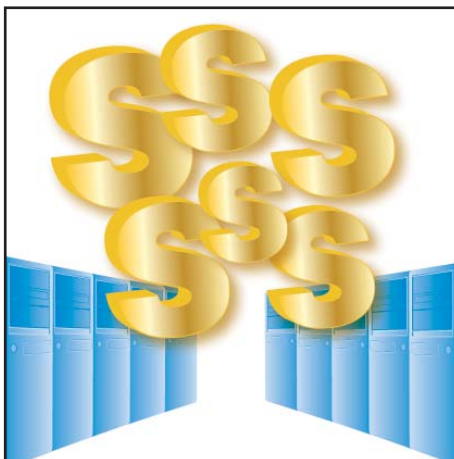


Schwerpunktthema

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S“

von Dr. Franz-Joachim Kauffels

Nach den Ereignissen Anfang August 2011 müssen wir uns auf eine in breiter Front verschärfte Wirtschaftslage einstellen, die letztlich der Preis für ein Leben über die Verhältnisse auf Pump ist. Unternehmen werden Kosten senken müssen und dabei rücken natürlich die ohnehin ziemlich ungeliebten „DV-Kosten“ in den Fokus. Gleichzeitig stehen wir im RZ vor einem erheblichen Umbruch. Die Kernfrage ist also, welche Elemente neuerer Technologie-Entwicklungen dazu geeignet sind, dauerhaft Kosten zu senken. Daraus lassen sich dann Forderungen ableiten, die man an neu zu beschaffende Komponenten unbedingt stellen muss.



Der Autor hat schon auf dem Netzwerk-Redesign Forum vor vielen Dutzend Zeugen einen Börsencrash beginnend mit Mai 2011 vorhergesagt. Während von Mai bis Mitte Juli eher ein leichtes Abschrubben überflüssiger Spekulationsbläschen stattfand, zeigt der Absturz der ersten August-Woche das gleiche Muster wie zu Beginn der Bankenkrise.

weiter auf Seite 11

Zweitthema

Glasfaserstecker, sinnvolle und weniger sinnvolle Auswahlkriterien

von Dipl.-Ing. Hartmut Kell

Glasfaserstecker stehen im Vergleich zu Kupferstecker seit den Anfängen der Verkabelungsnormen bis heute eher außerhalb des Fokus. Standardisiert wurden zunächst der ST-Steckverbinder (mehr im Sinne eines Bestands-schutzes) bzw. der SC-Steckverbinder,

beide mit kaum nennenswerten Übertragungstechnischen Unterschieden.

Ist es an der Zeit, in Anbetracht der steigenden Datenraten einen Wechsel der Steckertechnik zu vollziehen, gar eine Komplettablösung von vorhandenen In-

stallation? Der nachfolgende Artikel beschreibt die wichtigsten Anforderungen, die Unterschiede der alten Steckverbinder sowie der neuen Typen und gibt eine Empfehlung zu diesen Fragen.

weiter auf Seite 22

Sonderveranstaltung

Verkabelungssysteme für Lokale Netze, alles standardisiert, alles klar?

auf Seite 21

Aktueller Kongress

Geleit

Neue RZ-Infrastrukturen legen die Basis für virtualisierte TK-/UC-Lösungen

auf Seite 2

Standpunkt

Rechenzentrum Infrastruktur-Redesign Forum 2011

ab Seite 4

Drum prüfe, wer sich ewig bindet

auf Seite 20

Zum Geleit

Neue RZ-Infrastrukturen legen die Basis für virtualisierte TK-/UC-Lösungen

Analysiert man die System- und Infrastruktur-Anforderungen moderner Web-Architekturen und von zukünftigen TK- und Unified Communications-Lösungen inklusive Sprache und Video, dann stellt man schnell fest, dass diese Anforderungen fast identisch sind. Betrachtet man dazu die Entwicklung auf der Herstellerseite, dann liegt die Prognose nahe, dass TK- und UC-Architekturen innerhalb der nächsten fünf Jahre komplett virtualisiert werden. Die Vorteile für die Installation und den Betrieb sind zu groß, um auf diesen Schritt verzichten zu können.

Anders formuliert:

- Rechenzentrums-Infrastrukturen müssen sich anpassen, um dem aktuellen Bedarf von Web- und Realzeit-Anwendungen zu entsprechen.
- TK- und UC-System-Architekturen sind heute schon in vielen Fällen virtualisierbar, allerdings nur in sehr einfacher Weise. Die vollständige Virtualisierung unter Ausnutzung aller Architekturvorteile wird in den nächsten 2 bis 3 Jahren erfolgen.
- Beide Welten wachsen zusammen und werden in Zukunft auf gemeinsamen Infrastrukturen aufsetzen.

Wer ist von dieser Entwicklung betroffen:

- Die Planer und Betreiber von Rechenzentrums-, Web und TK-Infrastrukturen, also Netzwerken, Speichersystemen, Servern, Klimatisierung und Stromversorgung.
- Die Planer und Betreiber von TK- und UC-Lösungen. Diese müssen die sehr unterschiedlichen Ansätze der Hersteller analysieren und beobachten. Leider stellen speziell TK- und UC-Hersteller ihre Virtualisierungs-Architekturen nur sehr nebulös dar, detaillierte Analysen lassen sich nicht vermeiden. Insbesondere die Integrierbarkeit in eine gegebene RZ-Infrastruktur und der gegebenen Betriebs-Prozesse muss geprüft werden.

Neue System- und Software-Architekturen sind gefragt. Hier geht es nicht um Feigenblatt-Virtualisierung unwichtiger System-Bestandteile, hier geht es um wirklich neue Architekturen. Einige UC-Hersteller haben ihre Investitionen in Entwicklung



stark reduziert, dies ist ein deutliches Indiz dafür, dass diese Hersteller in den nächsten Jahren weiter zurückfallen werden.

Warum sind die Anforderungen von Web-Architekturen und zukünftigen UC-Lösungen so ähnlich? Warum kann heute eine RZ-Infrastruktur so aufgebaut werden, dass sie zukunftsorientiert ist und eine solide Basis für beide Bereiche bietet?

Um diese Frage zu beantworten macht es Sinn, die Kern-Eigenschaften moderner Web-Architekturen zu betrachten:

Parallelität

Web-Architekturen basieren nicht auf einem zentralen Prozess, der umso mehr Last generiert umso mehr Benutzer da sind. Die Grundidee der Web-Architekturen ist, dass sie über die Zahl paralleler virtueller Maschinen skalieren. Je mehr Benutzer eine Web-Applikation benutzen, je mehr virtuelle Maschinen werden gestartet. Die Last einer einzelnen virtuellen Maschine ist überschaubar und kann zum Beispiel einem CPU-Core zugeordnet werden.

Skalierbares Datenbank-Design

Im Rahmen eines parallelen Designs müssen Einzelprozesse vermieden werden, die zu viel Last erzeugen. Im Prinzip muss jeder Prozess parallelisiert werden können. Bei Datenbanken ist das naturgemäß bei dem schreibenden Prozess sehr schwierig und würde eine erhebliche Applikations-Logik erfordern. Ein einfacher Lösungsansatz sieht die Trennung zwischen schreibenden und lesenden Datenbanken mit einer entsprechenden

Synchronisation voraus. Lesende Datenbanken können in beliebiger Zahl parallel zueinander existieren.

64 Bit-Architekturen

64-Bit-Architekturen sind heute auch zunehmend für Low-End-Anwendungen normal. Der Vorteil liegt in dem größeren Hauptspeicher, so dass wichtige Teile der Applikation wie zentrale Datenbank-Tabellen nicht ausgelagert werden müssen. Dies vereinfacht die Architekturen und erhöht die Realzeit-Performance. Damit steigt die Fähigkeit virtueller Systeme, auch kritische Anwendungen umzusetzen, erheblich.

Prozesse sind virtuelle Maschinen

Jeder Kernprozess wird durch eine virtuelle Maschine abgebildet. Parallele Prozesse entstehen durch parallel ablaufende virtuelle Maschinen, die durch einen Broker gestartet und kontrolliert werden. Dies passt hervorragend zu der weiterhin zunehmenden Menge an Kernen auf CPU-Ebene.

Dynamische Lasten und hohe Benutzerzahlen

In dieser Architektur spielt die Zahl der Benutzer keine Rolle. Das System skaliert über die Zahl paralleler virtueller Maschinen. Die UC-Hersteller entwickeln quasi eine Software für alle Anwendungs-Größen.

Ausfallsicherheit

Neben den vielen Möglichkeiten, die Ausfallsicherheit in virtuellen Umgebungen zu steigern, ist natürlich die Erzeugung paralleler virtueller Maschinen ein zentrales Instrument der Ausfallsicherheit. Das passt durchaus zusammen. Man nutzt High Availability und Fault Tolerance, um die Prozesse abzusichern, die sich nicht parallelisieren lassen. Dazu zählt zum Beispiel die schreibende Datenbank, die Verbindungs-Information speichert (sofern Status-basierte Protokolle zum Einsatz kommen). Auch der Broker, der parallele Instanzen erzeugt und steuert, gehört dazu.

Schnittstellen und Latenz

In einer Architektur paralleler virtueller Maschinen gibt es zwei zentrale Werte, die die Leistung der Architektur bestimmen. Die Latenz in der Kommunikation zwischen virtuellen Maschinen und die Zeit zum Verlagern einer virtuellen Maschine auf ein anderes System. In der La-

Neue RZ-Infrastrukturen legen die Basis für virtualisierte TK-/UC-Lösungen

tenz gehen die Anforderungen in den Bereich von Millisekunden (kleiner 1 ms). Zur schnellen Verlagerung des RAM-Bereichs einer virtuellen Maschine muss der aktive virtuelle Speicher verlagert werden. Je langsamer dies geschieht, desto größer kann die zu verlagernde Datenmenge werden. 10 Gigabit-Ethernet kann man hier als das notwendige Minimum ansehen. Bei den Servern stellen jetzt die ersten Hersteller bereits auf 40 Gigabit um und erhöhen die internen Bandbreiten.

Bandbreitenbedarf

Man muss sich das Rechenzentrum der Zukunft als einen großen Rechner vorstellen. Das Netzwerk ist der bisherige Systembus. Ein wenig scheint hier die Infiniband-Idee durch. Gleichzeitig sinkt durch die Konzentration auf immer dichter gepackte Blade-Center und durch die Einführung von SAN-Lösungen auch für kleinere Anwendungsbereiche die Zahl der Schnittstellen. 10 Gigabit sind heute im RZ bereits normal, aber der Wechsel auf 40 und später 100 Gigabit ist im Sinne der sich abzeichnenden Skalierung in einem Zeitraum von 5 Jahren unvermeidbar.

Eine Rechenzentrums-Infrastruktur, die diese Anforderungen erfüllt, ist zur gleichen Zeit für die IT-Architekturen der Zukunft und für Sprach- und Video-Lösungen geeignet. Man muss also aus beiden Richtungen auf diese Anforderungen achten:

- Der Planer einer Sprach- und Video-Lösung muss auf die Anforderungen an Virtualisierung und die damit verbundenen Entwicklungen auf der Herstellerseite achten. Die vom TK- und UC-Hersteller angebotene Art der Virtualisierung sollte mit ein Prio-A Auswahlkriterium sein.
- Der Planer der Rechenzentrums-Infrastrukturen muss die Anforderungen sowohl der IT- und Web-Architekturen als auch der UC-Systeme erkennen und erfüllen.

Das Netzwerk wird zum Systembus

Ältere Netzwerk-Lösungen erfüllen die gegebenen Anforderungen je nach Größe der Umgebung nur zum Teil oder gar nicht. Das Problem liegt in der Verschaltung von Switch-Systemen. Bisher bevorzugen wir Baumstrukturen, in der Regel mit Spanning Tree abgesichert. Damit entstehen uneinheitlich lange Wege durch unnötig viele Switch-Systeme. Nicht nur, dass mit diesem traditionellen Design die Kosten um die 20% höher als mit einem modernen Design sind, wichtige Eigenschaften, wie der Betrieb paralleler Wege,

sind nicht umsetzbar. Der zunehmenden Parallelität auf der Server-Ebene muss zwangsläufig auch Parallelität im Netzwerk gegenüber gestellt werden.

Das neue Design unterstellt grundsätzlich nur noch 2 Hops und optimiert durch parallele Wege nicht nur die Bandbreite, sondern auch die Latenz. Kombiniert man das mit den Realzeitanforderungen einer virtualisierten TK- oder UC-Umgebung, dann erhält man mit modernen Netzwerken deutlich mehr Leistung für 20% weniger Kosten.

Trotzdem bleibt auch in diesem neuen Design das Thema Datenverlust bestehen. Ethernet kann eben von Hause aus nicht verlustfrei arbeiten. Allerdings erlauben neue Switching-Verfahren wie DCB den gezielten Schutz für besondere Anwendungen wie eben den TK- und UC-Bereich. Zwar wurde DCB ursprünglich für die Konsolidierung von Speicher-Lösungen mit FCoE geschaffen, doch sind sich die Anforderungen eines Realzeitsystems mit denen eines Speicher-Systems sehr ähnlich. In beiden Fällen erfordern moderne Technologien (zum Beispiel Einsatz von SSD-Laufwerken in einer Hierarchien von Tiers auf der Speicherseite) sowohl eine hohe Verfügbarkeit als auch einen sehr niedrigen Jitter und eine sehr geringe Latenz. Paketverluste sind für beide Anwendungsbereiche der Worst Case.

