da man bekanntlich die Gesetze der Physik nicht überlisten kann, dauert die Signalübertragung von Deutschland nach Indien eine bestimmte Zeit. Man kann die Strecke berechnen⁴:

- Frankfurt am Main bis Mumbai auf dem kürzesten Luftweg: 6.574,536 km
- Dieselbe Strecke auf dem kürzesten befahrbaren Landweg: 9.306 km

Wenn wir also von einem Kabelweg von 10.000 km ausgehen, lässt sich die Zeit für die optische Signalübertragung von Frank-

Kongress**Rechenzentrum Infrastruktur-Redesign
Forum 2011 - 07. - 10.11.11 in Königswinter**

Das ComConsult Rechenzentrum Infrastruktur-Redesign Forum 2011 ist die zentrale Veranstaltung des Jahres, in der hochqualifizierte Referenten die neuen Einflussfaktoren, Möglichkeiten und Strategien vorstellen und mit den Herstellern diskutieren, um Anregungen und praktische Hinweise für die Gestaltung des Rechenzentrums der jetzt beginnenden Zukunft zu erarbeiten.

Moderation: Dr. Franz-Joachim Kauffels, Dr.-Ing. Behrooz Moayeri
Preis: € 2.190,-* zzgl. MwSt. bzw. € 1.790,-* zzgl. MwSt. ohn Workshop

*Preise gültig bis zum 31.08.11 -
dann regulär € 2.390,- zzgl. MwSt. bzw. € 1.990,- zzgl. MwSt. ohne Workshop



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

⁴Siehe luftlinie.org

Wie viel Delay ist tolerierbar?

furt am Main nach Mumbai wie folgt berechnen:

$$t = \frac{10.000\text{km}}{200.000\text{km/s}} = 0,05\text{s} = 50\text{ms}$$

Ein kürzerer RTT-Wert als 100 ms ist also über optische Kabel nicht zu erreichen. Wenn also in High Frequency Trading Rechner an zwei Börsen in Frankfurt und Mumbai involviert sind, erfährt der eine Rechner frühestens nach 50 ms von einem Vorgang auf dem anderen Rechner (in Wirklichkeit frühestens nach ein paar Sekunden, denn zu einer solchen Meldung gehört mehr als nur eine unquitierte Übertragung). Diese Verzögerung ist unvermeidbar, auch wenn Milliarden Euro oder Dollar in neue Kabelwege durch neun Länder investiert werden.

Optimieren kann man also die WAN-Latenz durch die Vermeidung von Überlast (und den damit verbundenen Paketverlusten) bzw. durch die Wahl der direktesten möglichen Route. Ersteres kann man in Gestalt garantierter Netzdurchsatzes kaufen (es kostet nur eine Kleinigkeit). Wie der Provider den Durchsatz sicherstellt, muss ihm überlassen werden. Finanzielles überlassen wir einem geschickten Einkauf.

Letzteres, also die Wahl der optimalen Route, ist kniffliger. Denn dafür ist nicht nur der Provider zuständig, sondern auch der Kunde, wie wir sehen werden.

Ungünstige Wege

Für ungünstige Wege im WAN tragen nicht nur die Provider die Verantwortung. Man stelle sich eine Börse vor, die u. a. RZ-Standorte in Frankfurt am Main und New York unterhält. Von diesem Börsenverbund aus soll eine Verbindung zu einer indischen Börse mit einem RZ in Mumbai geschaffen werden. Da externe Verbindungen in der Regel über Komponenten der Perimetersicherheit geleitet werden, wird die Verbindung nicht direkt, sondern über Firewalls geführt. Weiterhin stelle man sich vor, dass die betroffenen Firewalls nicht in Frankfurt, sondern in New York aufgestellt sind.

Dann nämlich geht der Weg von Frankfurt zunächst nach New York, und erst von dort aus nach Mumbai, wie in der Abbildung 8 dargestellt. Der Weg wird mindestens verdoppelt und damit der RTT-Wert von 100 auf 200 Millisekunden erhöht, auch wenn man Milliarden für einen neuen Kabelweg von New York über die Arktis, Skandinavien, Russland, Kasachstan, Usbekistan, Turkmenistan, Afghanistan und Pakistan ausgibt (je weiter wir auf die-

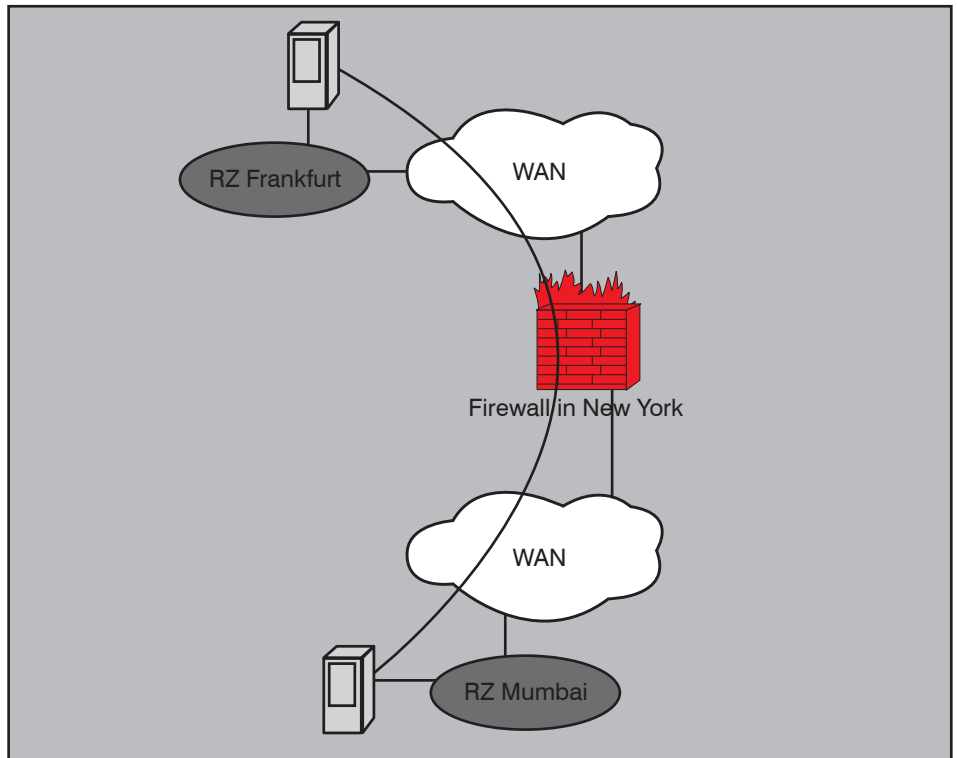


Abbildung 8: Ungünstiger Übertragungsweg

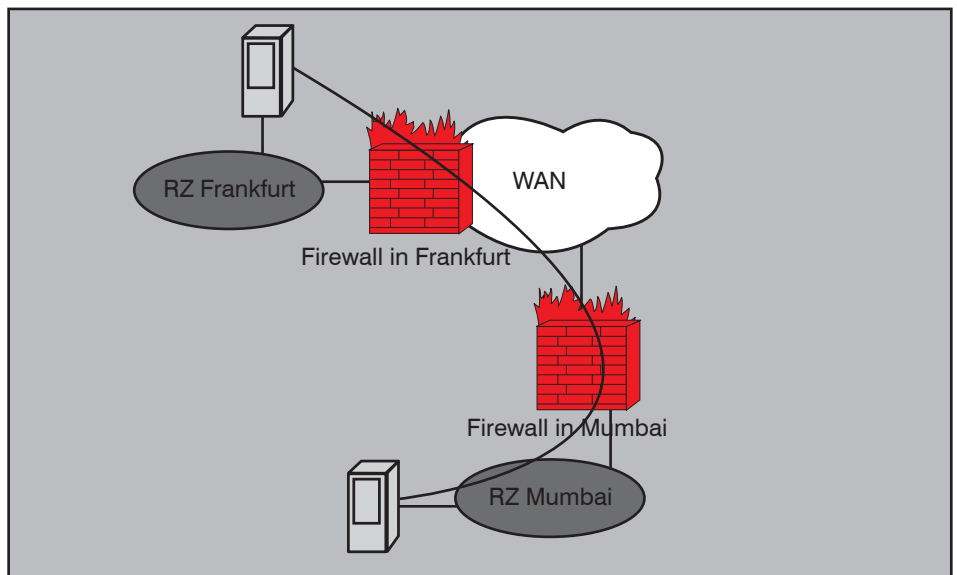


Abbildung 9: Günstiger WAN-Weg

ser Strecke kommen, umso höher sind die Kosten, nicht nur für die Kabelverlegung, sondern für deren militärische Absicherung). Die absolute Steigerung der Responsezeiten fällt noch stärker als 100 Millisekunden aus, weil die Transaktion sicherlich ein paar Durchläufe zwischen Frankfurt und Mumbai erfordern wird.

Um solche ungünstigen Wege zu vermeiden, müssen die RZ- und Netzbetreiber am besten für direkte Verbindungen sorgen. In unserem Fall heißt das: es muss

auch Firewalls in Frankfurt und/oder in Mumbai geben, damit im WAN der Übertragungsweg optimiert werden kann, am besten so, dass der kürzeste Weg zwischen Frankfurt und Mumbai genommen wird. (siehe Abbildung 9)

Im oben genannten Beispiel verursacht die aus Gründen der IT-Sicherheit vorgesehene Firewall-Inspektion den ungünstigen Weg. Es gibt aber auch andere denkbare Ursachen, zum Beispiel:

Wie viel Delay ist tolerierbar?

- Wo ist der Übergabepunkt zwischen den Netzen der jeweiligen WAN Provider in Deutschland und Indien?
- Sind Overlay-Strukturen wie MPLS realisiert, die ein Shortest Path Routing verhindern?
- Können Shortest-Path-Routing-Mechanismen anhand der IP-Adressen den wirklich kürzesten Weg berechnen?

Die Antworten auf diese Fragen sind nicht einfach. Wir werden anhand der dritten der oben gestellten Fragen sehen, dass neue RZ-Strukturen hier ein paar Unwägbarkeiten schaffen.

Auswirkungen der Virtualisierung

Man stelle sich vor, ein Client greife über das WAN auf eine virtualisierte Serverfarm zu, wie in der Abbildung 10 dargestellt. Der Weg geht über den Router am Client-Standort, das WAN, den Router am Standort eines Load-Balancers, den Load Balancer selbst und die Layer-2-Verbindung zwischen zwei RZ-Standorten bis zum aktiven Knoten in der virtualisierten Serverfarm. Je nach Latenz auf der Layer-2-Verbindung zwischen den beiden Rechenzentren ist der Weg signifikant ungünstiger als ein denkbarer direkter Weg vom Client durch das WAN zum Server.

Nun sind Layer-2-Verbindungen zwischen RZ-Standorten, die tausende Kilometer entfernt sind, eher ungewöhnlich, weil andere Restriktionen den Aufbau von Server-Clustern über so lange Wege erschweren. Aber bei einer mehrstufigen Schleuse, zum Beispiel externes Netz - äußere DMZ - innere DMZ - RZ-Netz, in jeder Stufe mit Load Balancing, Firewalling etc. versehen, können aus wenigen Millisekunden auf der Layer-2-Verbindung schnell maximale Latenzen entstehen, die eine Größenordnung höher liegen.

Die Hauptursache hierfür: Die verschiedenen Netz- und Virtualisierungsebenen „wissen“ nichts voneinander und können ihre Wegewahlentscheidungen nicht so treffen, dass sie unter dem Strich, also für die Strecke von einem Ende zum anderen (von Client zum physikalischen Server) mit der kürzesten Latenz verbunden sind.

Dieses „Manko“ wurde mit dem Aufbau der Protokollhierarchie auf den verschiedenen Ebenen des in Netzen üblichen Schichtenmodells bewusst in Kauf genommen. Die Wegewahl soll eben im Layer-2-Netz anhand der MAC-Adresse, im Layer-3-Netz anhand der IP-Adresse und auf höheren Schichten anhand von Ses-

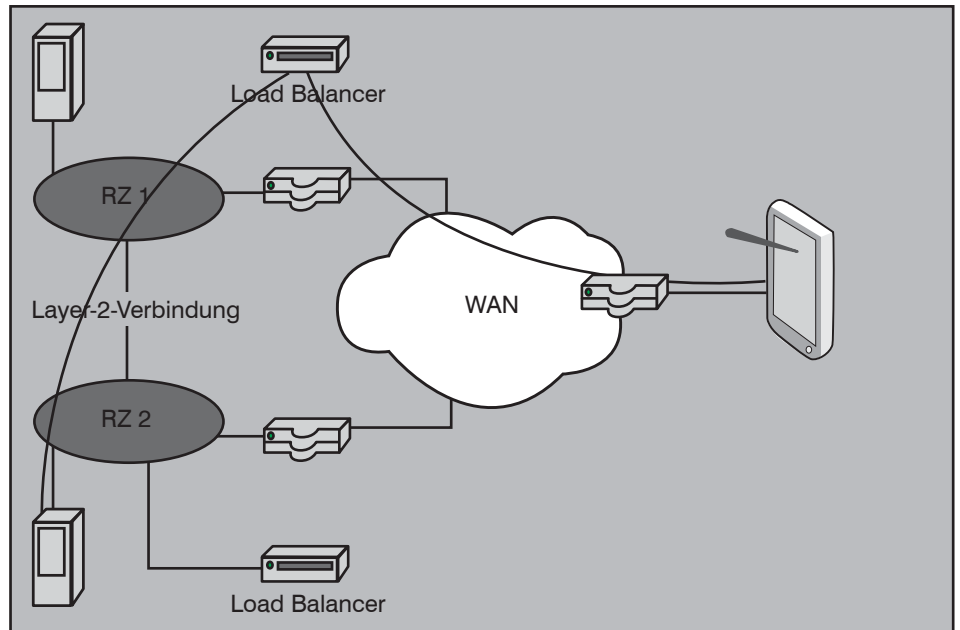


Abbildung 10: Ungünstiger Weg durch Virtualisierung

sion- und Applikations-bezogenen Merkmalen erfolgen, und Änderungen in einer Protokollschicht sollen möglichst keine Auswirkung auf andere Protokollschichten haben. Das ist seit vierzig Jahren das Erfolgsgeheimnis offener Netze. Somit haben wir es mit einem Segen zu tun, der zugleich ein Fluch sein kann, nämlich dann, wenn der Weg von einem Ende durch die Protokollschichten zum anderen Ende zu einem Schlingerkurs wird. Gibt es Abhilfe dagegen?

Overlay-Strukturen

Overlay-Strukturen können helfen, den Weg durch ein WAN zu optimieren. Eine

solche Overlay-Struktur ist in der Abbildung 11 dargestellt.

Der Client greift nicht direkt, sondern über ein Overlay-Netz mit Relays auf den Server zu. Die Relays sind möglichst engmaschig auf das WAN verteilt. In der Nähe des Clients sowie in der Nähe jedes der beiden RZ-Standorte ist jeweils ein Relay aufgestellt. Dann gibt es zwischen Client und Server im Overlay-Netz verschiedene Wege. Man könnte alle Pakete im Overlay-Netz über mehrere Wege übertragen. Die schnellste Kopie, die am Endziel ankommt, bestimmt den günstigsten Weg. Die anderen, langsameren Kopien werden verworfen. Man könnte

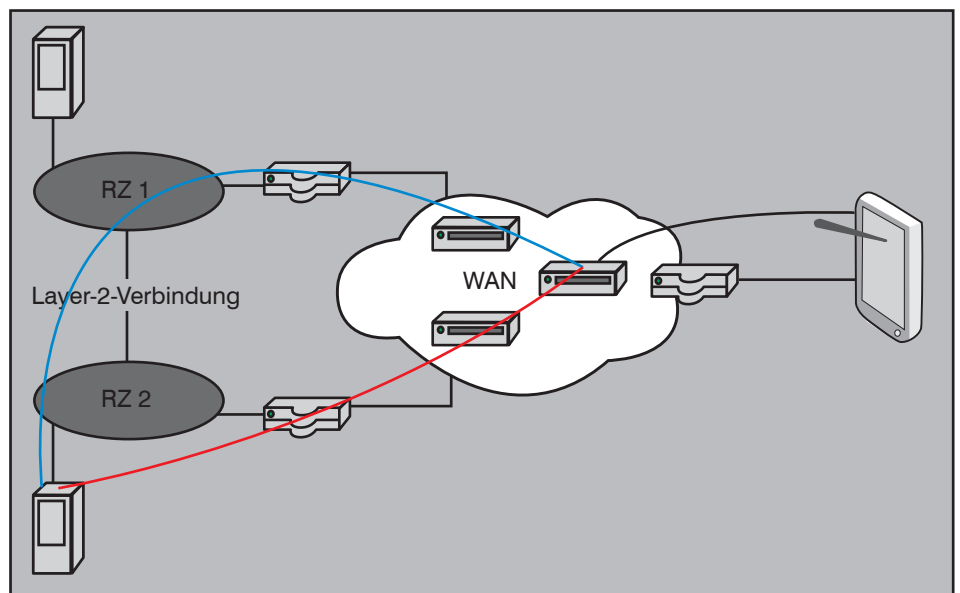


Abbildung 11: Overlay-Struktur

Wie viel Delay ist tolerierbar?

das für alle Pakete tun oder nur auf den Verbindungsaufbau beschränken. Erstes, wenn man auch eine Abhilfe gegen Paketverluste haben will, und Letzteres, wenn man befürchtet, dass überschüssige Kopien desselben Pakets Probleme verursachen könnten.

Genau eine solche Struktur nutzt die Firma Akamai, um den eigenen Kunden mit weltweiten Kommunikationsbeziehungen die Optimierung von Routen durch das globale Netz zu ermöglichen. Die Overlay-Struktur von Akamai ist im Internet erreichbar. In der Regel können die Akamai-Kunden durch ein „Umleiten“ von DNS-Einträgen auf die Akamai-Overlay-Struktur diese nutzen. Dann kann es sogar sein, dass ein Weg durch das Internet und die Akamai-Overlay-Struktur bessere Antwortzeiten aufweist als der Weg durch ein Virtual Private Network auf der Basis einer MPLS-Plattform.

Im Prinzip arbeiten globale Internet-Unternehmen wie Google nicht anders. Sie unterhalten ein engmaschiges Netz an Rechenzentren, die auf verschiedene Weltregionen verteilt sind. Die Inhalte werden zwischen diesen Rechenzentren repliziert (wie, das ist manchmal ein Geschäftsgeheimnis und grenzt an Zauberei). Der europäische Benutzer wird sicherlich nicht vom Fernen Osten aus bedient. Probieren Sie es aus: seit Kurzem liefert die Suchmaschine von Google sogar Zwischenergebnisse, noch während Sie die Stichworte eintippen! Das könnte man „Cut-Through Searching“ nennen.

LISP

Die Firma Cisco schlägt zur Optimierung von Übertragungswegen einen weiteren Weg vor: Locator/ID Separation Protocol (LISP). Wie der Name schon sagt, besteht die Grundidee darin, zwischen den Informationen über die Lokation eines Knotens im Netz und den Informationen über die Identität des Knotens zu unterscheiden. Das bisherige IP Routing macht diese Unterscheidung nicht. Für das IP Routing ist die IP-Adresse Identifikationsmerkmal und Lokationsinformation in einem. LISP arbeitet mit zwei verschiedenen Datenbasen für die Wegewahl:

- Routing Locators (RLOCs) für die Wegewahl im globalen Netz
- End Point Identifiers (EIDs) für die Erkennung von Sessions zwischen Endgeräten

Jeder LISP-fähige Router ist in der Lage, EIDs auf RLOCs abzubilden. Die Ab-

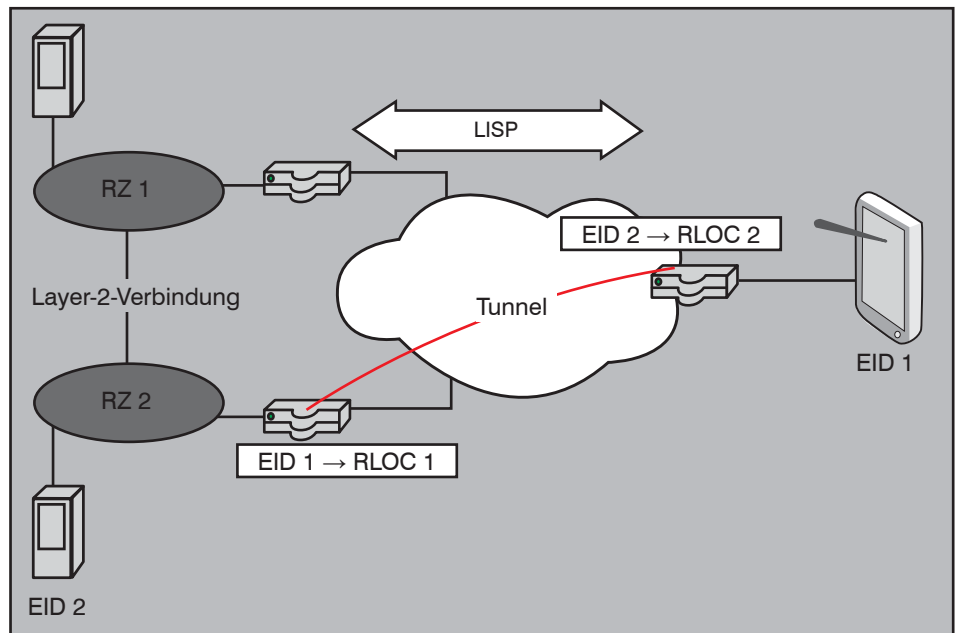


Abbildung 12: LISP

bildung erfolgt mithilfe von Informationen, die mittels LISP zwischen den Routern ausgetauscht werden. Diese bauen Tunnel zu anderen LISP-Routern auf. Beim Empfang eines Paketes an eine bestimmte EID ordnet ein LISP Router einen RLOC zur Ziel-EID zu, versieht das Paket mit der RLOC-Information und sendet das Paket durch den Tunnel zum Ziel-RLOC, d. h. zum LISP-Router in der Nähe der Ziel-EID. Dieser entfernt die RLOC-Information und leitet das Paket zur Ziel-EID weiter. Sowohl EID als auch RLOC entsprechen in ihrem Format herkömmlichen IP-Adressen. Das Schema ist in der Abbildung 12 dargestellt.

Wenn die Abbildung der EIDs auf die RLOCs unter Berücksichtigung der „Nähe“ eines LISP-Routers zum Ziel erfolgt, kann der Weg vom Client zum Server optimiert werden. Ohne LISP wären in der Abbildung 12 die Pfade über den oberen und den unteren linken Router aus der Sicht von IP Routing gleich berechtigt, sodass es durchaus dazu kommen kann, dass der ungünstige Weg über den oberen linken Router genommen wird. Mit LISP kann dafür gesorgt werden, dass der günstigere Weg beschritten wird. Dazu ist erforderlich, dass nicht nur WAN-Router, sondern auch Switches in den beiden Rechenzentren LISP unterstützen und LISP-Tunnel aufbauen können. Diese Switches können weiterhin reines Layer 2 Switching ausführen, jedoch mit der zusätzlichen Intelligenz versehen werden, LISP-Routen zu optimieren und LISP-Tunnel mit minimaler Latenz aufzubauen.