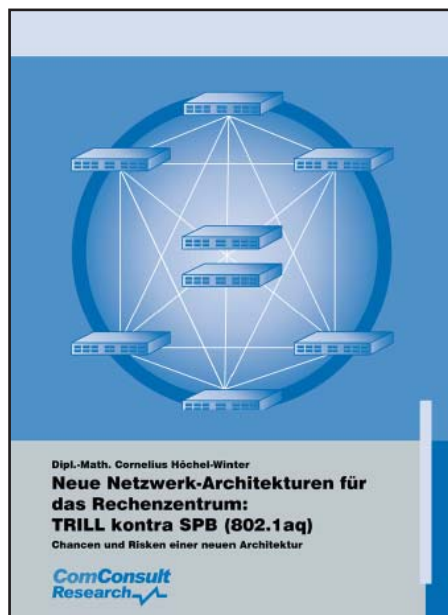


Abbildung 1: Neue Netzwerk-Architekturen für das Rechenzentrum: TRILL kontra SPB (802.1aq), neue Studie von ComConsult Research

Wer also nach einer Infrastruktur-Grundlage für sowohl neue Web-Architekturen als auch für TK-/UC-Lösungen sucht, der

kommt an einer Evaluierung von TRILL/SPB und DCB nicht vorbei. Die aktuelle Produktentwicklung auf dem Switching-Markt geht auch klar in diese Richtung. Allerdings ist bei einigen Produkten Vorsicht geboten, da relativ kleine Puffer-Bereiche zu einer Empfindlichkeit gegenüber Microbursts führen, die auch DCB nicht beheben kann. Hier sollte bei der Auswahl der Switch-Systeme darauf geachtet werden, dass genügend Puffer zur Verfügung steht.

Virtualisierungs-bewusste Infrastrukturen

Um die gegebenen Anforderungen zu erfüllen, muss es einen sauberen und gestaltbaren Übergang zwischen virtualisierten Servern und Netzwerken geben. Das ist bisher nicht der Fall. Anders formuliert: an der Realzeitfähigkeit traditioneller Schnittstellen zwischen VMs und physikalischem Netzwerk sind begründete Zweifel angebracht. Die Hersteller haben dieses Problem erkannt und die neuesten Produktvarianten zum Beispiel aus dem Hause Cisco erfüllen höchste Anforderungen an Realzeit-Fähigkeit (siehe UCS 2.0-Ankündigung). Allgemein gesprochen sind VEB und VEPA als Standards für virtuelle Umgebungen, in denen TK- und UC-Systeme umgesetzt werden, im Prinzip je nach Lastsituation unverzichtbar.

Zusammenfassend kann man sagen, dass der Bedarf für Datenapplikationen im RZ sich zunehmend mit den Anforderungen moderner TK/UC-Lösungen deckt. Wer beide Bereiche hinreichend beherrscht, wird ohne Probleme in der Lage sein, Infrastrukturen aufzubauen, die für beide Bereiche eine solide Basis legen. Bei den aktuellen TK/UC-Lösungen ist der Blick nach vorne erforderlich, da noch nicht alle Produkte den vollen Umfang einer virtualisierten Umgebung abdecken. Betrachtet man aber einen Zeitraum von 3 Jahren in die Zukunft, dann ist der Weg in eine vollständige Virtualisierung unter Nutzung aller Eigenschaften einer virtuellen Infrastruktur vorgezeichnet.

Unsere beiden Foren im November analysieren diese Entwicklungen aus beiden Richtungen. Das ComConsult RZ Infrastruktur-Redesign Forum 2011 analysiert und diskutiert die aktuellen Technologie-Entwicklungen im Rechenzentrum. Das ComConsult TK-, UC- und Video-konferenz-Forum stellt dem die aktuellen Entwicklungen aus diesem Bereich gegenüber. Die Frage der Anforderungen virtueller Implementierungen wird dabei eine besondere Rolle spielen.

Ihr
Dr. Jürgen Suppan

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.

Die weiterhin zunehmende Zentralisierung und Konsolidierung von Servern und Speicher in Kombination mit neuen Applikations-Architekturen stellt große Anforderungen an wirtschaftliche und zukunftsorientierte Infrastrukturen. Gleichzeitig nimmt die Abhängigkeit zwischen den Technologien weiter zu, so dass ein Rechenzentrum immer mehr als ein großes integriertes System betrachtet werden muss.

Das aktuelle Schlagwort im Rechenzentrum heißt „Technologie-Konvergenz“ und beschreibt diese zunehmende Integration und Verknüpfung der Technologien. Beispiele für diese Konvergenz sind:

- Aufbau von System- und Applikations-Architekturen, die in ihrer Leistung direkt vom hoch-performanten Zusammenspiel von Server, Netzwerk und Speicher abhängen
- Virtuelle Infrastrukturen mit mobilen Lasten und Verfahren für Hochverfügbarkeit und Disaster Recovery
- Aufbau Schrankinterner Kommunikationssysteme zwischen Blades und Servern sowohl für Daten als auch für Speicher mit großen Auswirkungen auf die Anbindung an externe Daten- und Speichernetze
- Hohe Abhängigkeiten zwischen virtuellen Maschinen auf physikalischen Servern und ihrer Integration in die externen Switch-Systeme
- Skalierung und Hersteller-übergreifende Integration von Speicher-Systemen mit direkter Abhängigkeit von der Kommunikations-Infrastruktur

Das ganze findet statt in einem Umfeld, indem sich die beteiligten Technologien in einem steten Wandel befinden:



- 15 Server des Jahres 2005 können heute durch einen einzelnen Server abgelöst werden, ein Ende dieser Entwicklung ist nicht absehbar
- Virtualisierung entwickelt sich immer weiter auch in Richtung komplexer Datenbank-Applikationen wie SAP
- Gigabit Netzwerke sind tot, wir stehen vor einer Welle neuer 10 Gigabit-Switches, 40 Gigabit wird in den nächsten 18 Monaten aktuell
- Neue Chip-Architekturen ermöglichen völlig neue Formen von Kommunikationssystemen
- Verfahren wie TRILL, SPB und DCB generieren ein völlig neues Verständnis vom Begriff Netzwerk, weg vom den bisherigen Baumarchitekturen
- SAN-Verfahren wie Auto-Tiering, Deduplizierung, Thin-Provisioning und Virtualisierung werden auch im Low-End-Bereich üblich. Skalierbarkeit wird immer mehr zur zentralen Herausforderung
- Die Diskussion um FCoE ist vorbei. Heute stehen FC, FCoE, iSCSI und pNFS gleichberechtigt nebeneinander und generieren hohe Anforderungen an die Infrastrukturen

Diese Entwicklung wirft eine ganze Reihe von Fragen und auch Problemen auf:

- Generell: was müssen Infrastrukturen im RZ in den nächsten 5 Jahren leisten?
- Welche Anforderungen entstehen an die Planung und den Betrieb von zunehmend abhängigen Technologien?
- Wie werden Server, Speicher und virtuelle Maschinen am besten in Netzwerke integriert?
- Welche Netzwerk-Technik wird sich durchsetzen? Welche der neuen Standards sind unbedingt erforderlich? Wie sinnvoll sind die neuen Chip-Architekturen?
- Brauchen wir Fabric-Technologien, die mehrere Switches in einem Netzwerk als ein System behandeln? Wo liegen die Vorteile?
- Wie sieht die SAN-Lösung der Zukunft aus? Auf welchem Kommunikationssystem wird die Anbindung an Server und die Kommunikation zwischen SAN-Komponenten im Rahmen von Speicher-Virtualisierung in Zukunft basieren?

Hier setzt unser hochaktuelles ComConsult RZ Infrastruktur-Forum 2011 an. Top-Experten diskutieren mit Ihnen die aktuellen Technologie-Entwicklungen und die Alternativen, die sich Ihnen bieten. Wir analysieren die wichtigsten Marktentwicklungen und stellen wichtige Produkte und Trends auf den Prüfstand. Das Forum ist in folgende Themenbereiche unterteilt:

- Anforderungen und System-Architekturen
- Physische Infrastruktur und Versorgung
- Server und Virtualisierung
- Netzwerke
- Speicher
- Sicherheit

Durch dieses Forum führen Sie Dr. Franz-Joachim Kauffels und Dr.-Ing. Behrooz Moayeri.

Versäumen Sie nicht, sich einen Platz in dieser herausragenden Veranstaltung zu sichern.

Rechenzentrum Infrastruktur-Redesign Forum 2011

Montag, den 07.11.2011

Themenblock A: Anforderungen und System-Architekturen

9:30 bis 10:30 Uhr

Keynote A: Aktuelle Anforderungen an moderne Rechenzentren: die Projektsicht

- Kapazität, Leistung, Latenz
- Struktur, Multimandantenfähigkeit
- SAN-Integration

Dr. Behrooz Moayeri, ComConsult Beratung und Planung GmbH

10:30 bis 11:30 Uhr

Keynote B: Zukünftige Anforderungen an RZs: die Technologiesicht

- Virtualisierung, Konsolidierung, Ost-West-Verkehr
- Cloud Computing und Cloud Storage
- Entwicklung der Basistechnologie: Prozessoren, Speicher, 40/100G
- Struktur: sind Data Center Fabrics die Zukunft?

Dr. Franz-Joachim Kauffels, freier Unternehmensberater

11:30 - 12:00 Uhr Kaffeepause

Themenblock B: physische Infrastruktur und Versorgung

12:00 bis 12:45 Uhr

Optische Anforderungen an ein praxisgerechtes DC Design:

- Güteklassen für LWL Stecker vor dem Hintergrund der begrenzten Dämpfungsbudgets beim 10, 40 und 100 GB/s und möglicher Rangierungen über mehrere LWL Strecken
- Konfiguration und Auskreuzungen von parallel optischen Verbindungen in Hinblick auf die Normierung und die Migration von 10GB/s bis 100GB/s
- Management von hohen Faserdichten auf begrenztem Raum

Stefan Ries, Reichle & De-Massari AG

12:45 bis 14:15 Uhr Mittagspause

14:15 bis 15:00 Uhr

Rauminfrastruktur, Versorgung und Verkabelung

- Energieversorgung, Kühlung, Schränke
- Energie-Effizienz im RZ

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

15:00 bis 15:45 Uhr

Anforderungen an elektrische Versorgungsstrukturen

- Typische Probleme
- Gesetzliche Auflagen
- Lösungsansätze

Peter Gabler, öffentlich bestellter und vereidigter Sachverständiger

15:45 bis 16:15 Uhr Kaffeepause

Themenblock C: Server und Virtualisierung

16:15 bis 17:15 Uhr

Virtualisierungstechniken: neuester Stand

- VMware vSphere 5.0
- VMware vs. Citrix vs. Microsoft
- Generelle Entwicklungstendenzen

Dipl.-Inform. Matthias Egerland, ComConsult Beratung und Planung GmbH

17:30 bis 18:00 Uhr

Streitfragen als Diskussionsbasis zur Happy Hour

- 10G ToR - Ein Irrweg der Industrie
- Data Center Fabrics - braucht man das wirklich?
- DCB - können wir auch ohne leben?
- Chipdesign - sehen in Zukunft alle Switches gleich aus?

Dipl.-Ing. Markus Nispel, Enterasys Networks Deutschland GmbH

Ab 18:00 Uhr Happy Hour

Dienstag, den 08.11.2011 - Aktuelle Entwicklungen

9:00 bis 10:00 Uhr

Planung eines RZ für das Jahr 2012 auf der grünen Wiese

- Server, Speicher, Netzwerke in Konvergenz
- Prozessautomatisierung, Automation, Standardisierung
- Bausteine des RZ im Jahr 2012
- Wie die Bausteine zusammen spielen
- Wichtige Trends
- Zentrale Infrastruktur-Elemente

Axel Simon, Thorsten Meudt, Hewlett-Packard Deutschland GmbH

10:00 bis 10:45 Uhr

Virtualisierung 2. Welle und die Auswirkungen auf Server:

- Die 2. Virtualisierungswelle: warum alles anders ist
- Virtualisierung von Tier 1 Anwendungen wie SAP
- RISC-x86 Migration: Architektur-Unterschiede und Auswirkung auf Tier 1
- Desktop-Virtualisierung
- Anforderungen an die Infrastruktur - Netz - Server - Storage
- Entwicklung bei den CPUs / Cores
- Server Auswahl, Rahmenbedingungen, etc.

Ulrich Hamm, Cisco Systems GmbH

10:45 bis 11:15 Uhr Kaffeepause

11:15 bis 12:00 Uhr

Anbindung virtualisierter Systeme

- Grundsätzliche Anbindungsalternativen
- Standards VEB, VEPA, ...
- Realzeitfähige Anbindungen (VNIlink, Direct Path)
- Aus der Sicht von VMware
- TRILL kontra SPB

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

Themenblock D: Netzwerke

12:00 bis 13:00 Uhr

Data Center Fabrics: Fluch oder Segen ?

- Motivation und Risiken

- Strukturelle Alternativen (Fat Tree, VC, ScaleOut)
- Proprietäre Virtual Chassis Techniken vs. Standards

Dr. Franz-Joachim Kauffels, freier Unternehmensberater

13:00 bis 14:30 Uhr Mittagspause

14:30 bis 15:15 Uhr

Ethernet Fabric

- Ethernet und SAN: Ausgangspunkt Vergangenheit
- Netzwerk-Design kontra SAN-Design
- Lossless LAN kontra FC
- Multipathing und Failover
- Die Idee der Fabric
- Wichtige Eigenschaften

Reinhard Lichte, Brocade Communications GmbH

15:15 bis 16:00 Uhr

Server-Anbindung im Jahr 2011

- Server-Virtualisierung und Auswirkung auf Netzwerk und Storage
 - Fabric Computing und IO Konsolidierung
 - Wo helfen die neuen Standards?
- Virtual und Physical Server Access
 - Adapter FEX, VM-FX
- Workload Mobility

Gerd Pflüger, Cisco Systems GmbH

16:00 bis 16:30 Uhr Kaffeepause

16:30 bis 17:15 Uhr

Scheitern traditionelle Netzwerke im RZ?