Bewertung von LISP

Die Latenzminimierung ist nur ein Seiteneffekt von LISP. LISP ist in erster Linie mit anderen Zielen entwickelt worden, zum Beispiel der Reduzierung der Größe von Routing-Tabellen, der Erleichterung eines Providerwechsels durch die providerunabhängige EID-Adressierung und einer sanften Migration von IPv4 zu IPv6.

Die Internet Engineering Task Force hat eine sehr aktive LISP-Gruppe⁵. Was allerdings auffällt, ist die überwältigende Präsenz von Cisco bei den Autoren der Drafts. Offensichtlich ziehen Cisco-Mitbewerber wie Juniper - wenn überhaupt - ein bisschen widerwillig mit.

Vor diesem Hintergrund ist noch nicht absehbar, ob sich LISP überhaupt durchsetzen wird und, wenn ja, wo die ersten Implementierungen stattfinden werden. In den letzten Jahren und Monaten hat die ehemals dominierende Marktmacht von Cisco nachgelassen, weshalb Cisco diesen neuen Ansatz nicht allein und nicht gegen massive Widerstände und Bedenken durchsetzen kann.

Es ist durchaus möglich, dass LISP zunächst innerhalb von Unternehmensnetzen Anwendung findet, denn LISP-Router sind zu herkömmlichen IP- Routern abwärtskompatibel. Ein LISP-Tunnel wird nur aufgebaut, wenn sich ein LISP-Partner finden lässt. Somit ist eine Mischung aus LISP- und herkömmlichen Routern, d. h. auch eine sanfte Migration möglich. Das ist ein wesentlicher Vorteil. LISP-Experimente können auch klein anfangen. Wenn sie die mit LISP verbundenen Ver-

⁵<http://tools.ietf.org/wg/lisp/>

Wie viel Delay ist tolerierbar?

sprechen einhalten, wird es die Nutznießer - das können Benutzer und IT-Verantwortliche in den Unternehmen sein - kaum interessieren, wie lange es bis zur abschließenden Standardisierung von LISP dauert. Erfolgreiche Technologien sind daran zu erkennen, dass sie sich schneller verbreiten als Standardisierungskomitees arbeiten.

Zusammenfassung

Dieser Beitrag behandelt aktuelle Entwicklungen mit dem Ziel der Minimierung von Latenzen in Netzen. Aus der Sicht des Autors sind Marketing Hypes und Substantielles in der aktuellen Latenzdiskussion vermischt. Der Beitrag ist ein Versuch, auf zwei verschiedenen Ebenen - auf der RZ-internen und auf der WAN-Ebene - die Sachverhalte und Konzepte ein-

zuordnen und zu bewerten. Während bei LAN Switches nach der Prognose des Autors die Diskussion um Ultra-Low Latency bald abebben wird, kann sich im Internet und WAN mit neuen Verfahren wie LISP einiges ändern. Hier ist auch der Handlungsbedarf zu sehen, wenn Cloud Computing, also der Zugriff auf Ressourcen „im Netz“ mit akzeptablen Antwortzeiten funktionieren soll.

Seminar

RZ-RZ-Kopplung - alles nur eine Frage der Bandbreite?

26.09.11 in Köln

Das Speichervolumen moderner Rechenzentren wächst schon seit Jahren exponentiell. Jetzt kommen Anforderungen synchroner Spiegelung und paralleler Datensicherungen in zwei Lokationen hinzu. Die Verwaltung der verteilten Speichersysteme geschieht dabei im Idealfall weiterhin zentralisiert.

All diese genannten Aspekte werden auf den Ebenen physikalische Verbindung, Netzwerk, Server, Speicher, Applikation unterschiedlich von den Herstellern adressiert. Dabei spielt die zur Verfügung stehende Bandbreite zwar auf der Kostenseite eine entscheidende Rolle. Wesentlicher für die Applikationen sind jedoch in der Regel die Latenz und die Funktionalität, die über die RZ-RZ-Verbindungsstrecke abgebildete werden können.

In diesem Seminar werden die aktuellen Techniken in diesen einzelnen Bereichen vorgestellt, technisch erläutert und für die richtige strategische Entscheidung zu einem Gesamtkonzept zusammengeführt.

Das Seminar geht unter anderem auf folgende Fragen ein:

- BSI / Basel III: wo kommen die Anforderungen zur RZ-RZ-Kopplung her und wie sehen sie aus?
- Dark Fiber, CWDM, DWDM: welche Kopplungsmöglichkeiten gibt es auf der physikalischen Ebene?
- Wie funktionieren diese Technologien? Welche Distanzen können hiermit überbrückt werden? Mit welchen Komponenten kann dies heute realisiert werden?
- Asynchrone / Synchroner Spiegelung: worin besteht der Unterschied? Welche Rolle spielt die physikalische Begrenzung durch die Lichtgeschwindigkeit? Welchen Einfluss hat das eingesetzte Übertragungsprotokoll?
- Wie sehen die in der Praxis ermittelten Datenübertragungsraten aus?
- Was bedeuten „Hochverfügbarkeit“ und „Fehlertoleranz“ auf Ebene der Server-Dienste? Welche technischen Anforderungen ergeben sich daraus?
- Welche Rolle spielt die Verschlüsselung auf den Weitverkehrsstrecken? Welche Techniken stehen hierfür zur Verfügung?
- Wie sind die aktuellen Netzwerk-Techniken Shortest Path Bridging (SPBM), Transparent Interconnection of Lots of Links (TRILL) sowie proprietäre Entwicklungen in diesem Bereich der Layer-2-Ausdehnung einzuschätzen (z.B. Cisco FabricPath, Juniper QFabric)?
- Welche weiteren Herausforderungen ergeben sich aus der Präsenz gleicher Subnetze an unterschiedlichen geografischen Standorten? Wie sind aktuelle Entwicklungen zu beurteilen, die diese Schwierigkeiten adressieren (z.B. Cisco Location ID Separator Protocol, LISP)?
- Wie geeignet sind die Speicherübertragungsprotokolle Fibre Channel (FC), iSCSI, Fibre Channel over Ethernet (FCoE), Network File System (NFS) für die Replikation von Speicherinhalten?
- Welche Anforderungen bringen Server- und Datenbank-Cluster wie der Oracle Real Application Cluster?
- Welche Möglichkeiten zur Lastverteilung gibt es über mehrere RZ-Standorte (z.B. Global Server Load Balancing)?
- Wie können diese technischen Maßnahmen geeignet in einer Leitlinie zum Betrieb redundanter RZ-Standorte dargestellt werden?
- Für welche Systeme sollte eine synchrone Spiegelung vorgesehen werden? Wo reicht Asynchronität aus? Was muss überhaupt nicht gespiegelt werden?

Referent: Dipl.-Inform. Matthias Egerland
Preis: € 890,- zzgl. MwSt.



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

Auf die Prozesse kommt es an

Der Standpunkt Sicherheit von Dr. Simon Hoff 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.

Inzwischen vergeht kaum noch einen Tag ohne neue Berichte über erfolgreiche Hacks, gefundene Datenlecks in Unternehmen und Behörden, kritische Schwachstellen in IT-Systemen, etc. Das Internet scheint zum kriegesischen Dschungel verkommen sein, in dem mehr oder weniger professionelle Söldner, Partisanen und Hobby-Krieger sich austoben. Die Statistiken der monatlichen Angriffe sprechen hier für sich und mit den bekannt gewordenen Sicherheitsvorfällen sehen wir nur die Spitze des Eisbergs.

Schauen wir uns einige Vorfälle genauer an: Vor einigen Tagen wurde in den Medien berichtet, dass Hacker vertrauliche Daten des Zolls erbeuten und Polizeirechner mit einem Trojaner infizieren konnten. Das Datenleck soll über Monate bestanden haben, ohne dass der Datenabfluss aufgefallen ist. Bei Rewe konnten Angreifer auf Kundendaten inklusive unverschlüsselter Passworte zugreifen. Die Hackergruppe Anonymous hat kürzlich bei Booz Allen Hamilton über einen unzureichend abgesicherten Server zehntausende militärische Zugangsdaten (E-Mail-Adressen und Passworte) erbeuten können. Die Passworte seien den Berichten zu folge hier zwar verschlüsselt gewesen, jedoch nicht mit besonders guter Qualität. Die Hacker waren dabei insbesondere über die Leichtigkeit des Angriffs überrascht.

Diese Liste ließe sich wahrscheinlich endlos fortsetzen, die paar Beispiele genügen jedoch, um festzustellen, dass Informationen mit offensichtlich hohem Schutzbedarf scheinbar öfter unzureichend geschützt (im Extremfall sogar im Klartext) auf vom Internet zugänglichen Rechnern liegen als es einem lieb ist.

Nehmen wir an, dass die angegriffenen Parteien selbstverständlich klare interne Sicherheitsrichtlinien haben, die festlegen, wie mit Daten mit einem hohen Schutzbedarf hinsichtlich der Vertraulichkeit (und anderer Grundwerte) umzugehen ist. Wiese so lagen dann in den genannten Beispielen vertrauliche Daten nur unzureichend abgesichert vor? Hier ist offenkundig bei der Umsetzung von Sicher-



heitsrichtlinien etwas schief gelaufen.

An dieser Stelle brauchen wir weniger neue Techniken, neue Bollwerke der Informationssicherheit in Form von noch mehr Firewalls oder anderen Werkzeugen. Wir brauchen eine effektive Verankerung der Informationssicherheit in den IT-Prozessen und in den Geschäftsprozessen.

Dies betrifft insbesondere die IT-Kernprozesse Change Management, Incident Management und Problem Management.

Wenn beispielsweise ein Change ansteht, bei dem eine neue Web-Anwendung eingeführt werden soll, könnte es doch eine Selbstverständlichkeit sein, dass im Rahmen des Change-Prozesses die Informationssicherheit eine Prüfung der Anwendung durchführt und erst nach Freigabe durch den Sicherheitsbeauftragten die Anwendung in Produktion geht. Wenn man dies ernsthaft betreibt, würde sofort auffallen, wenn z.B. Daten mit hohem Schutzbedarf nicht angemessen geschützt transportiert oder gespeichert werden sollen.

Außerdem brauchen wir lebendige Prozesse für die Informationssicherheit selbst. Hierzu gehört neben der Überwachung der Infrastruktur hinsichtlich Sicherheitsvorfällen die regelmäßige Prüfung der Infrastruktur hinsichtlich Schwachstellen, Umsetzung von Sicherheitsmaßnahmen sowie die Messung und Bewertung des erreichten Sicherheitsniveaus.

Dies erfordert eine schlagkräftige Sicherheitsorganisation und insbesondere einen Sicherheitsbeauftragten, der nicht nur strategische Richtlinienkompetenz hat, sondern auch im operativen Geschäft verankert ist. Natürlich kostet das Zeit und Geld. Nur was kann ein Schaden kosten?

Seminar

Sicherer Internetzugang, 07. - 09.09.11 in Aachen

Dieses Seminar identifiziert die wesentlichen Gefahrenbereiche und zeigt effiziente und wirtschaftliche Maßnahmen zur Umsetzung einer erfolgreichen Lösung auf. Alle wichtigen Bausteine werden detailliert erklärt und anhand praktischer Projektbeispiele und Übungen wird der Weg zu einer erfolgreichen Sicherheits-Lösung aufgezeigt.

In diesem Seminar lernen Sie

- wie aktuell die wichtigsten Bedrohungen aussehen und wie diese systematisch zu kategorisieren sind
- welche Kernbausteine eines sicheren Internetzugangs sich aus der Bedrohungslage ergeben
- wie Security Gateways (insbesondere Firewalls) arbeiten, welche Typen es gibt und wie Einsatzszenarien, Aufbau- und Betriebskonzepte aussehen
- worauf bei sicheren DMZ-Architekturen und Konfigurationen für Firewalls zu achten ist
- wie sich erweiterte Sicherheitsfunktionen wie IPS und Content Security integrieren lassen
- wie Sie Kommunikationsprotokolle und Netzwerke sinnvoll absichern
- was kryptografische Verfahren leisten, welche Verfahren wie sicher sind und wo diese Verfahren zum Einsatz kommen
- welche Methoden zur Absicherung von E-Mail-Kommunikation und Web-basierten Applikationen existieren

Referenten: Dipl.-Inform. Oliver Flüs, Dipl.-Inform. Andreas Meder

Preis: € 1.890,- zzgl. MwSt.



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

Zweitthema

Cloud Storage in der Analyse

Fortsetzung von Seite 1



Dr. Jürgen Suppan gilt als einer der führenden Berater für Kommunikationstechnik und verteilte Architekturen. Unter seiner Leitung wurden in den letzten 25 Jahren diverse Projekte aller Größenordnungen erfolgreich umgesetzt. Sein Arbeitsschwerpunkt ist die Analyse neuer Technologien und deren Nutzen für Unternehmen. Er leitet das internationale Labor von ComConsult-Research in Christchurch, das die Technologieentwicklung in Asien, Australien, den USA und Europa analysiert und für Kunden bewertet. Gleichzeitig ist er Inhaber der ComConsult Akademie, der ComConsult Technologie Information GmbH und der ComConsult Technology Information Ltd.

Dafür gibt es eine Reihe guter Gründe:

- Die Performance und Qualität der Speicher-Lösung ist von zentraler Bedeutung für die Gesamt-IT
- Die weitere Zentralisierung von Speicher und Rechenleistung schafft neue Herausforderungen
- Die Leistungsschere zwischen Prozessor- und Speicher-Leistung öffnet sich rasant schnell
- Speicher hat weiterhin ein starkes Wachstum
- Es gibt ein zunehmendes Problem der Managebarkeit von Speicher
- Backup- und Restore sind kaum noch handhabbar
- Es existieren sehr unterschiedliche Anforderungen an Speicher-Leistung seitens der Applikationen

Da hilft es wenig, wenn die Preise pro Gigabyte fallen. Die Summe aller Anforderungen generiert drängende Fragen zur Wirtschaftlichkeit und zum Betrieb von Speicher-Lösungen. Die Kostenprobleme liegen mehr und mehr im Betrieb und nicht im Einkauf. Die IT-Verantwortlichen sind hier klar unter Druck.

Hier setzt Cloud Storage an. Er verspricht: (siehe Abbildung 1)

- Beliebige Skalierbarkeit
- Sofortige Verfügbarkeit / Abrufbarkeit
- Klare Kostenstrukturen mit hoher Wirtschaftlichkeit

- Hohe Verfügbarkeit durch Speicherung in mehreren geografisch getrennten Regionen

Im Folgenden wollen wir deshalb die Frage analysieren, ob Cloud Storage in der Tat die bestehenden Anforderungen und Probleme der Unternehmen im Bereich Speicher erfüllt. Wir werden dazu eine Reihe von Kriterien entwickeln und auch priorisieren. Hier liegt in der Tat der Knackpunkt jeder Analyse zu Cloud Storage: in der Festlegung der Ziele und der Bewertung mit einer geeigneten Kriterienliste.

Ausgangslage

In der Vergangenheit war die Welt der Speicher klar getrennt in Direct Attached Storage DAS auf der Serverseite und Storage Area Networks SAN auf der Seite der großen Datenbank-Applikationen. Jeder Server hatte seine spezifische Aufgabe und sein DAS konnte der Aufgabe entsprechend ausgelegt werden. Auch das SAN konnte granular am Bedarf der Datenbank-Applikation optimiert werden, sei

es in Richtung Transaktionsrate oder im Bereich Multiusergrad oder der Verfügbarkeit. Auch diese traditionelle Situation war nicht frei von Problemen. Ein kontinuierliches Wachstum verbunden mit zunehmenden Backup- und Restore-Problemen speziell bei multinationalen Unternehmen mit kleinen Zeitfenstern sind die Basis vieler Diskussionen der letzten 10 Jahre.

Mit der Server- und Speicherkonsolidierung, die speziell mit dem Aufkommen von Virtualisierungstechnik eingesetzt hat, haben sich die Rahmenbedingungen für das Design einer guten Speicher-Lösung deutlich verändert. Der alte Zustand mit DAS am Server war auf Dauer nicht zu halten. Er war identisch mit einer gigantischen Verschwendung von Ressourcen, fehlender Flexibilität und Skalierbarkeit sowie völlig unkontrollierbaren und teuren Betriebsabläufen. Die logische Konsequenz ist die Trennung von Server und Speicher. Dabei werden die bisher verteilten DAS-Speicher zentralisiert und die Server virtualisiert (und zum Teil auch konsolidiert, so dass in Summe auch eine kleinere Zahl von virtuellen Maschinen gegenüber den bisherigen physikalischen Maschinen entsteht).

So einfach und plausibel diese Entwicklung klingt, sie hat ihre Tücken. Auf der Virtualisierungsseite führt der Bedarf nach Verfügbarkeit und Skalierbarkeit zum Aufbau von virtuellen Infrastrukturen, innerhalb deren sich virtuelle Maschinen nach Lastbedarf verschieben können. Auch werden VMs dynamisch hoch und runter gefahren, je nach Architektur der Applikationen. Dies führt auf der Speicherseite zu gewissen Fragezeichen in der Auslegung der Systeme. Zum einen kann sich hinter einer virtuellen Maschine alles Mögliche ausgehend von einem einfachen File Server über einen Email-Server hin bis zu

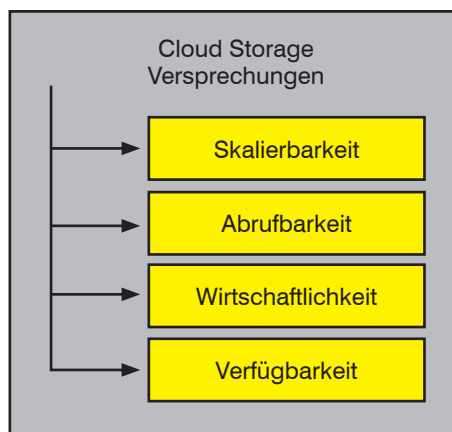


Abbildung 1: Cloud Storage Versprechen

Cloud Storage in der Analyse

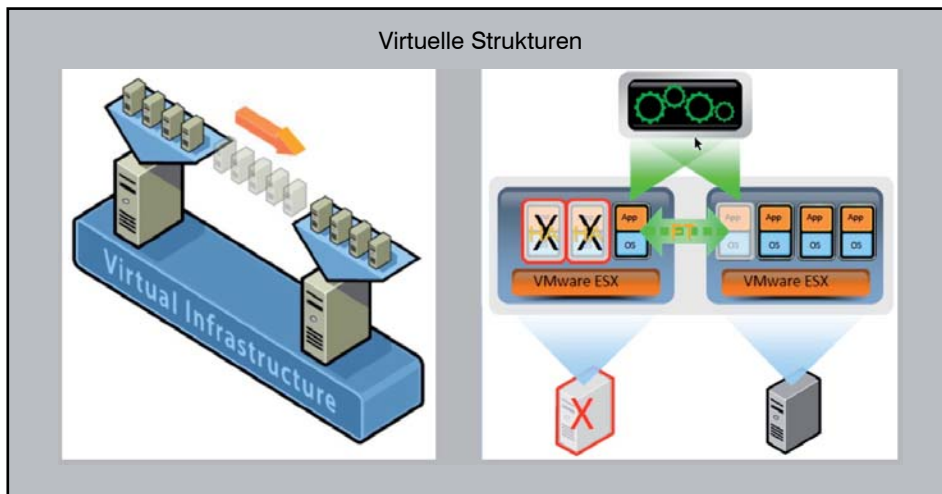


Abbildung 2: Virtuelle Infrastrukturen generieren neue Herausforderungen an Speicher
Quelle: VMware

einer mittelgroßen Datenbank-Applikation verbergen. Zum anderen bringen dynamische Lokationen Folgeprobleme mit sich, wenn wie bisher getrennte Netzwerke für Daten und Speicher aufgebaut werden müssen. (siehe Abbildung 2)

Aber auch auf der Seite des Speichers selber entsteht eine völlig neue Ausgangslage. So macht es kaum Sinn, den Bedarf eines einfachen File Servers mit einem High-End SAN abzudecken, das auf hohe Transaktionsraten und einen hohen Multiusergrad einer komplexen Datenbank-Anwendung ausgelegt ist. Die Kosten für Speicher würden dabei explodieren und die Ziele der Server-Konsolidierung würden auf der Speicher-Seite zerschlagen. Zwar gab es auf der Speicherseite immer schon den Bedarf eines Multitiering und des Aufbaus von Policies, um Daten entsprechend ihrer Anforderung im optimalen Tier zu speichern. Aber noch nie war die Situation so komplex wie heute.

Vereinfacht ausgedrückt hat die Konsolidierung von Servern und Speichern den Bedarf nach einem neuen Typ von Speicher generiert, der irgendwo zwischen DAS und SAN angesiedelt sein muss. Theoretisch gibt es diese Speicherformen schon lange in Form von iSCSI, NAS, NFS..., aber in Kombination mit Skalierung und Betrieb entstehen neue Herausforderungen.

Insgesamt kann man den aktuellen Bedarf nach Speicher in folgende Anforderungen unterteilen: (siehe Abbildung 3)

- Preiswert unter Einbeziehung der Betriebskosten
- Einfach zu handhaben und zu konfigurieren

rieren

- Effizient, Vermeidung unnötiger Reserven
- Leistungsverhalten optimiert am Bedarf

Anders formuliert, kann man diese Anforderungen auch anders darstellen:

- Die Lösungspreise pro TB oder PB müssen inklusive Betriebskosten fallen
- Wilde Mischungen aus iSCSI, NAS und SAN müssen vermieden werden, sie erhöhen die Betriebskomplexität
- Eine statische Zuordnung von Speicher zu VMs muss unterbleiben
- Es müssen pro Speicher mehr Benutzer als bisher mit besseren Antwortzeiten und mehr Transaktionen abgewickelt werden können

Die Lösungs-Perspektiven

Betrachtet man die Anforderungen, dann wird klar, dass eine Reduzierung auf den Preis pro TB am Bedarf vorbei geht. Betriebskosten, Flexibilität und Komplexität sind Bestandteile der Gleichung, die nicht unterschlagen werden dürfen. Auch die schnelle Gestaltbarkeit der Lösung ist ein wichtiges Thema. Betrachten wir die Geschwindigkeit, mit der VMware und Citrix ihre Architekturen verändern, dann ist auch klar, dass die Speicherseite flexibel bleiben muss. Ein statisches Binden an eine feste Lösung kann über mehrere Jahre hinweg nicht funktionieren. Auch ist auf der Serverseite weiterhin viel Dynamik zu beobachten. Intel hat gerade mit der Einführung der 22nm-3D-Transistoren ein neues Kapitel der Technologie aufgeschlagen, das so-

wohl auf der Seite der Server als auch der Speicher Änderungen nach sich ziehen wird. Bereits im typischen Lebenszeitraum einer Speicherlösung wird der Wechsel auf 12nm-Technologie kommen. Vergleicht man das mit der Situation von vor 3 Jahren, haben wir damit bereits wieder völlig veränderte Rahmenbedingungen.