- Herausforderungen im RZ
- Storage Network Fibre Channel
- Konvergenz der Netzwerke, was bedeutet das?
- Was treibt Storage over Ethernet?
- Wir brauchen eine neue Architektur: QFabric
- Datacenter Evolution

Dipl.-Ing. Alexander Steger, Juniper Networks

Programmübersicht: Rechenzentrum Infrastruktur-Redesign Forum 2011

Mittwoch, den 09.11.2011

Themenblock E: Speicher

9:00 bis 10:00 Uhr

Speichertechnologie:

Aktuelle Entwicklungen, Herausforderungen und Erkenntnisse

Dr. Behrooz Moayeri, ComConsult Beratung und Planung GmbH

10:00 bis 11:00 Uhr

SAN-Integration: neue Herausforderungen und Lösungen

- Enterprise Storage Management
- Herausforderung SSD-Speicher
- Aktuelle Trends

Dr. Georgios Rimikis, Hitachi Data Systems GmbH

11:00 bis 11:30 Uhr Kaffeepause

Themenblock F: Sicherheit und Management

11:30 bis 12:30 Uhr

Anforderungen an die Überwachung im dynamischen, virtualisierten Rechenzentrum

- Monitoring virtualisierter Rechenzentren
- Einfluss der Change-Automation auf CMDB und CMS Systeme
- Service Level Management im Cloud Context
- Einbindung der Incident- und Problem-Management Prozesse
- Auswirkungen auf das Business Service Management

Ralf Horstmann, ComConsult Kommunikationstechnik GmbH

12:30 bis 14:00 Uhr Mittagspause

14:00 bis 14:45 Uhr

Sicherheitstechnologien im RZ

- Fluch und Segen von Zonenarchitekturen
- Informationssicherheit trotz maximaler Konsolidierung bei Virtualisierung
- Server-based Computing und Desktop-Virtualisierung
- Absicherung von SAN und NAS
- Data Loss Prevention im RZ

Dr. Simon Hoff, ComConsult Beratung und Planung GmbH

14:45 bis 15:30 Uhr

Trouble Shooting im RZ

- Reaktive Protokollanalyse versus Langzeitaufzeichnung: was braucht der RZ-Betreiber?
- Beitrag der Mess- und Analysetechnik zum Service Level Management
- Wie aufwändig ist die Langzeitaufzeichnung?
- Flächendeckende Verteilung von Taps: ist das notwendig?
- Was leisten Matrix-Switches?
- Produktsituation

Dr. Joachim Wetzlar, ComConsult Beratung und Planung GmbH

15:30 bis 16:00 Uhr Kaffeepause

16:00 bis 16:45 Uhr

RZ-Fernkopplung

- Aktuelle Technologien
- Herstellerkonzepte
- Weit wandernde VMs

Dipl.-Inform. Matthias Egerland, ComConsult Beratung und Planung GmbH

Donnerstag, den 10.11.2011 - Optional: Ein-Tages-Intensiv-Training/Workshop - 9:00 - 15:30 Uhr

Beginn 9:00 Uhr Die Workshops laufen parallel. Bitte wählen Sie bei Ihrer Anmeldung eines der zwei angebotenen Themen aus!

Workshop 1

Data Center Fabrics

Dr. Franz-Joachim Kauffels, freier Unternehmensberater

Workshop 2

Die neuen Standards in der Analyse, TRILL kontra SPB

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

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 € 2.390,-- zzgl. MwSt.

ohne Workshop am letzten Tag

☐ vom 07.11. - 09.11.11 in Königswinter zum Preis € 1.990,-- zzgl. MwSt.

☐ Bitte reservieren Sie mir ein Zimmer

vom _____ bis _____ 11

☐ inklusive Report "RZ Netzwerk-Infrastruktur Redesign"

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

☐ inkl. Report „Neue Netzwerk-Architekturen für das Rechenzentrum: TRILL kontra SPB (802.1aq)“

zum Sonderpreis von € 310,-- zzgl. MwSt.

Vorname _____

Nachname _____

Firma _____

Telefon/Fax _____

Straße _____

PLZ, Ort _____

eMail _____

Unterschrift _____



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

Kongress

TK-, UC- und Videokonferenzforum 2011

21.11. - 24.11.11 in Königswinter

Die ComConsult Akademie veranstaltet vom 21.11. - 24.11.11 ihren Kongress „TK-, UC- und Videokonferenzforum 2011“ in Königswinter.

Dieses hochaktuelle Forum analysiert aktuelle Trends, neue Technologien und Produkt-/Hersteller-Strategien im Bereich TK, UC und Videokonferenztechnik. In diesem Jahr stehen drei Themen im Mittelpunkt des Forums:

Wie viel UC braucht TK?

Telefonieren muss Jeder, aber wie viel Unified Communication wird wirklich benötigt und welche Alternativen der Umsetzung gibt es?

Wir analysieren:

- Wo stehen integrierte Lösungen, die TK und UC aus einem Guss liefern?
- Wie sinnvoll sind Ergänzungs-Lösungen, die mehr TK-orientierte Installationen durch eine externe UC-Lösung ergänzen?
- Welche Rolle spielen Web-basierte Dienste wie Microsoft Office 365, die ein umfangreiches UC-Portfolio auf Webbasis liefern?
- Wie sinnvoll ist die Nutzung von Web-Diensten wie Facebook, Google+ oder auch Skype als Ersatz für eine UC-Lösung?

Welche Bedeutung haben mobile Teilnehmer in Zukunft und wie werden sie integriert?

Der Anteil von Kommunikations-Teilnehmern mit mobilen Endgeräten nimmt permanent zu. iPhone und Android-Phones auf der einen und iPads und Android-Tablets auf der anderen Seite werden für viele mobile Mitarbeiter das schnurgebundene Büro-Gerät verdrängen. Sie bieten einen erheblichen Kommunikationsumfang und speziell die Tablet-Endgeräte sind als Kommunikations-Endpunkte mit integriertem Video ausgelegt.



Wir analysieren:

- Welche Rolle spielen mobile Endgeräte in Zukunft?
- Wie können sie sinnvoll in eine Gesamtlösung integriert werden?
- Welche Clients gibt es im Moment und welche Vor- und Nachteile haben sie?
- Wohin entwickelt sich die Client-Technik auf mobilen Geräten?
- Was leisten heutige Produkte in diesem Bereich?
- Wie sieht es um die Sicherheit dieser Lösungen aus?
- Was passiert mit privaten iPads und iPhones, die dienstlich genutzt werden sollen?

Welche Rolle wird die Videokonferenz in Zukunft haben?

Videokonferenztechnik ist seit Jahren auf dem Übergang von einer Technologie für Wenige hin zu einer Massentechnologie für Alle. Mit dem zurzeit erfolgenden Übergang zur HD-Konferenztechnik für fast alle Endpunkte entsteht ein gemeinsamer Basisstandard. Trotzdem skalieren viele Produkte immer noch nicht wirtschaftlich.

Wir analysieren für Sie:

- Wo steht Videokonferenztechnik und was sind die dominierenden Zukunftstrends?
- Videokonferenztechnik für Alle, was bedeutet das?
- Integration TK, UC und Videokonferenztechnik, wie sieht das aus?
- Welche Endgeräte werden in Zukunft dominieren?
- Welche technischen Infrastrukturen sind gefordert?
- Wo stehen die führenden Anbieter?

Das ComConsult TK-, UC- und Videokonferenzforum 2011 bietet Top-aktuelle Information und Analysen mit ausgewählten Experten. Eine ausgewogene Mischung aus Analysen, Hintergrundwissen und Projekterfahrungen in Kombination mit Produktbewertungen und Diskussionen liefert das ideale Umfeld für alle Planer, Betreiber und Verantwortliche solcher Lösungen. Zögern Sie nicht, sich rechtzeitig einen Platz in dieser herausragenden Veranstaltung zu sichern.

Durch das Forum führen Sie Petra Borowka-Gatzweiler und Cornelius Höchel-Winter.

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

Dipl.-Math. Cornelius Höchel-Winter ist Leiter des Testlabors der ComConsult Technologie Information GmbH. In dem Labor werden regelmäßig Messungen und Evaluierungstests neuester Hard- und Softwareprodukte durchgeführt und ausgewertet. Herr Höchel-Winter besitzt langjährige Erfahrung in der Konzeptionierung, im Aufbau und Betrieb von Windows- und Unixnetzen.

TK-, UC- und Videokonferenzforum 2011

Frühbucherrabatt bis 30.09.11

**TK-, UC- und
Videokonferenzforum 2011
21.11. - 24.11.11 in Königswinter**

Wir bieten Ihnen exklusiv eine Vorbuchungsphase für das ComConsult TK-, UC- und Videokonferenzforum 2011 bis zum 30.09.2011 für eine rabattierte Teilnahmegebühr an.

ComConsult TK-, UC- und Videokonferenzforum 2011
zum Preis bei Buchung bis 30.09.11 von € 2.190,-- zzgl. MwSt.
statt regulär € 2.390,-- zzgl. MwSt. (4-Tages-Veranstaltung) bzw. € 1.790,-- zzgl. MwSt.
statt regulär € 1.990,-- zzgl. MwSt. (3-Tages-Veranstaltung)

Die Buchung innerhalb der Frühbucherphase ist verbindlich, kann aber jederzeit auf einen anderen Mitarbeiter Ihres Unternehmens übertragen werden.

Fax-Antwort an ComConsult 02408/955-399

Anmeldung TK-, UC- und Videokonferenzforum 2011

Ich buche den Kongress
TK-, UC- und
Videokonferenzforum 2011

mit Workshop am letzten Tag

☐ vom 21.11. - 24.11.11 in Königswinter
zum Preis € 2.190,--* zzgl. MwSt.

ohne Workshop am letzten Tag

☐ vom 21.11. - 23.11.11 in Königswinter
zum Preis € 1.790,--* zzgl. MwSt.

*gültig bis zum 30.09.11

☐ Bitte reservieren Sie mir ein Zimmer

vom _____ bis _____ 11



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

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift

Neuer Report

Neue Netzwerk-Architekturen für das Rechenzentrum: TRILL kontra SPB (802.1aq)

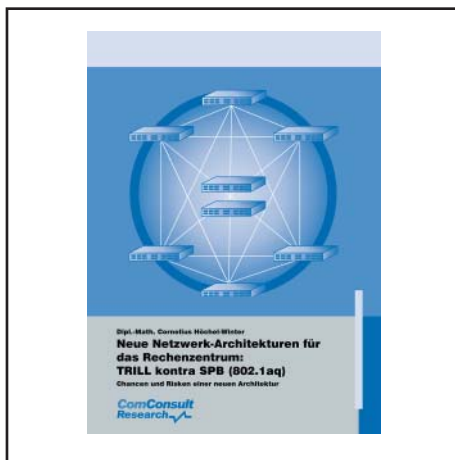
Chancen und Risiken einer neuen Architektur

Die ComConsult Research veröffentlicht Mitte September 2011 ihren neuen Report "Neue Netzwerk-Architekturen für das Rechenzentrum: TRILL kontra SPB (802.1aq) - Chancen und Risiken einer neuen Architektur".

Die Argumente für die neuen Netzwerk-Architekturen im Rechenzentrum sprechen eine deutliche Sprache:

- deutlich verbesserte Performance
- höhere Verfügbarkeit
- 20% weniger Kosten

Ausgangspunkt für die neuen Verfahren ist die Erkenntnis, dass die traditionellen Ethernet-Switching-Architekturen in einem modernen Rechenzentrum mit Layer-2-Strukturen nur sehr schlechte Ergebnisse bringen. Schlecht ausgenutzte Ressourcen, schwer skalierbare Kapazitäten, hohe Umschaltzeiten im Fehlerfall und eine komplexe Architektur sind nicht akzeptierbar. Seit Jahren arbeitet die Normung bereits an neuen Verfahren, die diese Probleme lösen sollen. Diese Verfahren revolutionieren das Netzwerk-Design im Rechenzentrum. Leider



hat es in der Entwicklung eine Parallelentwicklung zwischen IEEE und der IETF gegeben, so dass wir jetzt in der Situation sind, zwei konkurrierende Verfahren zu haben:

- TRILL seitens der IETF
- 802.1aq von IEEE

Die meisten Hersteller unterstützen bisher TRILL, obwohl es gute Gründe in Richtung

IEEE 802.1aq gibt. Der Anwender sollte deshalb die Eigenschaften beider Verfahren kennen und unter Abwägung der Vor- und Nachteile für seine Umgebung eine Entscheidung treffen.

Hier setzt dieser wichtige Report an, der die Arbeitsweise beider Verfahren erklärt und die Unterschiede aufzeigt. Er legt die Basis für eine fundierte Entscheidung für das für Ihre Umgebung optimal geeignete Produkt.

Sie erfahren in diesem Report:

- warum diese neuen Verfahren benötigt werden
- warum sie mehr Leistung bringen
- warum sie mehr Verfügbarkeit bringen
- wieso Sie 20% einsparen können
- wie sich TRILL und IEEE 802.1aq unterscheiden
- welche Vor- und Nachteile beide Verfahren haben

Dieser Report ist von essentieller Bedeutung für jeden Planer, Betreiber und Entscheider von Netzwerken im Rechenzentrum.

Fax-Antwort an ComConsult 02408/955-399

Bestellung

Neue Netzwerk-Architekturen für das Rechenzentrum: TRILL kontra SPB

Ich bestelle den Report

☐ **Neue Netzwerk-Architekturen für das
Rechenzentrum: TRILL kontra SPB
(802.1aq)**

zum Subskriptionspreis von nur € 310,-*
zzgl. MwSt. und Versand

*gültig bis 15.09.11



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

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift

Aktuelle Neuerscheinungen bei ComConsult-Study.tv

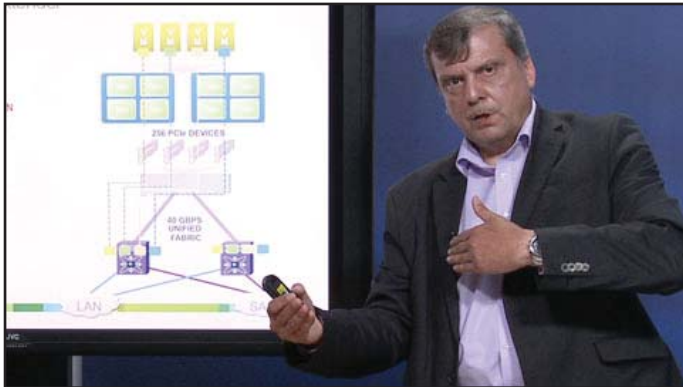
Themenbereich: Desktop, Server, Speicher

Cisco Unified Computing System

Referent: **Ulrich Hamm**

Zeit: 00:26:12

Preis: Kostenlos



Ulrich Hamm erklärt die Architektur des UCS und die Neuerungen der Version 2. Erfahren Sie, worin sich UCS von anderen Server-Lösungen unterscheidet.

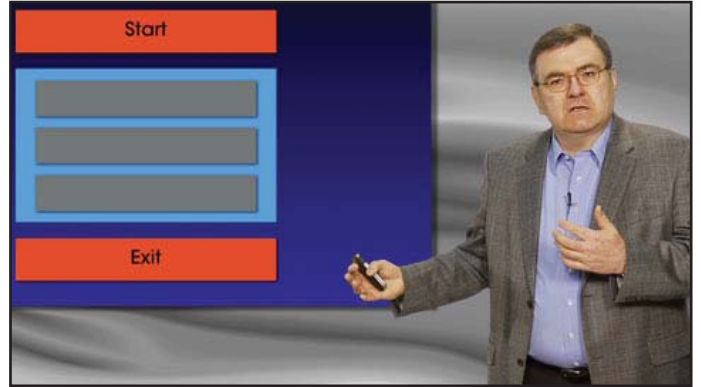
Themenbereich: Betrieb und Architekturen

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.

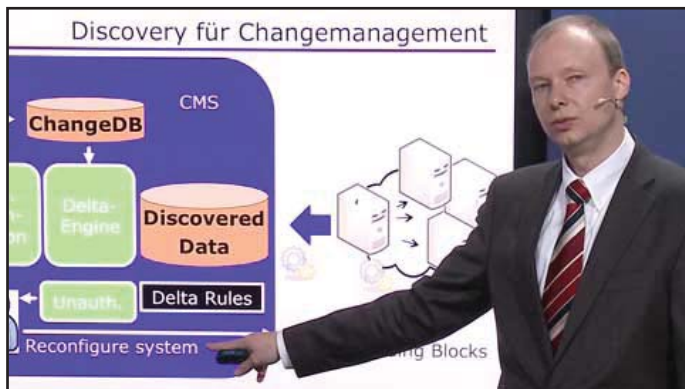
Themenbereich: Betrieb und Architekturen

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.

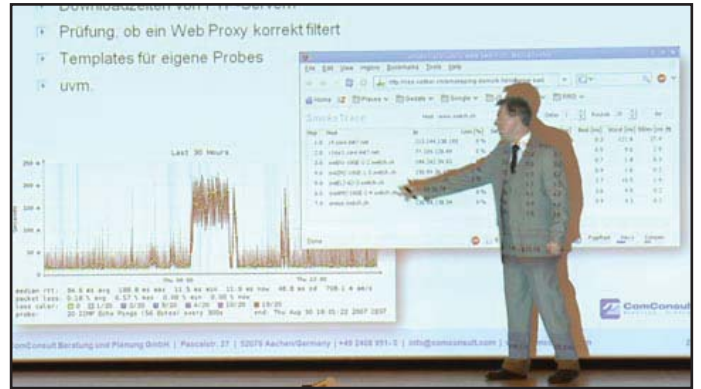
Themenbereich: Netzwerke

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.

Schwerpunktthema

Anforderungen an RZ-Kompo- nenten vor dem Hintergrund einer proble- matischen Wirtschafts- lage: die Goldenen sechs „S“

Fortsetzung von Seite 1



Dr. Franz-Joachim Kauffels ist einer der erfahrensten und bekanntesten Referenten der gesamten Netzwerkszene (über 20 Fachbücher und unzählige Artikel) und bekannt für lebendige und mitreißende Seminare.

Am 6.8.2011 hat dann auch noch Standard & Poors das Rating der USA von AAA auf AA+ herabgesetzt und damit die Tür für einen weiteren Absturz der börsennotierten Unternehmenswerte weit aufgestoßen. Verbunden mit ein wenig Panik könnte sich der DJII schnell unter 8000 Punkte verkrümmeln. Wegen der längst nicht gelösten EU-Schuldenkrise zieht das aber weitere Kreise.

Selbst wenn man es insgesamt optimistischer sieht, kann es ja nichts schaden, sich auf das Schlimmste einzustellen.

Besonders in einer solchen Situation gibt es einige Irrtümer, die sich schnell verfestigen, aber letztlich genau zu dem führen, was man grade nicht möchte: zusätzliche Kosten und Risiken.

In einem solch kurzen Artikel kann ich nur die meiner persönlichen Ansicht nach größten Irrtümer benennen. Es gibt noch viel mehr. Aber jeder einzelne führt zu grundsätzlichen Erkenntnissen, die bei der Planung der Entwicklung eines RZs von fundamentaler Bedeutung sind. Ich nenne sie die goldenen sechs „S“.

Grober Irrtum 1: Personalabbau und Cloud Computing sparen Kosten

In einer Krise fallen Entscheidungsträgern auf höheren Ebenen sofort ein bis zwei Dinge ein, die völlig kontraproduktiv sind: Personal-Abbau und Cloud-Computing.

Meine Erfahrung zeigt, dass die deutschen Unternehmen und Organisationen im letzten Jahrzehnt die Qualität der im RZ eingesetzten Mitarbeiter erheblich gesteigert haben. Außerdem wurde die Anzahl der Mitarbeiter längst nicht in dem Maße erhöht, wie die Aufgaben gewachsen sind. Ein weiterer Kahlschlag ist ohne jede Diskussion verantwortungslos. Es bleibt ja dann auch die Frage, wie man entlassene Mitarbeiter ersetzen soll, wenn es wieder aufwärts geht.

Top-Manager, die sofort mit Entlassungen winken, sind die gleichen, die in die Cloud-Falle tappen. In verschiedenen Publikationen haben wir Funktionen und Möglichkeiten des Cloud Computings kritisch beleuchtet /1/, /2/. Im Prinzip verbinden Top-Manager aufgrund geschickter Reklame der interessierten Anbieter die Hoffnung einer Margenschöpfung: der Anbieter von Cloud-Leistung baut große RZs mit günstigen Randbedingungen. Dadurch kann er zunächst anscheinend wirklich Leistungsangebote erstellen, die für ihn geringere Kosten erzeugen als für das durchschnittliche RZ eines Unternehmens. Diese Leistungen möchte der Anbieter mit einem Aufschlag verkaufen. Das ist legitim. Der Top-Manager hofft nun, diese Leistungen zu geringeren Kosten als die ihm durch ein eigenes RZ entstehenden einkaufen zu können. Das ist ein gefährlicher Irrtum, denn:

- es können in erheblichem Maße gesetzliche Bedingungen verletzt werden

- Datenschutz und Datensicherheit können neuen Gefährdungen ausgesetzt sein
- das Unternehmen muss einem Provider, den es gar nicht kennt, völlig vertrauen
- vertragliche Absicherungen können Schall und Rauch sein, z.B. wenn der Provider einen absichtlichen Konkurs herbeiführt
- der Cloud-Provider kann nach einem gelungenen Start die Preise anheben
- der Cloud-Provider kann die Abhängigkeit des Unternehmens von ihm nutzen, um überproportionale Preiserhöhungen durchzusetzen, wenn er sieht, dass das Unternehmen sein eigenes RZ abgebaut hat und ein Providerwechsel für das Unternehmen teuer würde
- und selbst ohne alle diese Risiken ist die Margenschöpfung so schwach, dass sie einem erheblichen Wechselkursrisiko unterliegt und je nach Umfang dagegen gehedged werden muss, was wiederum Kosten nach sich zieht.

Philosophisch betrachtet gehorcht „Cloud-Computing“ dem gleichen Modell wie eine Studenten-Kommune oder eine LPG. Eine Anzahl Personen tut sich zusammen und beschafft gemeinschaftlichen Wohnraum oder Maschinen. Dadurch sinken die Kosten für den Einzelnen gegenüber einer Individual-Lösung. Das gibt aber auch gleich das Stichwort

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S

für den wichtigsten Problemkreis: Individualität. Eine Kommune kann nur solange funktionieren, wie alle Teilnehmer vergleichbare Interessenlagen haben. Möchte jemand in der LPG auf einmal statt Weizen Weintrauben anbauen, hilft ihm die Gemeinschaft nicht mehr. Und genau dieses Problem tritt auch bei Unternehmen oder Organisationen auf. Es gibt keine zwei Unternehmen, die wirklich gleiche Aufgaben haben. Also benötigen sie auch individuelle Lösungen. Natürlich können sich wie bisher auch Unternehmen mit wirklich identischen Geschäftsfeldern zusammentun und z.B. gemeinsam ein RZ betreiben, wie z.B. Sparkassen. Aber der generalisierte Cloud-Ansatz wird mit Sicherheit nicht zufriedenstellend funktionieren.

Grober Irrtum 2: Skalierbarkeit ist Nebensache

Eigentlich sprechen wir ja ziemlich häufig über „Skalierbarkeit“ oder über das „Skalierungsproblem“. Die Umsetzung der dabei gewonnenen Erkenntnisse ist in der Praxis jedoch sehr unterschiedlich, besonders wenn man auf Speichersysteme blickt.

Es gibt nämlich durchaus die Tendenz, Speicher für einen Mehrjahreszeitraum auf einen Schlag zu beschaffen. Sie resultiert aus der in der Vergangenheit gemachten Erfahrung, dass man einen einmal eingesetzten Plattentyp in Extremfällen schon nach Monaten nicht mehr nachkaufen kann.

Betrachtet man nur einen monolithischen Speicher oder ein unflexibles Array ohne weitere nachgelagerte Intelligenz, kann diese Perspektive durchaus berechtigt sein.

Alle wichtigen Hersteller bieten aber heute hochflexible, modulare und in weitem Bereich skalierbare intelligente Speichersysteme an. Beispiele wären Symmetrix von EMC, Lefthand von HP, XIV von IBM oder FAS von NetApp, um nur einige zu nennen.

Ein Kunde, der einmal begonnen hat, mit einem solchen System zu arbeiten, erwartet mit Recht zweierlei:

- Möglichkeit der Erweiterung der Grundkapazität auch nach Jahren ohne Änderung bei Nutzung oder Betrieb
- Nahtlose Integration neuer technologischer Möglichkeiten (z.B. SSD) bei Verfügbarkeit und Bedarf
- Problemloser gemeinschaftlicher Betrieb auch unterschiedlicher technolo-

gischer Generationen von Systemen des Herstellers (aktuell wären die Änderungen bei der neuen XIV-Generation zu nennen)

Alle genannten und alle mit diesen vergleichbaren modularen Speichersysteme unterscheiden sich in ihrer spezifischen Arbeitsweise. Allen gemein ist aber die Einführung einer intelligenten Zwischenebene, die man auch als „Speicher-Virtualisierung“ bezeichnen kann. Die wesentliche Funktion der Speicher-Virtualisierung ist die Abstraktion der von den Anwendungs- und Systemprogrammen durchgeführten Speicherzugriffen von ihrer tatsächlichen Implementierung durch eine Speichersystemhierarchie (z.B. in den Stufen SSD, schnelle Platten, langsame billige Platten, Bänder). Darüber hinaus kann es noch weitere Zusatzfunktionen geben, z.B. im weiten Bereich der Ausfallsicherung.

Konzeptionell zieht die Technologie der Speichersysteme damit sozusagen auf das nach, was allgemein hinsichtlich der Server-Virtualisierung als vorteilhaft angesehen wird.

Wirklich spannend ist aber, dass verschiedene Hersteller (in unterschiedlichem Maße) die grundsätzliche Funktionalität der Speicher-Virtualisierung auch schon bei kleinen Einstiegs-Konfigurationen anbieten.

Die in diesem Rahmen erzielte Skalierbarkeit birgt die Möglichkeit zu erheblichen wirtschaftlichen Vorteilen, wenn man von der „Einmalbeschaffung“ Abstand nimmt. Die Preisentwicklung bei Servern und Netzwerk-Komponenten kann grob durch Moore's Law abgeschätzt werden, solange die Chips die wesentlichen Bestandteile darstellen. Moore's Law bezieht sich

aber nur auf die auf einer Fläche unterzubringenden Transistorfunktionen. Daher könnten wir es hinsichtlich der Speicher theoretisch nur für SSDs anwenden, die aber nur einen kleinen Teil der Speicherkapazität repräsentieren.

Nach einigem Suchen findet man aber ein anderes genau auf Plattenspeicher zugeschnittenes Gesetz: Komorowski's Law. Matthew Komorowski hat die Preise der führenden Festplattenhersteller über Jahrzehnte ausgewertet und nicht nur einzeln aufgeschrieben /3/, sondern auch mathematisch zusammengefasst /4/.

In der Abbildung 1 sieht man, wie sich die Kosten pro Gigabyte zwischen 1980 und 2009 entwickelt haben. Bitte beachten Sie die logarithmische Skala für die Kosten!