Trotzdem kann man wichtige Elemente eines Lösungspaketes identifizieren:

- Im Vordergrund steht das Speicher-Management, angefangen von der Konfiguration und Einbindung in die virtuelle Welt von Citrix und VMware und endend in der Überwachung und dynamischen Optimierung
- Auto-Tiering wird auf jeden Fall eine große Rolle spielen
- Self-Provisioning, so dass VMs dynamisch den Speicher nach ihrem Bedarf skalieren können, ist ein Muss
- Dies geht einher mit Thin-Provisioning. Server dürfen keine unnötigen Ressourcen blockieren
- Erhöhung der Leistung durch Einbindung neuer Technologien wie SSD verbessert die Handhabbarkeit, da die Zahl der notwendigen Speicher-Systeme reduziert werden
- SSD-Technologie wird mit dem Wechsel auf 22nm und später 12nm-Prozesse eine stark zunehmende Bedeutung in der Speichertechnologie haben. Sie erhöhen die mögliche Zahl von Servern pro Speichersystem in Kombination mit einer deutlichen Steigerung der möglichen Benutzer und Applikationen
- Kompression und Deduplizierung reduzieren den Bedarf an Speicher und auch die Anzahl der notwendigen Spei-

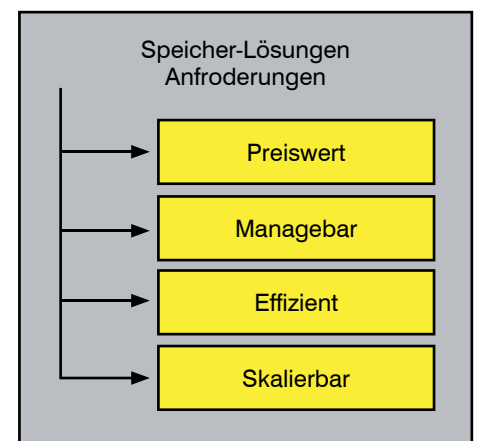


Abbildung 3: Anforderungen an Speicher

Cloud Storage in der Analyse

cher-Systeme. Zwar hängt die generell von dem Typ der vorhandenen Daten ab, doch gerade im Umfeld der zu konsolidierenden Server werden typischerweise viele Daten vorliegen, die mit hoher Wirkung komprimiert und dedupliziert werden können (diese Technologien haben ihre Tücken und erfordern ein detailliertes Verständnis auch der verschiedenen möglichen Varianten, an dieser Stelle sei auf den Artikel von Matthias Egerland im Netzwerk-Insider Ausgabe Mai hingewiesen)

- Flexible Schnittstellen angefangen in der freien Wahl zwischen Block- und File-Zugriff und endend im freien Einsatz des Protokolls, sei es iSCSI oder FC
- Bandbreite und Latenz im Zugriff auf die Speichersysteme werden immer wichtiger. Simple Gigabit-Technik im RZ-LAN und zur Anbindung von Speicher ist tot, jede Investition in diese Richtung ist pure Geldverschwendung. 10 Gigabit Ethernet und in den nächsten Monaten zunehmend 16 Gigabit FC sind das Maß der Dinge. Auf Dauer wird auch das nicht skalieren. Der Schritt nach 100 Gigabit wird in 3 bis 5 Jahren erfolgen.

Anders formuliert, egal wie die Lösung aussehen wird, sie wird immer ein Paket sein, das aus unterschiedlichen Komponenten geschnürt wird, um den Bedarf eines Unternehmens wirtschaftlich optimal abzudecken.

Ist die Forderung nach flexiblen Schnittstellen gleichbedeutend mit dem Verzicht auf reine iSCSI oder NAS-Lösungen? Nein, solange die eingesetzten Systeme einen signifikanten Umfang und eine strategische Größe erreichen. Wildwuchs speziell mit kleinen Systemen unter dem vorgeschobenen Argument der Kosteneinsparung sollte vermieden werden. Auf jeden Fall sollten auch die iSCSI oder NAS-Lösungen die anderen aufgeführten Anforderungen erfüllen. Auf keinen Fall darf ein Systemmangel bei iSCSI wie eine nicht ausreichende Verfügbarkeit zu einem unnötigen zusätzlichen Tier in der Gesamtlösung führen. Es gibt aber eine Reihe von Produkten wie EMC VNXe oder HP P4000 (früher LeftHand Networks), die diese Anforderungen erfüllen. Speziell HP zeigt mit dem P4000-System, dass iSCSI-Lösungen sehr wohl einen hohen Grad an Gestaltbarkeit und Skalierbarkeit haben können. (siehe Abbildung 4)

Die Nutzung von Kompression und Deduplizierung gewinnt im Markt immer mehr an Bedeutung. Speziell bei den Entry-Level-Lösungen muss abgewogen werden, ob eine Einsparung von wenigen Tausend

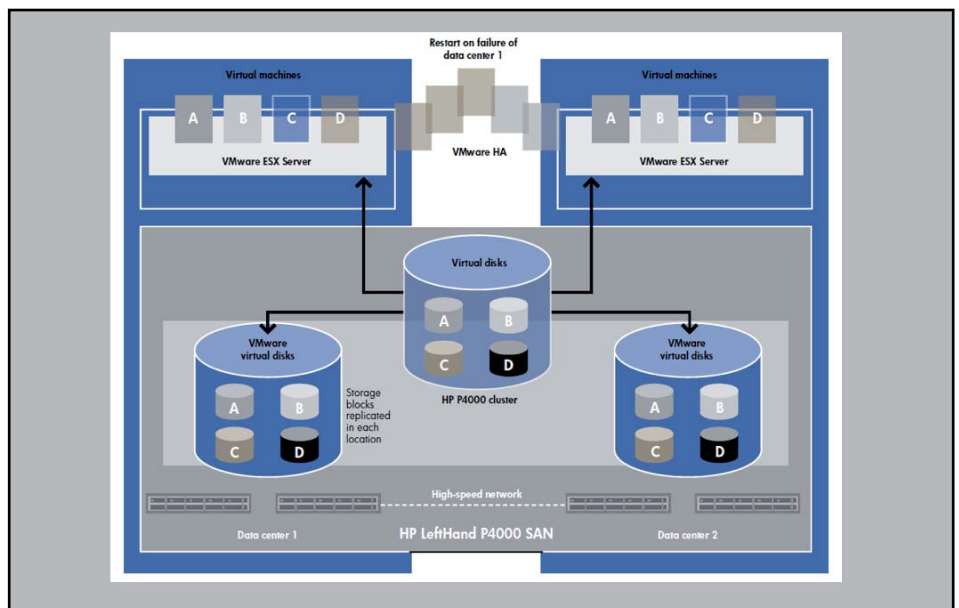


Abbildung 4: HP P4000, Beispiel für eine strategische iSCSI-Lösung

Euro pro System den Verzicht auf diese Funktionen rechtfertigt. EMC hat hier mit VNXe die Hürde für den Einstieg in diese Technologien deutlich gesenkt. (siehe Abbildung 5)

Wo ordnet sich jetzt Cloud Storage ein?

Um diese Frage zu beantworten, müssen wir die Cloud Storage Angebote mindestens in zwei Leistungsklassen unterteilen:

- Basis-Speicher
- Value-Added Speicher

Wie immer man die Bewertung vornimmt, es muss klar bleiben, über welchen Typ von Angebot man ein Urteil abgibt.

Beispiele für die sehr unterschiedlichen Angebote sind:

- Die Mutter aller Cloud-Storage Lösungen: der Amazon Simple Storage Service S3
- Das HiDrive Angebot der Telekom-Tochter Strato
- Nirvanix mit seinem weltweit verteilten Speichernetzwerk
- Value-Added Angebote wie Dropbox und SugarSync
- Angebote, die irgendwie dazwischen liegen wie Box.net

Kongress

Rechenzentrum Infrastruktur-Redesign Forum 2011 - 07. - 10.11.11 in Königswinter

Das ComConsult Rechenzentrum Infrastruktur-Redesign Forum 2011 ist die zentrale Veranstaltung des Jahres, in der hochqualifizierte Referenten die neuen Einflussfaktoren, Möglichkeiten und Strategien vorstellen und mit den Herstellern diskutieren, um Anregungen und praktische Hinweise für die Gestaltung des Rechenzentrums der jetzt beginnenden Zukunft zu erarbeiten.

Moderation: Dr. Franz-Joachim Kauffels, Dr.-Ing. Behrooz Moayeri
Preis: € 2.190,--* zzgl. MwSt. bzw. € 1.790,--* zzgl. MwSt. ohn Workshop

*Preise gültig bis zum 31.08.11 -
dann regulär € 2.390,- zzgl. MwSt. bzw. € 1.990,- zzgl. MwSt. ohne Workshop



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

Cloud Storage in der Analyse



Abbildung 5: Das Imperium schlägt zurück, EMC greift den Entry-Level-Markt mit VNXe an

Wie kann man so unterschiedliche Angebote sauber bewerten? Zum einen muss man klar sehen, dass im Sinne der klassischen Speicher-Lösungen nur die Basis-Speicherdienste wie Amazon, HiDrive oder Nirvanix eine Rolle spielen. Value-Added Angebote müssen wir davon trennen und gesondert bewerten. Zum anderen müssen wir ausgehend von der beschriebenen Bedarfssituation die Kriterien benennen, die uns wichtig sind.

Unsere Kriterien

Wir haben unsere Bewertungs-Kriterien nach ihrer Wichtigkeit geordnet. Die daraus entstandene Reihenfolge mag auf den ersten Blick überraschen, bei näherer Betrachtung sollten die Gründe für diese Wichtung aber klar werden. (siehe Abbildung 6)

Priorität 1: Service- und Management-Integration / Prozess-Integration

Mit dem Wechsel von monolithischen hin zu vernetzten System-Architekturen steigen die Abhängigkeiten zwischen den Technologien und der Bedarf an Verfügbarkeit für die einzelnen Knoten steigt gegenüber der nicht-vernetzten Lösung deutlich an (da sich Verfügbarkeiten in einem Netzwerk multiplizieren). Die Frage also, wie der Betriebszustand einer Technologie überwacht werden kann, wie im Störfall eine Ursache festgestellt werden kann und wie die Auswirkung auf Geschäftsprozesse ermittelt wird, ist dementsprechend elementar. Beispielsweise muss bei einer Anfrage an den UHD und einer Beschwerde über schlechte Performance in jedem Fall weiterhin eine geord-

nete Ursachenermittlung möglich sein.

Eine Cloud-Lösung egal welcher Art löst deshalb auch nicht bestehende interne Betriebsabläufe komplett ab. Im Gegenteil generiert sie neue Schnittstellen hin zu den etablierten Betriebsprozessen, die sauber definiert und abgedeckt werden müssen.