Mathematisch ausgedrückt gilt:

$$\text{Cost} = 10 \exp. (-0,2502(\text{year}-1980) + 6,304)$$

Mit dieser Formel können Sie schon mal die nächsten Angebote überprüfen. Vereinfacht gesagt bedeutet Komorowskis's Law die **Verdopplung der Speicherkapazität pro Kosteneinheit alle 14 Monate**.

Wie kann man diese Erkenntnis jetzt umsetzen?

Ohne in irgendeiner Weise über die Technik zu argumentieren, kann man sagen, dass in einer Verteilung einer in Speicher geplanten Investition ein erhebliches Sparpotential liegt.

Wir stellen uns jetzt einmal vor, ein Unternehmen habe seinen Speicherbedarf für die nächsten fünf Jahre extrapoliert. Das Ergebnis wäre ein Speichersystem, das aktuell unter Ausnutzung aller Rabatt-

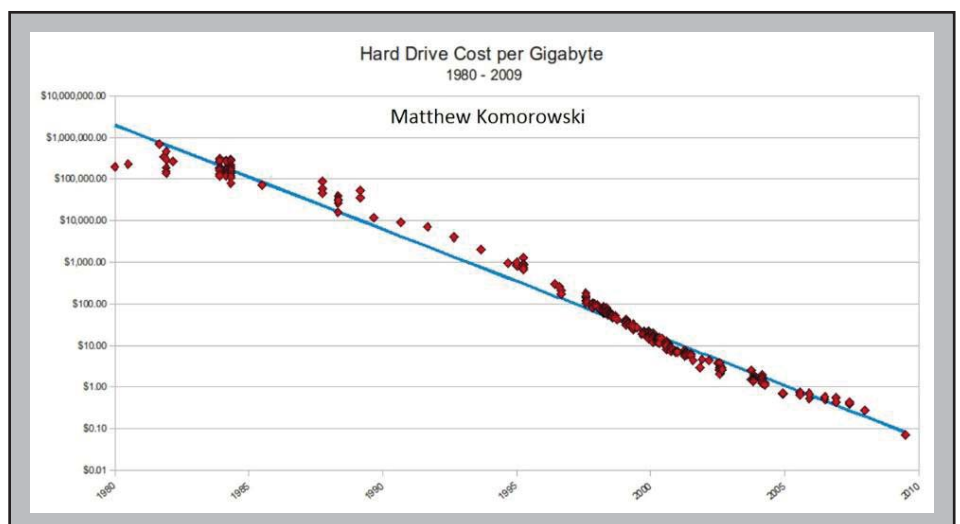


Abbildung 1: Kosten pro Gigabyte

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S

te 1 Million Euro kostet. Es ist für dieses Beispiel völlig unerheblich, ob es sich dabei um einen einzigen großen Speichertopf oder um ein modernes, strukturiertes mehrstufiges System handelt.

Schafft ein Unternehmen das Speichersystem jetzt auf einen Schlag an, ist die Million weg.

In Abbildung 2 zeigen wir, was passiert, wenn das Unternehmen nicht den ganzen Speicher auf einmal kauft, sondern fünf Jahre lang jeweils ein Fünftel der beabsichtigten Kapazität. Es geht mir hier um das Grundprinzip, es wären auch andere Teilungen möglich, die aber jetzt nicht weiterführen.

Jedes Jahr wird die Speichermenge, die man kaufen möchte, nach Komorowski's Law billiger, wobei wir natürlich hier keine Wechselkurschwankungen berücksichtigen können.

Interessant ist jetzt, dass nach 5 Jahren der Speicher so wie geplant zusammengekauft worden ist, aber statt 1.000.000 Euro oder Dollar zusammen nur 628.000 gekostet hat. Wir sparen also einen Ferrari einzig und alleine durch eine Verteilung der Anschaffungsmenge.

Dabei wurde davon ausgegangen, dass das Geld in Form eines Haufens vorhanden ist und einfach der Reihe nach ausgegeben wird. Das ist natürlich extrem vereinfacht und unrealistisch.

Deshalb stellen wir in der Abbildung 3 drei weitere Fälle dazu:

- das Geld ist vorhanden, der zwischenzeitlich stehengebliebene Betrag wird mit 3% verzinst.
- das Geld ist nicht vorhanden, sondern muss genau dann, wenn es benötigt wird, aufgenommen werden. Der Kredit kostet jeweils 6% oder 8% und läuft jeweils über 5 Jahre.

Selbst im zweiten Fall bietet die verteilte Beschaffung erhebliche Vorteile und mildert selbst hohe Kreditkosten. Sogar ein kompletter „Speicher auf Pump“ ist bei verteilter Beschaffung noch über 20% billiger als ein direkt aus Eigenmitteln finanziertes Komplettsystem!

Ein solches Ergebnis erzielt man nur bei einem erheblichen Preisverfall der zu beschaffenden Komponenten, also wie z.B. einem Verfall um den Faktor 10 (oder 90%) in 5 Jahren.

Jetzt sind natürlich folgende Sätze unverzichtbar: Die Modellrechnung enthält in

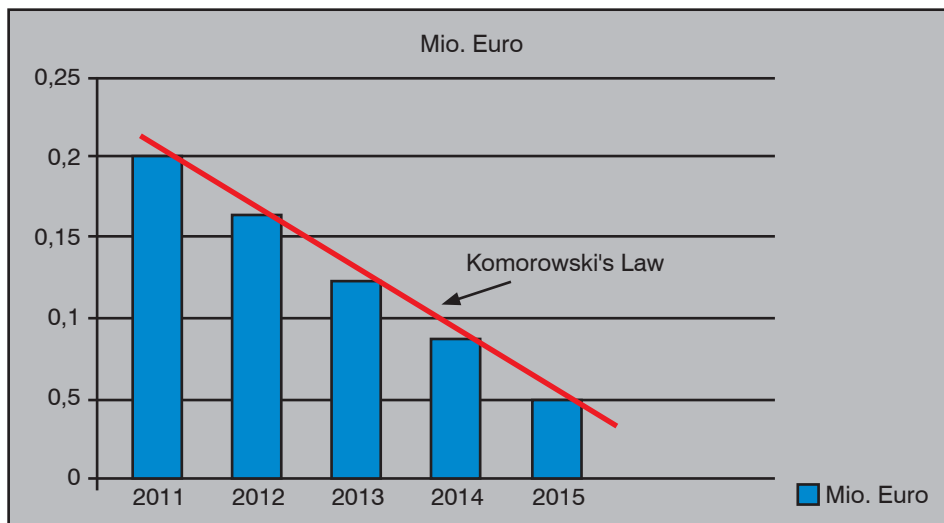


Abbildung 2: Verteilung der Anschaffungsmenge über 5 Jahre

die Zukunft extrapolierte Aussagen. Mögliche Abweichungen durch eine erhebliche Inflation der in den Berechnungen benutzten Währungen wurden nicht berücksichtigt. Autor und Herausgeber haften nicht für den tatsächlichen Eintritt der unter Modell-Annahmen getroffenen Ersparnispotentiale.

Aber: was kann bei einer verteilten Beschaffung wirklich passieren? Eine Explosion der Platten-Preise hat es in den Jahren von Komorowski's Law nicht gegeben. Warum sollte sie angesichts des extremen, sich bei schlechteren Bedingungen tendenziell noch verschärfenden Preiswettbewerbs der Hersteller also gerade in den nächsten fünf Jahren kommen?

Schließlich gewinnt ein Unternehmen ja auch Flexibilität. Es kann notwendig sein, von den geplanten 100 % Endkapazität in einem Jahr 25 statt 20% und dafür im nächsten nur 15% zu installieren.

Fassen wir zusammen: die einmalige Beschaffung von Speicher für einen Mehrjahreszeitraum ist finanztechnischer Unfug.

An dieser Stelle können wir aber auch sofort die erste wichtige Eigenschaft, die an anzuschaffende Komponenten und Konzepte zu stellen ist, nochmals nennen: Skalierbarkeit.

Grober Irrtum 3: DAS ist billiger als SAN

Über die Notwendigkeit von SANs in virtuellen Umgebungen, die verschiedenen Technologien (iSCSI, FCoE, FC, IB) und die Konvergenzproblematik haben wir alle in den letzten Jahren viel geschrieben. Ca. 33 qm Ausführungen dazu kommen alleine von diesem Autor in /5/.

Dennoch hält sich das Gerücht, dass ein SAN teuer sei und man besser die Platten

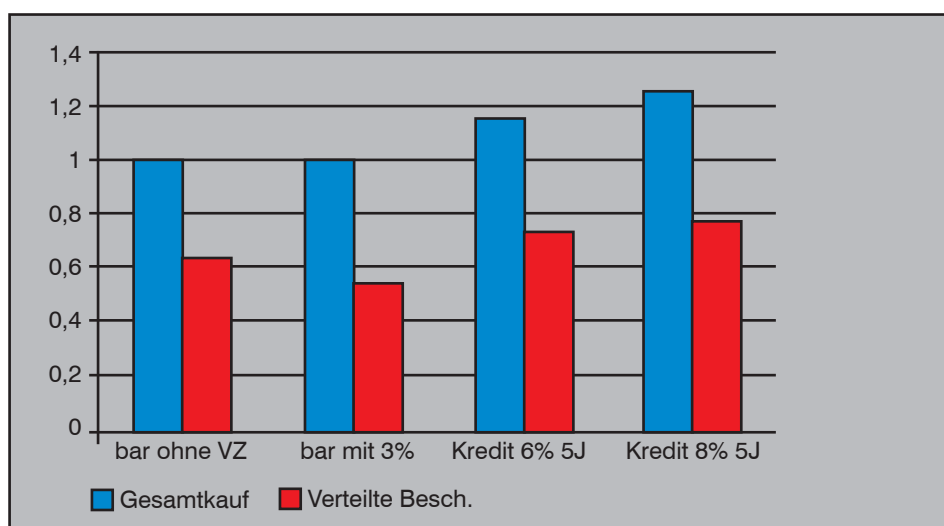


Abbildung 3: Effekt der Verteilung der Anschaffungskosten unter verschiedenen Randbedingungen

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S

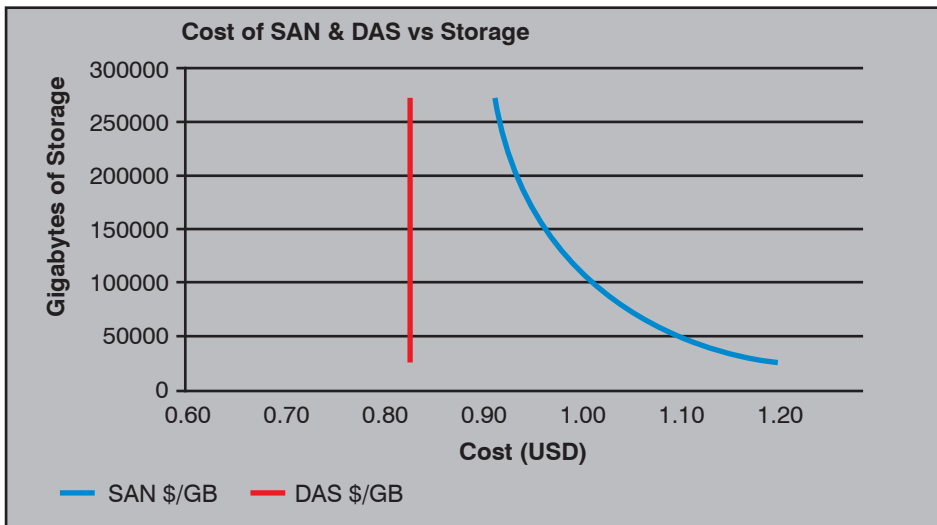


Abbildung 4: Cost of SAN & DAS vs Storage

direkt an die Server hängen sollte, hartnäckig.

Eine Kostenuntersuchung /6/ mit Speicher-Hardware von Dell brachte eine andere Wahrheit ans Licht. Wenn man bestimmte systemspezifische Einzelheiten ausblendet, ist das Ergebnis übertragbar.

In der Abbildung 4 sieht es natürlich danach aus, dass das SAN teurer ist. Die Graphik ist etwas ungewöhnlich gezeichnet, weil die y-Achse die Gigabytes und die x-Achse den durchschnittlichen Preis pro Gigabyte wiedergibt.

Es entspricht aber genau unseren intuitiven Erwartungen, dass (bei einem Kauf an einem Tag) der Preis für das Gigabyte bei DAS gleich bleibt, egal wie viel man kauft. Bei der Untersuchung der Kosten wurde Speicher immer in Blöcken von 146 oder 300 GB gekauft, weil zu diesem Zeitpunkt die 500 und 1000 GB-Platten langsamer waren.

In einem SAN können aber auch die langsameren Platten hoher Kapazität gut eingesetzt werden, weil für die üblichen Anwendungsfälle die im SAN mögliche Parallelverarbeitung greift. Der Nachteil eines SANs liegt natürlich in zunächst hohen Basis-Kosten durch die notwendige Beschaffung der Kommunikationskomponenten, im konkreten Fall wurde FC benutzt.

Die Abbildung 4 zeigt aber auch (logischerweise), dass die relativen Kosten für ein SAN mit der Gesamtkapazität fallen, weil eben die Basis-Kosten auf immer mehr Laufwerke umgeschlagen werden können.

Dennoch sieht es so aus, als sei das SAN teurer. Dabei wird aber eine wesentliche

Fähigkeit des SANs übersehen, nämlich die, den Speicherplatz wesentlich effektiver zu nutzen. Eine Umfrage von Theln-

foPro.net bei den Fortune 1000 Unternehmen hat ergeben, dass DAS nur zwischen 30 und 40% ausgenutzt wird. Ein SAN wird aber aufgrund seiner Flexibilität wesentlich höher ausgenutzt, z.B. zu 80%. (siehe Abbildung 5)

Bezogen auf eine 45 TB-Konfiguration mit 1000GB SATA-Drives ergibt sich unter Berücksichtigung der durchschnittlichen Auslastung Abbildung 6.

Diese Argumentation ist statisch, wir vergleichen dabei DAS und SAN nur hinsichtlich der Kosten pro GB und das auch nur für eine Klasse Platten, nämlich SATA. Jedes Unternehmen kann aber für seine Anforderungen und Auslastungsgrade einen Schnittpunkt finden, bei dem ein SAN letztlich günstiger als DAS ist.

Vorteile der SANs, die ein DAS nie erreichen wird, sind darüber hinaus:

- zentralisiertes Management,

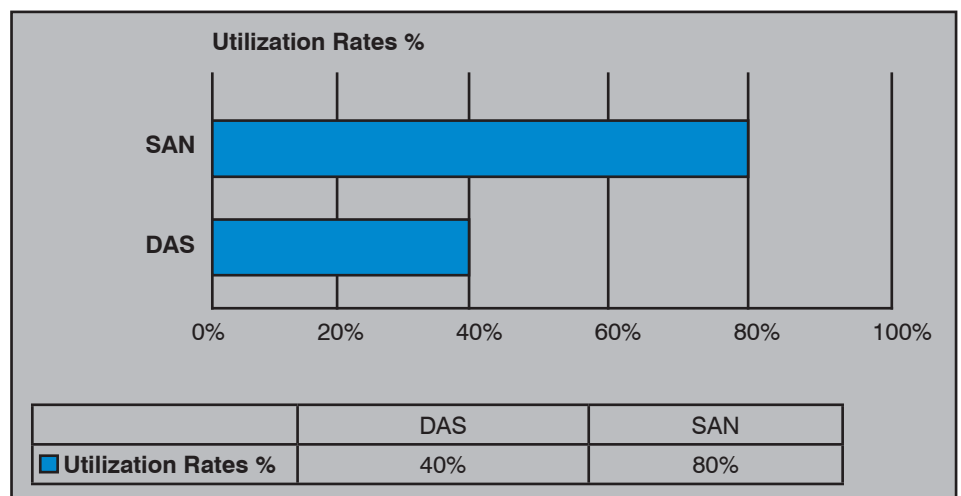


Abbildung 5: Utilization Rates in %

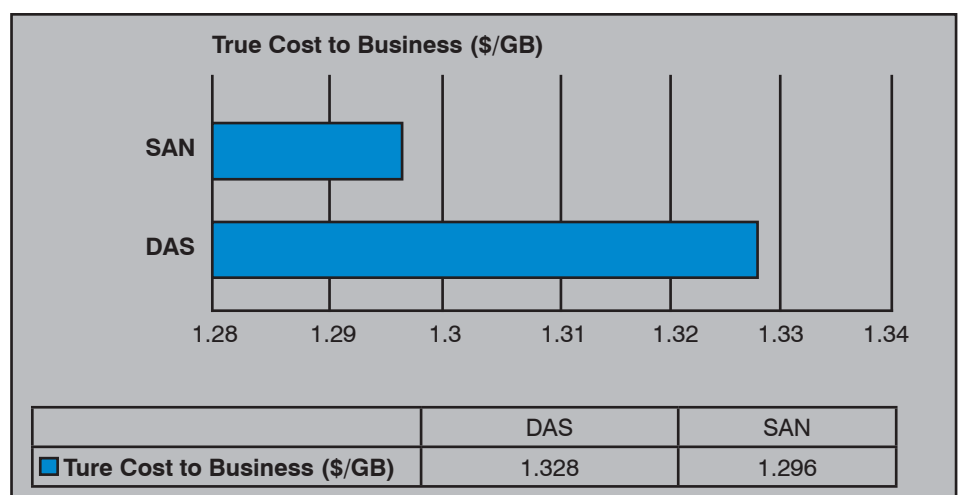


Abbildung 6: True Cost to Business in \$/GB

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S

- hohe Sicherheit,
- sinnfälliger Verwaltung der Speicher-Ressourcen,
- einheitliche Darstellung und Implementierung spezieller Storage Services, wie periodische Backups,
- Unterstützung des Betriebs effektiver Benutzungsniveaus der Speicher-Ressourcen.

Ein SAN ist natürlich von sich aus per Definitionem skalierbar und erfüllt damit die im letzten Abschnitt genannte Forderung perfekt.

Grober Irrtum 4: ScaleUp ist eine sinnvolle Strategie

Langsam kommen wir von den Speichern zu den Netzen. Hier gibt es ein verbindendes Argumentationsglied, nämlich die Strukturdiskussion. Niemand käme heute noch ernsthaft auf den Gedanken, zwei riesige Plattenstapel als die Speicherstrategie eines Unternehmens zu definieren. Aber genau das sollen Unternehmen nach Ansicht der überwiegenden Mehrzahl der Hersteller bei ihrem RZ-Netz machen: sie erklären immer noch zwei (oder 4 – 8) maximal fette Core-Switches zum Kern eines Netzes.

Scale-Up bedeutet, die Aufgaben immer größeren monolithischen Maschinen zu übertragen. Eigentlich wissen wir schon, dass das nicht der richtige Weg ist. Neben Skalierungsproblemen kommen Kostenexplosion und Management-Probleme auf uns zu.

Scale-Out bedeutet die horizontale Skalierung durch Clustering multipler kleinerer Maschinen.

In einem Scale-Out Design werden viele Elemente, die nach Industriestandards arbeiten, mit einer Clustering-Software zu einem größeren Gebilde zusammengebunden, die für eine vereinheitlichte Systemsicht sorgt. Da kleinere Komponenten nach Industriestandards viel schneller hergestellt werden können als große monolithische Switches, dauert es bei einem solchen Design längst nicht so lange, bis neue Industriestandards auch in Form von Produkten bei den Kunden ankommen.

Bei Speichern ist ScaleOut längst Gang und Gäbe, siehe Abbildung 7.

Ein Produktbeispiel wäre das Data ONTAP 8 Cluster-Mode-System von NetApp. Die Abbildung 8 zeigt eine Übersicht.

Durch Architektur und Struktur ergeben sich erhebliche Vorteile im Betrieb, in Abbildung 9 nur am Beispiel des Global Namespace exemplarisch dargestellt.

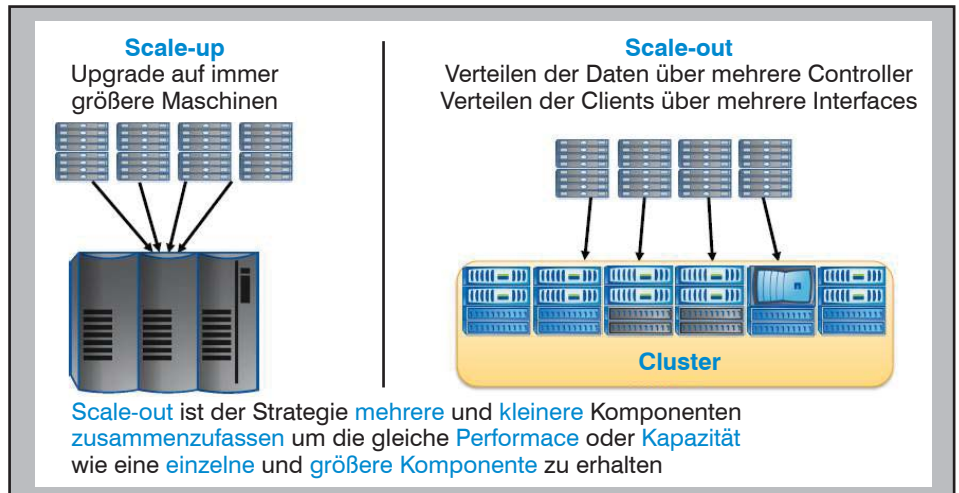


Abbildung 7: ScaleOut

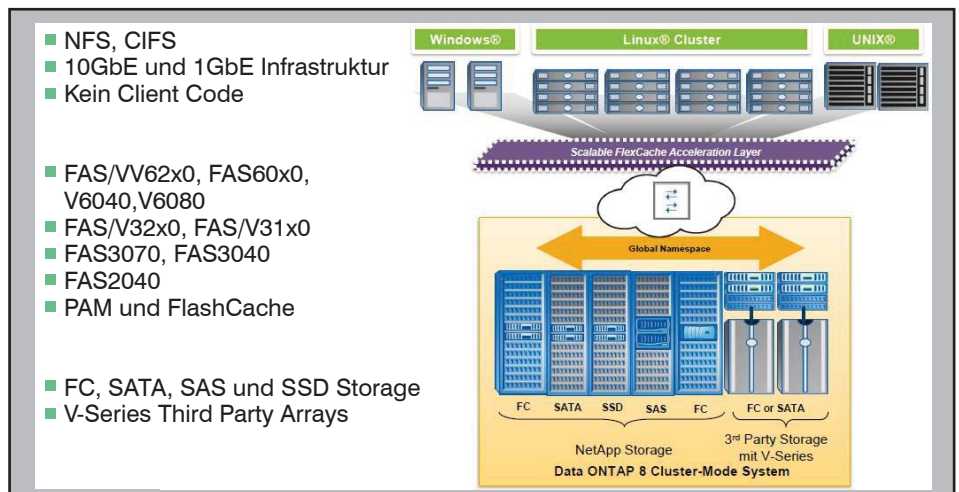


Abbildung 8: Data ONTAP 8 Cluster-Mode Systemarchitektur

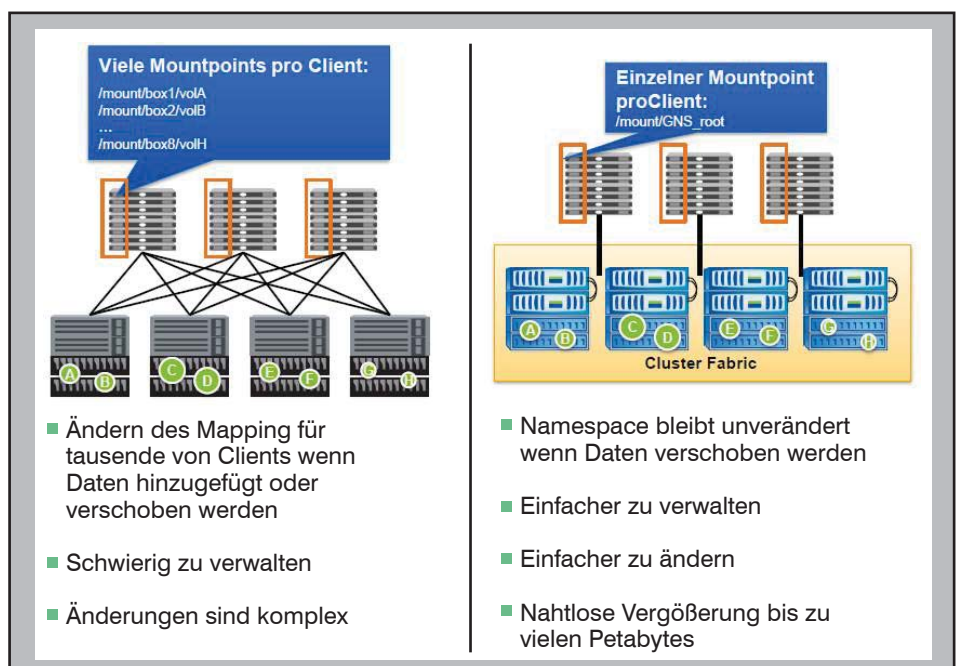


Abbildung 9: Vorteile des Global Namespace

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S

Der Global Namespace vereinfacht dann natürlich auch das Multi-Tiering. (siehe Abbildung 10)

Alle Hersteller, die in dieser Liga spielen, haben ähnliche oder vergleichbare Funktionen.

Spannend ist aber auch, dass die Summe kleinerer Systeme immer erheblich preisgünstiger ist als ein großes System. Nehmen Sie einen beliebigen Netzwerk-Hersteller. Die Kosten für einen ordentlichen 10 GbE TOR oder Bladeswitch liegen zwischen 500 und 900 US\$ pro Port. An einem Core-Switch kostet der 10 GbE-Port zwischen 2500 und 5000 US\$ pro Port. Selbst wenn wir alle möglichen Marktverwerfungen, asiatische Reimporte und Rabatte hinzunehmen, ändert sich an diesem Verhältnis nichts.

Denken Sie jetzt einmal an ein Fat Tree Design, wie es von praktisch allen Netzwerk-Herstellern sogar noch als neue Strategie gepusht wird und die Kosten, die dadurch entstehen, dass es auf der Spine-Ebene große Core-Switches geben muss und man zu den Kosten der Ports an den Leaf-Switches, die die eigentliche Versorgung der Server vornehmen noch die Kosten für die Ports der gesamten Infrastruktur hinzurechnen muss. Wie viele Ports man im Rahmen der Infrastruktur braucht, hängt vom Konzentrationsgrad auf den Leitungen ab. Da dieser individuell verschieden ist, können wir hier nicht weiter rechnen, aber mit dieser Anleitung können Sie das für Ihr Netz ja selbst.

Die Abbildung 11 ist sozusagen eine „Gegendarstellung“ zu den üblichen Herstellerfolien und zeigt ein ScaleOut-Netz für die leistungsfähigsten HP-Server. Am interessantesten ist das, was Sie nicht sehen: kein Core Switch!