Damit entstehen eine Reihe technischer und organisatorischer Anforderungen an die Integration von Cloud-Storage:

- Die externe Lösung muss je nach Art der Nutzung in die interne technische Management-Lösung als Element-Manager integriert werden können. Störungen in der Cloud müssen in qualifizierter Form automatisch an die Unternehmens-Konsolen gemeldet werden

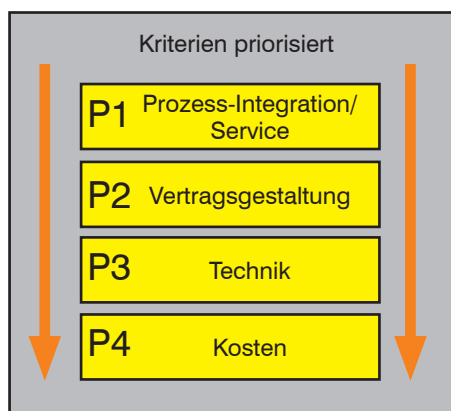


Abbildung 6: Bewertungs-Kriterien

- Es muss eine organisatorische Schnittstelle geben, um Störungen analysieren zu können, insbesondere wenn Ursachen nicht eindeutig sind. Speziell Performance-Störungen erfordern mehr als nur eine Black-Box-Lösung zur Evaluierung von Ursachen und Betroffenheit
- Es muss einen immer erreichbaren externen Service-Ansprechpartner geben und klar zugesagte Response-Zeiten (siehe ITIL für weitere Details), dies muss schriftlich fixiert sein und gepflegt werden

Die hohe Bedeutung dieses Kriteriums wurde durch den Ausfall bei Amazon überdeutlich. Die Kommunikations-Politik von Amazon auch größeren Kunden gegenüber lässt bisher keine geordneten Service-Abläufe zu. Amazon steht damit nicht alleine. Service ist bei vielen Cloud Providern ein düsteres Kapitel. Aus unserer Sicht sind Cloud Lösungen, die nicht klar und geordnet in wohldefinierte Serviceabläufe eingebunden werden können, nicht nutzbar. Wir empfehlen dringend im Rahmen einer Evaluierung auch die Service-Qualität zu testen.

Unverzichtbare Mindestanforderungen an die Service-Schnittstelle des Cloud-Providers sollten sein:

- Verfügbarkeit einer telefonischen Hot-Line
- Email und/oder Chat-Service
- Definierte Abläufe im Krisenfall

Priorität 2: Vertragsgestaltung

Wie immer, wenn Leistungen nach Außen vergeben werden, muss der Inhalt der Leistung und die Art der Leistungserbringung in einem Vertrag klar festgelegt werden. Je nach Art der Cloud Storage-Lösung stellt sich damit eine klare Frage:

- Reicht ein nicht auf den Bedarf des eigenen Unternehmens angepasster Standard-Vertrag von der Webseite des Erbringers
- Ist eine individuelle Vertragsgestaltung erforderlich

Es liegt nahe anzunehmen, dass, wenn die Cloud Storage Lösung nur irgend eine Relevanz für das Unternehmen hat, ein Standardvertrag nicht ausreichen wird. Speziell unsere Anforderung mit Priorität 1 erfordert die schriftliche Festlegung von Service-Schnittstellen und Ablaufprozeduren. Die im Web üblichen Standard-Verträge decken das in der Regel nicht ab. Hat die Cloud-Lösung keine Relevanz, dann

Cloud Storage in der Analyse

muss die Frage gestellt werden, warum sie überhaupt angegangen wird.

Mit der Vertragsgestaltung gehen eine Reihe weiterer Detailfragen einher. Spezielle Anforderungen an den Vertrag entstehen durch eine mögliche Speicherung von Personendaten. Aber auch so simple Fragen wie ein Gerichtsstand im Falle der Auseinandersetzung müssen beachtet werden.

Die meisten der von Cloud Providern angebotenen Standard-Verträge beinhalten eine einfache Form von SLA in Kombination mit der Zusage von Verfügbarkeiten. Diese Angaben sind weitestgehend wertlos, wenn die mit einem Verstoß gegen die Zusagen verbunden Kompensationen keinen wirklichen Wert darstellen.

In jedem Fall sollte die Frage des geeigneten Vertrags mit einem Cloud Provider unter Hinzuziehung eines auf diesen Bereich spezialisierten Anwalts diskutiert werden.

Ohne Frage sollte ein europäischer Anbieter, besser noch ein deutscher, mit lokalem Gerichtsstand und individueller Vertragsgestaltung bevorzugt werden.

Priorität 3: Technische Fragen

Die viel diskutierte Technik von Cloud-Lösungen taucht bei uns erst als Priorität 3 auf. Das hat den einfachen Grund, dass wir die Prioritäten 1 und 2 als KO-Kriterien sehen. Werden sie nicht erfüllt, macht ein Einstieg in technische Diskussionen aus unserer Sicht keinen Sinn.

Die technischen Fragen müssen für Basis-Dienste und Value-Added Dienste getrennt betrachtet werden. Naturgemäß wird bei den Value-Added-Diensten der jeweilige „Value“ eine zentrale Rolle spielen müssen.

Für die Basis-Dienste unterteilen wir die technischen Fragen in die Bereiche

- Generelle Eigenschaften
- Performance
- Verfügbarkeit
- Sicherheit

Auch wenn sich die Anbieter im Detail unterscheiden, treten sie im Kern mit ähnlichen Angeboten auf:

- Wählbare Anzahl von Admin und User-Accounts mit sehr verschiedenen ausgeprägter Verwaltung der individuellen Benutzerrechte
- NAS-vergleichbarer Speicher mit Zugriff über SMB/CIFS, FTP/SFTP/FTPS, Web-

DAV, Rsync oder spezielle Clients

- Vereinzelt Einbindung in die Clients als Laufwerk (Beispiel: OpenDrive für Windows), darüber dann Nutzung von Standard-Software für Backups etc.
- Bei den großen Anbietern auch Unterstützung der 3rd-Party-Clients (zum Beispiel Zugriff auf Amazon S3 über gängige FTP-Clients wie Transmit unter MAC OS)
- Spezielle Apps für Mobile Endgeräte (in der Regel Apple iOS und Android)
- Shared Links für Dateien und/oder Folder
- Zugriff über verschlüsselte Kommunikation, zum Beispiel SSL-256 oder SSH
- Vereinzelt Email-Upload oder DVD Send-In-Service
- Option von Backups mit zusätzlich wählbaren Parametern (Zeitpunkte, Aufbewahrungsdauer, Versionen-Verwaltung, ...)
- Option für höhere Verfügbarkeit durch Speicherung auf mehreren Knoten in regional verteilten Lokationen gegen Aufpreis (siehe Nirvanix)
- Option für Multi-Media-Fähigkeit (siehe Amazon, wobei die Option unsinnig ist, wenn die Performance stimmt und sowieso alle Datentypen unterstützt werden, Beispiel Strato)
- Unterstützung der Speicherung bereits verschlüsselter Folder (zum Beispiel

Truecrypt)

- Zugang naturgemäß über das Internet, aber zum sinnvollen Betrieb sollten geeignete Bandbreiten bereitstehen. Einige Anbieter werden vermutlich auch in der Lage sein, Gigabit-Links zu füllen. Optionen zum Bandbreiten-Management entweder durch den Client oder auf der Seite des Internet-Routers sind dringend zu empfehlen
- Synchronisation ist typischerweise eher im Bereich der Value-Added-Dienste wie Dropbox oder Sugarsync zu finden, vereinzelt bieten aber auch die Anbieter der Basis-Dienste diese Option (Beispiel Strato)
- Früher gab es zum Teil Beschränkungen in der maximalen Dateigröße, der maximalen Upload-Kapazität pro Tag oder den unterstützten File-Typen. Dies hat sich überholt und normalerweise sollte der Service keine Limitierungen beinhalten

Einen guten Überblick über „typische“ Angebote findet sich auf den Webseiten von Amazon Web Services AWS und Strato. Speziell Strato ist als deutscher Anbieter und als Tochter der Telekom ein interessanter Anbieter (auch wegen der erreichbaren Performance und der kürzeren Signallaufzeiten, siehe später unter Performance), der die Details seines Angebots in den letzten Monaten auch immer weiter ausgebaut hat. Zur Zeit fehlt noch eine Apple iOS-APP, diese ist aber angekündigt (allerdings muss sie auch mit den anderen iPad-Applikationen kombiniert werden können). (siehe Abbildung 7)

Kongress

Rechenzentrum Infrastruktur-Redesign Forum 2011 - 07. - 10.11.11 in Königswinter

Das ComConsult Rechenzentrum Infrastruktur-Redesign Forum 2011 ist die zentrale Veranstaltung des Jahres, in der hochqualifizierte Referenten die neuen Einflussfaktoren, Möglichkeiten und Strategien vorstellen und mit den Herstellern diskutieren, um Anregungen und praktische Hinweise für die Gestaltung des Rechenzentrums der jetzt beginnenden Zukunft zu erarbeiten.

Moderation: Dr. Franz-Joachim Kauffels, Dr.-Ing. Behrooz Moayeri
Preis: € 2.190,--* zzgl. MwSt. bzw. € 1.790,--* zzgl. MwSt. ohn Workshop

*Preise gültig bis zum 31.08.11 -
dann regulär € 2.390,- zzgl. MwSt. bzw. € 1.990,- zzgl. MwSt. ohne Workshop



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

Cloud Storage in der Analyse



Abbildung 7: Strato HiDrive

Der Bereich Performance dominiert durch die Zugriffszeiten, die auf jeden Fall durch die Laufzeiten entstehen. Je nach Lokation des Rechenzentrums des Cloud Providers werden hier mehrere Hundert Millisekunden entstehen. Die lokalen Zugriffszeiten auf der Cloud-Seite selber werden dabei unerheblich sein. Damit ist Cloud Storage für die meisten Arten von Block-Zugriff im Rahmen von Datenbank-Applikationen ungeeignet (es gibt auch so gut wie keine Angebote in diesem Bereich). Auch Server-Applikationen sollten auf dieser Basis nicht installiert und betrieben werden (das ist nicht das Ziel dieser Services).

Somit kommen als Anwendungsbereiche primär in Frage:

- File-Zugriff für verteilte Standorte im Rahmen von Textverarbeitung, Email, ...
- Backup

Der Kern des Interesses an Cloud Speicher scheint momentan im Bereich Backup als zusätzlichen Tier zu liegen. Standort-übergreifende Datei-Speicherungen sind eher die Domäne von SharePoint oder ähnlichen Team-Speicher-Lösungen und werden dementsprechend eher im Value-Added-Bereich angesiedelt sein.

Die Nutzung von Cloud-Speicher als zusätzlichen Backup-Tier hat zudem den Vorteil, dass die Einbindung in normale Service-Prozesse an Bedeutung verliert, da keine laufenden Anwendungen betroffen sind und somit keine Service-Prozesse

unter Einbeziehung des Anwenders stattfinden.

Im Bereich Verfügbarkeit geben die Anbieter sowohl Verfügbarkeitswerte an als auch Optionen zur Erreichung einer noch höheren Verfügbarkeit. Da diese SLA-Angaben in der Regel nicht an wirksame Kompensationen gebunden sind, betrachten wir sie als weitestgehend wirkungslos. Speziell der Datenverlust bei dem im Mai aufgetretenen Amazon-Desaster macht klar, dass eine Cloud-Lösung ein lokales Backup nicht ersetzen kann. Da für viele Kunden aber die Cloud u.a. spannend ist, weil darüber ein neues Standort-neutrales Backup-Tier erreicht werden kann, widerspricht sich das natürlich. Maximal kann hier momentan eine Art Disaster Recovery Option gesehen werden.

Im Bereich Sicherheit wurde bei mehreren Anbietern speziell außerhalb der EU deutlich, dass es einen Zugriff auf die Daten seitens der Betreiber geben kann und auch muss. Das US-Recht sieht einen Zugriff auf Daten durch Gerichts-Beschluss vor. Von daher sind die amerikanischen Provider gezwungen, diesen Zugriff technisch möglich zu machen. Dies beinhaltet auch Provider und Dienste, die im Cloud Speicher eine Verschlüsselung unter Schlüsselverwaltung durch den Provider beinhalten. Wer die Sicherheit seiner Daten will, kommt an einer eigenen Verschlüsselung mit der Schlüsselverwaltung im Unternehmen nicht vorbei. Damit werden naturgemäß eine ganze Reihe von Clients wertlos, da beim Zugriff auf Daten erst eine Entschlüsselung erfolgen muss.

Betrachten wir Value-Added-Lösungen, so dienen diese in der Regel einem Spezialzweck (eben dem Value entsprechend). Dominant sind hier Synchronisations-Lösungen wie Dropbox und SugarSync, die Desktop-PCs und mobile Endgeräte integrieren. Insbesondere für die mobilen Endgeräte sind diese Dienste im Moment die erste Wahl bei der Integration ins Unternehmen. Von daher sieht auch die Frage der Kunden-seitigen Verschlüsselung momentan schlecht aus, da dazu ein geeigneter Client auf der Seite der mobilen Endgeräte vorhanden sein müsste. Dies schränkt die in diesem Dienstbereich nutzbaren Daten ein.

Aus unserer Sicht sollten bei der Auswahl derartiger Synchronisations-Produkte auf der technischen Seite beachtet werden:

- Eignung für den Unternehmenseinsatz. Viele dieser Produkte kommen aus dem Consumer-Markt, Funktionalitäten wie Benutzer-Verwaltung sind in der Regel sehr mager ausgeführt
- Benutzer-Verwaltung mit Einrichtung von Benutzer-Gruppen und Quotas
- Zugang über SSL-256 und Speicherung in verschlüsselter Form
- Wie werden Ordner zur Synchronisation ausgewählt, beliebig oder eine fest vorgegebenes Directory
- Versionenverwaltung: wie weit reicht die zurück
- LAN-Sync
- Ausblendung von Ordnern aus dem Sync für individuelle Geräte
- iOS und Android App
- Schnittstellen zu anderen Apps unter iOS und Android
- Sync-Performance
- Belastung der lokalen Ressourcen (CPU und Netzwerk-Bandbreite)
- File oder Folder-Sharing ohne den entsprechenden Client installieren zu müssen
- Aktivitäts-Nachrichten zu gesharten Ordnern
- Ereignis-Übersicht
- Zentraler File-Manager mit detaillierter Übersicht über Ordnergrößen

Cloud Storage in der Analyse

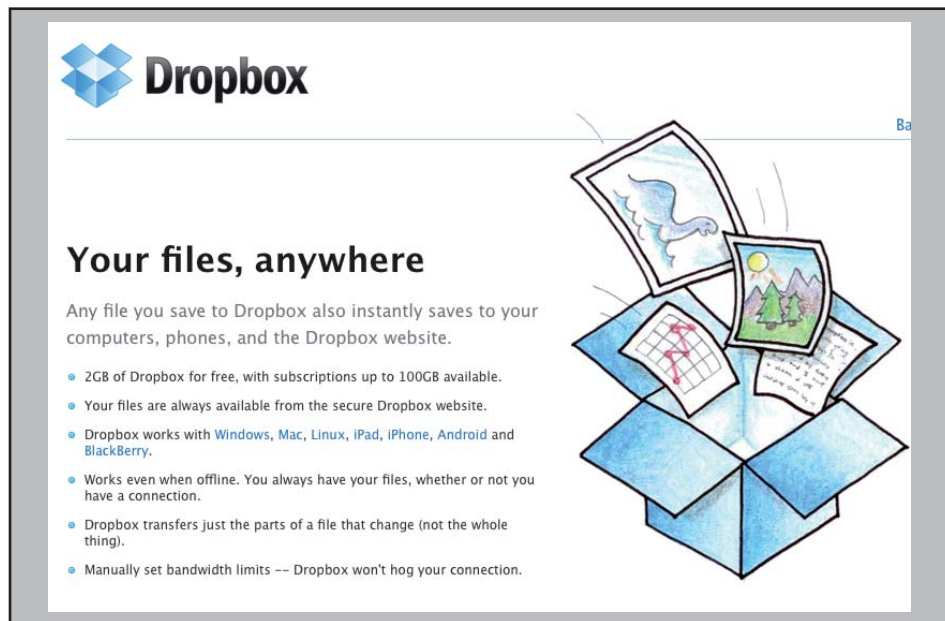


Abbildung 8: Marktführer im Bereich Value-Added: Dropbox

Ohne Frage wird Apples neuer iCloud-Dienst den Markt verändern, allerdings nur im Bereich der reinen Apple-Welt (nach bisherigen Erkenntnissen). Die automatische Synchronisation von Dokumenten, Fotos, Musik, Kalendern, ... inklusive der Bereitstellung der APIs für Dritt-Applikationen als kostenloser Dienst wird Dropbox, SugarSync und Co. unter Druck setzen. Apple bringt damit das lang ersehnte File-System für das iPad, aber eben als Cloud-Dienst.

Priorität 4: Der Preis

Einige Leser werden sicher von der Einstufung des Preises in unserer Prioritätenliste überrascht sein. Die Logik ist einfach. Sind die Prioritäten 1 bis 3 nicht abgedeckt, macht die Lösung keinen Sinn, sei sie auch noch zu preiswert.

Generell entsteht bei der Preis und Gesamtkostenbewertung das Problem bzw. die Frage, mit welchen internen Kosten hier verglichen werden soll. Das Spektrum der intern möglichen Speicher-Lösungen ist weit und reicht technisch von einfachen Open Source iSCSI- oder NFS-Lösungen bis hin zu komplexen SANs. Auch sollte jeder Vergleich nicht auf der Basis des Bestands erfolgen, sondern gegen aktuelle Produktangebote gefahren werden, da sich die Preis und die Produktdetails in den letzten Jahren dramatisch verändert haben. Ein typisches Beispiel ist hier die Preisgestaltung des neuen EMC VNX-Produktes.

Auch bei der Kostenbetrachtung muss ansonsten unterschieden werden zwischen Basis-Speicher-Diensten, die quasi Lauf-

werke bereitstellen, und Value Added Diensten. Wir starten mit der Betrachtung der Basis-Speicher-Dienste und gehen danach auf den Value-Added-Bereich ein.

Ein Kernproblem in vielen veröffentlichten Kostenvergleichen zwischen lokalen Diensten und Cloud Diensten liegt in der Annahme, dass lokale Dienste durch die Cloud abgelöst werden. Dies ist in der Regel nicht der Fall. Im Gegenteil bleiben die bisher vorhandenen Betriebsabläufe und Infrastrukturkosten in vielen Fällen beste-

hen, während durch die Integration des Cloud Dienstes neue Kosten entstehen. Diese liegen nicht nur im Cloud Dienst selber sondern speziell in der Integration des Dienstes in den Betrieb. Gleichzeitig bleiben auch die vorhandenen Infrastrukturen wie Netzwerke erhalten und somit auch deren Betriebskosten. Damit stellt sich die Frage, ob unter diesem Blickwinkel der Cloud-Dienst überhaupt wirtschaftlich sein kann.

Im Bereich Speicher beobachten wir nicht nur seit 10 Jahren und mehr einen kontinuierlichen Verfall der Preise pro TB, sondern gerade im Moment das Aufkommen einer neuen Produktklasse, die speziell auf die Anforderungen der Server-Konsolidierung mit Virtualisierung optimiert ist. Diese Produktklasse erfüllt typischerweise Bedingungen wie:

- Niedrige Kosten pro TB
- Einfache Handhabung
- Gute Integration in die bestehenden Architekturen
- Bereitstellung aller Sonderfunktionen wie Kompression, Deduplizierung und Thin-Provisioning

Auch wenn momentan keine Betriebskostenrechnungen für diese neue Art von Lösungen vorliegt, kann doch mit einiger Wahrscheinlichkeit angenommen werden, dass diese vergleichsweise niedrig sind (man verzeihe uns diese Art der Spekulation, aber es spricht doch einiges dafür).

Auf der Cloud-Seite unterscheiden sich



Abbildung 9: Cloud Storage

Cloud Storage in der Analyse

die Kosten je nach Anbieter erheblich. Im amerikanischen Markt ist es üblich, sowohl pro belegtem Speicher zu bezahlen als auch für Ein- und Auslagerungen. Damit werden die zukünftigen Kosten weitestgehend unkalkulierbar. (siehe Abbildung 9)

Aus unserer Sicht sollte ein Cloud-Storage-Dienst folgende Preismerkmale erfüllen, um auf Dauer kalkulierbar und wirtschaftlich zu sein:

- Preis pro bereitgestelltem oder belegten GB / TB
- Anzahl Admins und User über den Preis gestaltbar
- Trennung zwischen der Anzahl User/Admins und dem Umfang des belegten Speichers
- Keine Berechnung von Ein- und Auslagerungen, Flatrate im Sinne, dass dies im Grundpreis enthalten ist
- Optionspreise für die im technischen Bereich aufgeführten Zusatzoptionen

Fazit

Vergleicht man Cloud Storage mit unse-

rem drängenden Bedarf im Bereich Speicher, dann wird klar, dass dieser im Wesentlichen weder angesprochen noch abgedeckt wird. Man könnte auch sagen, dass die Diskussion um Cloud Storage von den wirklichen Problemen ablenkt und das Potenzial hat, die wichtigen Projekte zu stören.

Schon die Diskussion der Bewertungs-Kriterien macht deutlich, dass aus unserer Sicht am Reifegrad und an der Nutzbarkeit der angebotenen Cloud-Storage-Lösungen erhebliche Zweifel bestehen.

Value-Added-Dienste leben zur Zeit davon, dass es keine Alternativen zur Einbindung bestimmter mobiler Endgeräte gibt (speziell des Apple iPads). Trotzdem ist ihre Eignung für den Unternehmens-Einsatz nur sehr bedingt gegeben, da elementare Grundfunktionen im Bereich der Benutzerverwaltung häufig fehlen. Zudem werden unsere wesentlichen Kriterien auch hier nicht erfüllt. Hier bleibt für das iPad und die iPhones abzuwarten, wie Apples neuer iCloud-Dienst im Markt einschlägt. Naturgemäß deckt er die Windows-Welt nicht mit ab, das wird für viele Unternehmen Grund genug sein, weiterhin bei Dropbox und Co nach Lösungen zu suchen.

Im Bereich der Basis-Dienste ist die Liste der nicht erfüllten Kriterien lang. Nimmt man die aufgeführten Kriterien ernst, dann ist spätestens bei der Service-Frage und der Vertragsgestaltung bei vielen Anbietern schon das Ende erreicht.

Zusammengefasst:

- Bedarf Basis-Dienste: nein
- Kriterien-Erfüllung Basis-Dienste: nein
- Bedarf Value-Added: ja
- Kriterien-Erfüllung Value-Added: nur bedingt

Trotzdem wird es je nach Unternehmen immer wieder einzelne Einsatzgebiete für Cloud-Storage geben. Gerade im Bereich der Versorgung der internationalen Standorte multinationaler Konzerne gibt es immer wieder den Fall, dass der Weg in die Cloud die schnellste und greifbarste Lösung ist. Allerdings wird dieser Weg mit vielen Kompromissen und Nachteilen erkauft werden, die wir im Detail in der Besprechung unserer Prioritätsliste benannt haben. An dieser Stelle verweisen wir auch auf neue Produkte, die aus dem globalen Umfeld der Private Clouds stammen und eine gezielte Verteilung von Information über Standortgrenzen hinweg erlauben, EMC ATMOS ist ein Beispiel dafür.