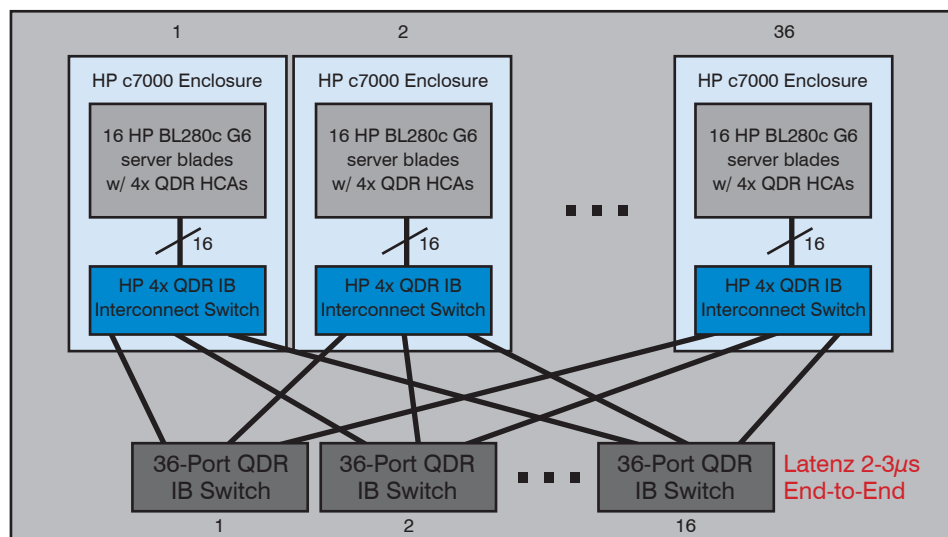


Abbildung 11: HP ScaleOut Cluster

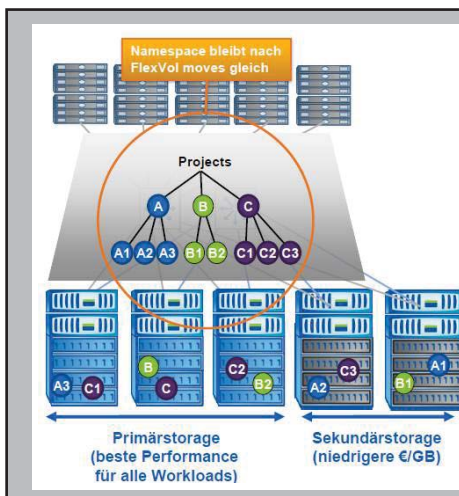


Abbildung 10: Tiered Storage

Die Diskussion der Folgekosten durch Platzbedarf der Core-Switches, Stromversorgung und Kühlung unterlassen wir jetzt, aber da ist auch Potential drin.

Es gibt viele Beispiele für Scale-Out-Lösungen im RZ wie Cluster von Web-Servern, geclusterte Anwendungen und Datenbanken sowie geclusterten File- oder Blockspeichern. Sie alle zeigen, dass Scale-Out ein erfolgreiches konstruktives Konzept ist. Alle diese Beispiele verbessern Leistung und Kapazität mit weniger Ressourcen. Der Erfolg einer Cluster-Lösung hängt von ihrer Software-Architektur ab. Dabei stehen folgende Fragen im Vordergrund:

- Wie nahe kommt das Verhalten des Clusters einem monolithischen System?
- Wie steht es um die Linearität der Skalierbarkeit und deren Grenzen?

■ Tiered Storage

- Verwalten von mehreren Tiers im gleichen Namespace
- Verschieben von Daten ohne Downtime auf andere Speicher-/Controllerklassen
- Beispiele
 - Referenzdaten
 - DR Mirror Ziel
 - Skalierbare Archive

- Wie verhält sich das geclusterte System bei Fehlern?

Es gibt sogar die These, dass die bisherigen Scale-Up-Lösungen bald an konstruktive Grenzen stoßen werden. Die Core-Switches müssen einerseits immer mehr Aufgaben übernehmen, andererseits wachsen Datenrate und Konzentrationsgrad laufend, das heißt, die Mehraufgaben müssen immer schneller erledigt werden. Irgendwann wird auch noch so parallelisiertes CMOS dabei an Grenzen stoßen. Also muss man auf einen anderen VLSI-Herstellungsprozess zurückgreifen, wie etwas GaAs. Dadurch werden die Core-Switches wahrscheinlich nochmal deutlich schneller, aber bestimmt auch nochmal wesentlich teurer, ganz abgesehen vom Energieverbrauch.

Dieser Scale-Up Ansatz ist eigentlich nur eine Krankheit von Ethernet. Bei anderen Netztypen wie InfiniBand und in gewissem Umfang Fibre Channel hat man schon früh erkannt, dass es so nicht geht und das Scale-Out Konzept angewendet. Eine typische InfiniBand-Umgebung besteht üblicherweise aus einer sehr leichtgewichtigen Switching-Struktur, die folgende Eigenschaften aufweist:

- Fähigkeit der Arbeit in fast beliebigen Mesh-Topologien
- Unterstützung paralleler, multipler Wege
- Fabric- und I/O-Partitionierung
- Zentrale Discovery
- Policy Management

In den letzten Jahren hat IEEE aber an Verbesserungen von Ethernet gearbeitet, die dem Ethernet in Form des Converged Enhanced Ethernet InfiniBand-ähnliche Eigenschaften verleihen, wie wir das ja schon oft diskutiert haben.

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S

Grober Irrtum 5: Standardisierung ist halb so wichtig

Die konventionelle vielstufige Strukturierung eines RZ-Netzes ist ein Auslaufmodell. Eine nahe liegende Möglichkeit ist die Schaffung so genannte Fat Trees oder die Nutzung eines Virtual Chassis. Beide Konzepte haben aber den schwerwiegenden Nachteil, dass sie auf teuren fetten Core-Switches beruhen. Genau das ist es aber, was uns die meisten Hersteller heute als neue Strategie verkaufen möchten.

Alle Data Center Fabrics nach diesen Strategien haben mehr oder minder die gleiche zweistufige Struktur: eine Stufe besteht aus möglichst preiswerten Switches für den Anschluss der Endgeräte. Die zweite Stufe besteht dann aus Core-Switches. Beide Stufen zusammen bilden einen latenzarmen Fat Tree, der auch manchmal als Clos-Netz bezeichnet wird. Machen Sie sich mal den Spaß und gehen auf die verschiedenen Websites der Hersteller. Dann sehen Sie sogar weitestgehend identische Fundamentaldiagramme. Eine Gruppe der Hersteller verwendet in der Leaf-Stufe sogar die gleichen Trident + - Chips von Broadcom, eine andere Gruppe benutzt selbstentwickelte Chips. Das ist auch nicht weiter zu kritisieren, denn der Trident+ Chipset ist sehr gut. Die Hersteller haben nur das (selbstgemachte) Riesenproblem, dass sie sich auf dieser Ebene nicht weiter von den anderen Herstellern unterscheiden. Cisco hat jetzt auch schon einen Low Latency Switch (Nexus 3000), wo sie sich die vergebliche Mühe gemacht haben, den Trident extra zu kastrieren, damit man es nicht so schnell sieht.

Also müssen sie sich hinsichtlich der Differenzierung an einer anderen Stelle austoben und da bleibt nur die Spine-Ebene mit den Core Switches. Das ist aber auch wieder dumm für die Hersteller, denn die Aufgabe der Spine-Ebene ist so klar definiert, dass es hier an und für sich auch nur wenige Möglichkeiten gibt.

Aber, aktuell haben sie dann doch etwas gefunden: das Virtual Chassis-Konzept. Eine Spine-Ebene wird ja nicht mit ein oder zwei Core-Switches auskommen wie auf den Bildern, sondern eine größere Anzahl von Core Switches sozusagen zu einem sinnvollen Ganzen zusammenschalten.

Also hat jetzt jeder Hersteller ein eigenes Virtual Chassis Konzept auf dieser Ebene. Darin haben sie Übung, denn das gibt es bei verschiedenen Herstellern schon lange. Mit diesem Konzept wollen sie die Kunden an sich binden und ignorieren da-

bei vollkommen, dass man die notwendigen Funktionen auch auf standardisierter Basis mit PLSB, TRILL oder M-LAG (Multi-Chassis Link Aggregation) machen könnte. Sie gehen sogar so weit, dass sie auf Nachfrage entweder völligen Quatsch erzählen („Fabric Path ist die Trill-Realisierung von Cisco“ O-Ton Cisco Referent auf dem Netzwerk-Redesign-Forum 2011), die Realisierung von Standards diffus in die Zukunft verlegen (in Art des US-Notenbankchefs Bernanke: „Growth should pick up soon“, Rede vom 22.6.2011) oder einfach auf mehrfaches Nachsetzen ungerne zugeben, dass es zwar irgendwann die Realisierung von Standards geben könnte, der Kunde dann aber mit deutlichen Leistungseinbußen gegenüber dem proprietären Produkt zu leben hätte.

Persönlich finde ich das eine Unverschämtheit, denn verschiedene der genannten Hersteller liefern für den Provider-Bereich durch und durch standardisierte und auch mit Höchstleistung kooperationsfähige Systeme. Etwas anderes würden ihnen die Provider nämlich vor die Füße werfen. Aber mit den Corporate-Kunden kann man es ja machen (oder wenigstens versuchen).

An dieser Stelle möchte ich den Leser daran erinnern, dass es in der Vergangenheit schon zwei fulminante Abflüge gab, mit denen niemand gerechnet hätte: (das alte) 3Com hatte auf einmal keine Lust mehr, Switches zu bauen und hat sich freiwillig verabschiedet und Nortel Networks hat die größte Pleite in der Geschichte der Kommunikations-Ausrüster hingelegt. In beiden Fällen waren die Kunden die Dummen.

Der Autor beobachtet den finanziellen

Hintergrund von Netzwerk-Herstellern schon seit mindestens 20 Jahren. Es geht dabei primär um die Frage, ob ein Kunde bei einer Bindung an einen Hersteller damit rechnen muss, dass dieser im Nebel verschwindet. Die aktuelle Situation der Hersteller ist alles andere als rosig. Aufgrund des Konkurrenzdrucks werden die Margen immer kleiner und sie haben einfach erhebliche Probleme, sich darauf einzustellen. Durch die Bank haben alle börsennotierten Hersteller, die nur Netzwerk-Komponenten herstellen, auch schon in den ersten Monaten des Jahres 2011 völlig gegen den allgemeinen Aufwärtstrend (siehe Google, Apple oder IBM) erheblich an finanzieller Substanz verloren und zwar in einer Größenordnung von ca. 30%. Zu nicht gelisteten Herstellern kann ich leider auch nichts sagen. Die Aktien verschiedener Hersteller werden mit einem sehr spekulativen KGV (30 – 40, realistisch wäre 7 - 11) gehandelt, sie müssten Jahrzehnte zweistellig wachsen, um das zu rechtfertigen.

Abgesehen davon, dass es schon die ersten Entlassungen gab, sind die möglichen Auswirkungen einer erneuten Finanzkrise nicht wirklich abzuschätzen. Besonders kleinere Hersteller könnten unter die Räder kommen. Sie würden dann wahrscheinlich wegen bestehender Patente aufgekauft. Was das aber dann für den Kunden bedeutet, ist sehr fraglich.

Ein weiteres Problem für alle Hersteller kann man namentlich benennen: es heißt Huawei. Bei den Provider-Systemen konnte sich der chinesische Konzern weltweit bereits auf Platz 2 vorarbeiten, die anderen Netze sind sicher auch bald „dran“. Aktuell greift man aber erst mal Apple mit Smartphones zu einem Drittel des iPho-

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: € 2.390,- zzgl. MwSt. bzw. € 1.990,- zzgl. MwSt. ohne Workshop



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

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S

ne-Preises an. Huawei ist im Besitz seiner Mitarbeiter, die Aktien unterliegen aber strengen Auflagen.

Huawei konnte im Provider-Bereich nur deshalb so wirkungsvoll punkten, weil alle Produkte streng nach Standards aufgebaut sind. Es gibt keine proprietären Sonderlocken. Man geht lieber über den Preis als Differenzierungsmerkmal.

Was bedeutet das für die unternehmens-eigenen RZ-Netze? Ganz klar die Notwendigkeit einer strengen Ausrichtung an Standards, so wie es die Provider schon längst vormachen. Alles andere kann richtig teuer werden!

Das führt auch auf die Frage zurück, ob man das RZ-Netz nicht ganz anders organisieren könnte und den in anderen Bereichen der DV erfolgreichen Ansatz des Scale-Out auch auf Netze übertragen kann. HP (mit Voltaire), und IBM machen das schon für ihre High End Blade-Server. Alle Netzkomponenten sind dabei streng nach Standards aufgebaut, lediglich die notwendige Steuerungssoftware nicht, weil es dafür (noch) keine Standards gibt.

Grober Irrtum 6: in einem Switch ist viel drin

Einem Netzwerk-Hersteller mit einer schwächlichen Quartalsbilanz ist vor wenigen Tagen in seiner Not nichts anderes eingefallen, als einen 24-Port 1-Giga-

bit Ethernet Switch als besonders preiswerten Switch ab 998,- US\$ anzubieten. Das ist schon ein extremer Standpunkt, denn auf entsprechenden Portalen findet man solche kleinen Switches auch schon für unter 100,- Euro. Vielleicht ist ja in dem Switch dieses Herstellers viel drin, wir wissen es nicht.

Tatsache ist, dass aufgrund der neueren Entwicklungen bei den Switch-ASICs in einem Switch, der im RZ als Access-Switch positioniert wird, tatsächlich nicht mehr viel „drin“ ist.

Möchte man abschätzen, was wirklich mit den neuen Produkten der Switch-Hersteller verbunden ist und sich der Antwort auf die Frage, wie es denn in den nächsten Jahren weitergeht, annähern, ist ein Blick in die Chips unerlässlich.