Report

Public und Private Clouds in der Analyse

Dieser Report von Dr. Jürgen Suppan und dem Team von ComConsult Research analysiert: was leisten Public Clouds, werden sie sich durchsetzen, wie reif ist die Technologie, wo steht die Private Cloud, macht sie überhaupt Sinn, welche Vor- und Nachteile bieten diese Technologien, welche Anforderungen an Infrastrukturen entstehen, wer ist wie betroffen?

Dieser Report von Dr. Jürgen Suppan und dem Team von ComConsult Research analysiert:

- Was leisten Public Clouds?
- Werden sie sich durchsetzen?
- Wie reif ist die Technologie?
- Wo steht die Private Cloud?
- Macht sie überhaupt Sinn?
- Welche Vor- und Nachteile bieten diese Technologien?
- Welche Anforderungen an Infrastrukturen entstehen?
- Wer ist wie betroffen?



Die Ergebnisse des Reports sind sehr gemischt. Zum einen können Cloud-Dienste identifiziert werden, die sofort und unmittelbar einen Zugewinn bringen. Der Report verdeutlicht auch, welche IT-Bereiche in Zukunft besonders durch die Cloud beeinflusst werden. Es werden im Report auch kritische Bereiche identifiziert, die mindestens im Moment gemieden werden sollten. Der Report leitet daraus den Bedarf zur Entwicklung eines Unternehmens-spezifischen Cloud-Portfolios, quasi der optimalen Cloud-Dosis, ab. Nutzen Sie diesen hochaktuellen Report, um ihre optimale Cloud-Dosis zu bestimmen.

Autor: Dr. Jürgen Suppan

Preis: € 398,- zzgl. MwSt. und Versand



Bestellen Sie über unsere Web-Seite www.comconsult-research.de

Aktuelle Veranstaltungen

Sicherer Internetzugang, 07.09. - 09.09.11 in Aachen

Das Internet hat sich zu der entscheidenden Plattform für moderne Kommunikation und Geschäftsfelder entwickelt – trotz aller mit der damit verbundenen weitgehend unkontrollierten globalen Vernetzung einhergehenden Bedrohungen für IT-Infrastruktur und Daten. Der Anschluss an dieses Kommunikationsmedium muss daher so gestaltet sein, dass unkalkulierbare Risiken vermieden werden, ohne Nutzungspotenziale zu verschenken. Dieses Seminar identifiziert die wesentlichen Gefahrenbereiche und zeigt effiziente und wirtschaftliche Maßnahmen zur Umsetzung einer erfolgreichen Lösung auf. Alle wichtigen Bausteine werden detailliert erklärt und anhand praktischer Projektbeispiele und Übungen wird der Weg zu einer erfolgreichen Sicherheits-Lösung aufgezeigt. Preis: € 1.890,- zzgl. MwSt.

IPv6: Planung, Migration und Betrieb, 12.09. - 14.09.11 in Siegburg

Der Wechsel von IPv4 auf IPv6 wird für die meisten Unternehmen und Behörden in den nächsten Jahren unvermeidbar kommen. Dabei liefert IPv6 nicht nur ein neues Adress-Konzept sondern auch ein völlig verändertes Betriebs-Szenario. DHCP und auch DNS müssen neu durchdacht werden. Naturgemäß sind auch Firewall-Installationen und NAT von einer IPv6-Umstellung betroffen. Preis: € 1.890,- zzgl. MwSt.

IT-Projektmanagement Kompaktseminar, 12.09. - 14.09.11 in Aachen

Ein Projekt stellt an einen Projektleiter hohe Anforderungen. In diesem Kurs vervollständigen Sie praxisnah Ihre Kenntnisse aus der gesamten Bandbreite des Projektmanagements: Der Kurs umfasst sowohl Administratives, wie Planen und Überwachen des Projekts, als auch Softskills, wie Moderation von Projektsitzungen und Präsentation von Information. Denn die in der Regel nur „lose“ unterstellten Projektmitarbeiter müssen überzeugend auf Basis einer strukturierten Planung geführt werden. Und jede Chance, sich und sein Projekt erfolgreich zu präsentieren, ist zu nutzen! Preis: € 1.890,- zzgl. MwSt.

Lokale Netze für Einsteiger, 12.09. - 16.09.11 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,- zzgl. MwSt.

Sonderveranstaltung: UC - Cisco versus Microsoft - Wer hat die bessere Unified-Communications-Lösung?, 16.09. - 16.09.11 in Düsseldorf

Die Sonderveranstaltung UC - Cisco versus Microsoft analysiert die bestehenden UC-Lösungen von Cisco und Microsoft und stellt die spannende Frage, wer die bessere Lösung hat. Auch die erkennbaren Weiterentwicklungen der nächsten Jahre werden dabei berücksichtigt. Bewertet wird nicht nur die rein technische Funktionalität, sondern die Veranstaltung gibt auch Einschätzungen zum strategischen Einsatz sowie zur Zukunftssicherheit der Produkte ab. Preis: € 890,- zzgl. MwSt.

Netzzugangskontrolle: Technik, Planung und Betrieb, 19.09. - 21.09.11 in Köln

Dieses 3-tägige Seminar vermittelt den aktuellen Stand der Technik der Netzzugangskontrolle (Network Access Control, NAC) und zeigt die Möglichkeiten aber auch die Grenzen für den Aufbau einer professionellen NAC-Lösung auf. Schwerpunkt bildet die detaillierte Betrachtung der Standards IEEE 802.1X, EAP und RADIUS. Dabei wird mit IEEE 802.1X in der Fassung von 2010 und mit IEEE 802.1AE (MACsec) auch auf neueste Entwicklungen eingegangen. Preis: € 1.890,- zzgl. MwSt.

Trouble Shooting in vernetzten Infrastrukturen, 20.09. - 23.09.11 in Aachen

Dieses Seminar vermittelt, welche Methoden und Werkzeuge die Basis für eine erfolgreiche Fehlersuche sind. Es zeigt typische Fehler, erklärt deren Erscheinungsformen im laufenden Betrieb und trainiert ihre systematische Diagnose und die zielgerichtete Beseitigung. Dabei wird das für eine erfolgreiche Analyse erforderliche Hintergrundwissen vermittelt und mit praktischen Übungen und Fallbeispielen in einem Trainings-Netzwerk kombiniert. Die Teilnehmer werden durch dieses kombinierte Training in die Lage versetzt, das Gelernte sofort in der Praxis umzusetzen. Als Protokoll-Analysator-Software kommt Wireshark zum Einsatz. Einer Verwendung selbst mitgebrachter Analyse-Software, mit deren Bedienung der Teilnehmer vertraut ist, steht nichts im Wege. Preis: € 2.290,- zzgl. MwSt.

Datenschutz- und steuerrechtliche Aspekte von Cloud Computing, 26.09. - 27.09.11 in Stuttgart

Das Schwerpunktthema der diesjährigen CeBit war Cloud Computing. Jederzeitige Verfügbarkeit der Daten und Kosteneinsparungen lassen Cloud Computing verlockend erscheinen. Wer jedoch Daten in fremde Hände geben will, muss sich über Datenschutz und Datensicherheit erhebliche Gedanken machen, da in Deutschland der Auftraggeber von IT-Dienstleistungen unabhängig von der vertraglichen Regelung die Haftung gegenüber Dritten in Bezug auf Datenverlust, unbefugter Nutzung oder sonstiger Datenschutzverletzungen übernehmen muss. Darüber hinaus gibt es zahlreiche Rechtsvorschriften, die die Speicherung in bestimmten Ländern oder den Datenzugriff aus diesen Ländern an bestimmte Voraussetzungen knüpfen oder sogar ganz verbieten. Preis: € 1.490,- zzgl. MwSt.

TCP/IP intensiv und kompakt, 26.09. - 30.09.11 in Stuttgart

LAN-, WLAN- und WAN-Netzwerke sind heutzutage IP-Netze, und ein Verzicht auf Nutzung des IP-basierten Internet undenkbar. Auch für früher nur mit herstellerspezifischen Protokollen in Verbindung gebrachte Anwendungsgebiete wie Telefonie oder Produktionsumgebungen gibt es mittlerweile geeignete IP-basierte Lösungen. Hersteller und Dienstleister versuchen den Eindruck zu vermitteln, die Nutzung sei kinderleicht, fast schon plug and play - man trägt ein paar Adressen ein (wenn überhaupt), und es kann losgehen. Falsch! Preis: € 2.490,- zzgl. MwSt.

Zertifizierungen

ComConsult Certified Network Engineer

Lokale Netze

12.09. - 16.09.11 in Aachen
05.12. - 09.12.11 in Aachen
06.02. - 10.02.12 in Aachen
16.04. - 20.04.12 in Aachen
03.09. - 07.09.12 in Aachen
12.11. - 16.11.12 in Aachen

TCP/IP intensiv und kompakt

26.09. - 30.09.11 in Stuttgart
27.02. - 02.03.12 in Berlin
07.05. - 11.05.12 in Hamburg
17.09. - 21.09.12 in Düsseldorf

Internetworking

17.10. - 21.10.11 in Aachen
12.03. - 16.03.12 in Aachen
11.06. - 15.06.12 in Aachen
22.10. - 26.10.12 in Aachen

Paketpreis für alle drei Seminare € 6.720,-- zzgl. MwSt. (Einzelpreise: je € 2.490,--)

ComConsult Certified Trouble Shooter

Trouble Shooting in vernetzten Infrastrukturen

20.09. - 23.09.11 in Aachen
14.02. - 17.02.12 in Aachen
12.06. - 15.06.12 in Aachen
23.10. - 26.10.12 in Aachen

Trouble Shooting für Netzwerk-Anwendungen

18.10. - 21.10.11 in Aachen
20.03. - 23.03.12 in Aachen
26.06. - 30.06.12 in Aachen
04.12. - 07.12.12 in Aachen

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

ComConsult Certified Voice Engineer

Session Initiation Protocol Basis-Technologie der IP-Telefonie

14.11. - 16.11.11 in Bonn
26.03. - 28.03.12 in Stuttgart
18.06. - 20.06.12 in Bonn
29.10. - 31.10.12 in Bonn

Umfassende Absicherung von Voice over IP und Unified Communications

17.11. - 18.11.11 in Aachen
12.03. - 13.03.12 in Bonn
11.06. - 12.06.12 in Köln
01.10. - 02.10.12 in Düsseldorf

IP-Telefonie und Unified Communications erfolgreich planen und umsetzen

26.09. - 28.09.11 in Stuttgart
28.11. - 30.11.11 in Köln
27.02. - 29.02.12 in Berlin
07.05. - 09.05.12 in Hamburg
24.09. - 26.09.12 in Bonn
26.11. - 28.11.12 in Bonn

Optionales Einsteiger-Seminar:

IP-Wissen für TK-Mitarbeiter

10.10. - 11.10.11 in Berlin
13.02. - 14.02.12 in Düsseldorf
16.04. - 17.04.12 in Bonn
10.09. - 11.09.12 in Berlin

Basis-Paket: Beinhaltet die drei Basis-Seminare
Grundpreis: € 4.740,-- zzgl. MwSt. statt € 5.270,-- zzgl. MwSt.

Optionales Einsteigerseminar: Aufpreis € 1.090,-- zzgl. MwSt. statt € 1.490,-- zzgl. MwSt.

Impressum

Verlag:
ComConsult Technology Information Ltd.
ComConsult Research
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