In früheren Publikationen hatte ich ja schon exemplarisch zwei Chips vorgestellt, die die Welt der Datennetze bereits erheblich verändert haben: den Low Latency Switch von Fulcrum und den Converged Controller von Broadcom.

Der Low Latency Switch ASIC von Fulcrum hat dazu geführt, dass niemand mehr wagt, neue Netzkomponenten mit hoher Latenz anzubieten. Wie im Artikel /7/ von Herrn Dr. Moayeri anschaulich dargestellt, wird die neue Komponentengeneration für die allermeisten Anwendungsfälle in modernen Unternehmen und Organisationen völlig ausrei-

chend sein. Lediglich in Einzelfällen wird man auf noch reaktionsschnellere Schaltungen zurückgreifen müssen. In diesen Fällen kommt es aber dann auch kaum auf die Kosten an.

Der Broadcom BCM75840S NetXtreme II® Converged Controller vereint sozusagen alle Funktionen, über die wir in den letzten Jahren bei RZ-Netzen gesprochen haben. Es ist ein „10 Gbps Quad Port iSCSI, FCoE, TOE und RDMA PCI-SIG SR-IOV x8 PCI-Express® 3.0 ready Controller“. Dieser Controller der sechsten Generation (von 10 GbE Controllern) ist design für hochvolumigen 10 GbE LOM-Verkehr (LAN on Motherboard, was bedeutet, dass man ihn nicht nur in einem Switch, sondern z.B. auch direkt auf einer Blade-Server Karte einbauen kann) und er beherrscht eben alle genannten Protokolle. Nebenbei hat er dann auch noch viele Diskussionen, die wir in der Vergangenheit hatten, durch die Wahrheit des Silikons geschlichtet. (siehe Abbildung 12)

Das ist nur ein einziges Beispiel. Besuchen Sie doch einfach mal die Homepage des Herstellers www.broadcom.com und Sie werden Controller für wirklich jeden Bedarf finden, also auch solche mit integrierter Sicherheitstechnologie, besonders sparsame Controller, die man auch extrem dicht packen kann und solche, die die Signale besonders auffrischen („Eye-Opener“-Technologie) und damit längere Distanzen auf z.B. WAN/Campus SMF-Verbindungen erlauben.

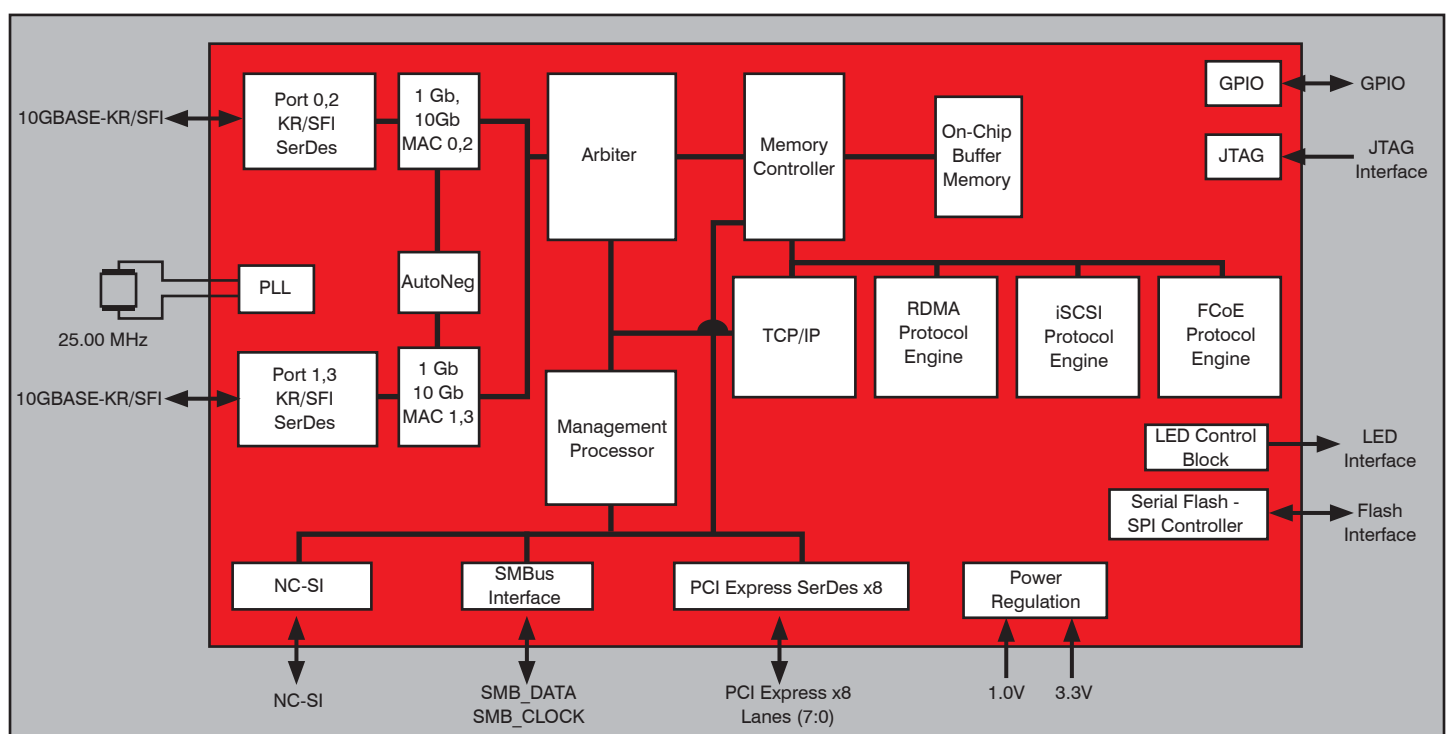


Abbildung 12: Broadcom BCM75840S Blockdiagramm

Quelle: Broadcom

Anforderungen an RZ-Komponenten vor dem Hintergrund einer problematischen Wirtschaftslage: die Goldenen sechs „S“

Nach Moore's Law können wir damit rechnen, dass sich die Anzahl der Transistorfunktionen, die auf einem zur Verfügung stehenden Platz realisiert werden können, ca. alle 18 Monate verdoppelt. Wozu wird man das aber nutzen? Natürlich könnte Fulcrum in diesem Zeitraum einen 48-Port Chip entwickeln und Broadcom die Funktionenvielfalt auf 8 Ports erweitern. Das wird auch geschehen, denn zum Zeitpunkt der Manuskripterstellung sind die Fulcrum Chips schon ein Jahr auf dem Markt und die Broadcom Chips haben auch schon mindestens ein Jahr auf dem Buckel, wenn sie jetzt ganz offiziell für jedermann zu beziehen sind.

Aber, das wäre alles zu harmlos, denn manchmal macht die Evolution auch einen Sprung. Bei integrierten Schaltkreisen ist das genau dann der Fall, wenn es gelingt, die nächste Stufe der „Miniaturisierung“ in der Fertigung zu erreichen. Bei den Prozessoren war es der Übergang zur 35 bzw. 40 nm Technologie, die die Implementierung der Nehalem-Architektur, wie wir sie heute mit den Intel® XEON®-Prozessoren kennen, erst möglich gemacht hat.

Bislang waren Switch-ASICs demgegenüber eher grob gestrickt. Das hat sich aber jetzt geändert.

Der Trident von Broadcom ist der erste kommerziell verfügbare Switch-ASIC in 40nm-Technologie. Es ist ein 64-Port 10 GbE-Switch-ASIC, den man aber wegen der Vorgaben des „skalierbaren Ethernet“ auch als 16-Port 40 GbE-Switch oder gemischt z.B. als 48-Port 10 GbE + 4 Port 40 GbE-Switch benutzen kann, was z.B. Juniper macht. Natürlich beherrscht der Chip die gesamte gewohnte Menagerie von VLANs und IPv6 über DCB bis hin zu Sicherheitsfunktionen.

Geht es noch weiter? Natürlich, der 64-Port Chip ist nur ein weiterer Schritt auf dem Weg, das ganze RZ-Netz in einen oder einige wenige Chips zu packen. Denken Sie einmal an die Scale-Out Architektur. Die NetXtreme®-Controllern besitzen schon die grundsätzliche Fähigkeit, auf sehr einfache Weise mittels der KR-Schnittstellen untereinander verbunden zu werden. Von Herstellern wie Voltaire (jetzt: Mellanox) gibt es schon länger Scale-Out-Software, die es ermöglicht, eine praktisch beliebig große Gruppe von Switches (in dem Fall noch vollständige Geräte) so zusammenzuschalten, dass sie wie ein großer Switch agieren.

Letztlich ist es durchaus denkbar, eine derartige Software mit allen ihren Leistungsmerkmalen wie z.B. dynamische Zuordnung von QoS-orientierten Daten-

flüssen zwischen VMs komplett auf einen Prozessor zu packen, der entweder ein eigenständiger Controller ist oder einfach mit auf dem Switch-ASIC sitzt. Ferner gibt es keinen Grund, derartige Funktionen nicht auf mehreren Prozessoren so zu etablieren, dass sie beliebig skaliert.

Wie sieht unter diesen Bedingungen ein RZ-Netz mit vollem Funktionsumfang und sagen wir 1000 10 GbE-Anschlüssen aus? Nach aktuellem Stand braucht man dazu nur noch 2 X 16 Trident-Chips! Selbst wenn man sie platzverschwendend verbaut und zusätzliche Prozessoren und SSD-Speicher einbaut, reicht ein einziges Rack dazu vollständig aus. Wie man sofort sieht, ist der Schwarze Peter jetzt in Richtung der Anschlusstechnik verschoben. Ohne Verwendung von MPO und Konzentration wird man hier nicht weiterkommen. Es gibt ja heute schon Transceiver, die über ein 12-Pol MPO 10X10 GbE bzw. 1X100 GbE übertragen können. Aber selbst von diesen würde man noch 100 benötigen. Auf der anderen Seite haben wir jetzt schon Server-Blades mit 2 X 40 GbE I/O-Leistung.

Ich bin fest davon überzeugt, dass das RZ-Netz, wie wir es heute kennen, binnen der nächsten fünf Jahre als eigenständige Komponente weitestgehend verschwinden wird und seine Funktionen auf integrierte Switch/Controller-ASICs auf den Server-Blades verteilt werden.

Eigentlich wollte ich ja nur darstellen, dass in Switches nicht mehr viel drin ist. Man benötigt die 1U-Bauform eigentlich nur noch, um die altmodischen Stecker verwenden zu können.

Und binnen der nächsten 2 bis 3 Jahre sollte ein voll ausgestatteter 24 X 10 GbE-Switch deshalb auch unter 100,- Euro zu haben sein.

Also, geben Sie nicht zu viel Geld für überteuerte Switches aus.

Konsequenzen für die Unternehmen: die Goldenen sechs „S“

In diesem Artikel konnten nur einige Bereiche angesprochen werden, die einen erheblichen Einfluss auf die Wirtschaftlichkeit eines RZs haben. Es gibt sicher noch andere wie Platzbedarf oder Energie-Effizienz.

Der aufmerksame Leser wird die Server vermisst haben. Hier machen Unternehmen eigentlich die wenigsten Fehler. Sie kaufen mehr oder minder genau die Server, die sie benötigen und orientieren sich dabei sehr stark an den Vorgaben aus der Anwendung.

Bei Speichern und Netzen sieht das anders aus. Hier besteht wirklich die Neigung, mit dem Holzhammer auf Spatzen zu werfen. Die Verwendung von Kanonen würde ja noch eine grobe Zielorientierung erfordern.

Deshalb werde ich nicht müde, immer wieder darauf hinzuweisen, dass die bisher meist getrennten Fachspezialisten (Server, Speicher, Netz) unbedingt zusammenarbeiten müssen. Nur zusammen können sie eine präzise Vorstellung davon entwickeln, was wirklich benötigt wird.

Fassen wir die wesentlichen Anforderungen an neu zu beschaffende Komponenten bzw. wichtige Stichworte nochmals zusammen:

- **Skalierbarkeit** (für zeitlich verteilte Anschaffung)
- **SAN** (für intelligente, flexible Speicherlösungen)
- **ScaleOut** (für intelligente Architekturen)
- **Standards** (für Unabhängigkeit von unübersichtlichen Entwicklungen)
- **Silicon** (für niedrige Kosten)
- **Sparsamkeit**

Das sind die Goldenen sechs „S“.

Der letzte Punkt soll einfach daran erinnern, dass man bei allen Funktionen und Geräten, die man so einkauft, genau überlegen muss, ob man sie denn tatsächlich auch wirklich benötigt. Denn letztlich stecken Unternehmen viel Geld in zukünftigen Elektronikschrott.

Literatur/Quellen

- /1/ Dr. Jürgen Suppan, Report Cloud Computing, ComConsult Research 2011
- /2/ Dr. Jürgen Suppan, Cloud Storage in der Analyse, Netzwerk-Insider August 2011
- /3/ <http://ns1758.ca/winch/winchest.html>
- /4/ M. Komorowski, A history of Storage Cost, 2009, www.mkomo.com/cost-per-gigabyte
- /5/ Dr. F.-J. Kauffels, Report RZ-Netzwerk-Infrastruktur-Redesign, ComConsult Research 2011
- /6/ Przemek Radzikowski, SAN vs DAS A Cost Analysis of Storage in the Enterprise, capitalhead.com, 2010
- /7/ Dr. B. Moayeri, Wie viel Delay ist tolerierbar? Netzwerk-Insider August 2011

Drum prüfe, wer sich ewig bindet

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

Haben sie schon einmal über die Beschaffung von Analysetechnik nachgedacht? Nein, es geht nicht um den Ersatz für Ihren portablen Protokollanalytiker (Sniffer oder ähnliches), sondern um Geräte zur Hochleistungs-Paketzeichnung. Wenn Sie beispielsweise in der komplexen Welt Ihrer Rechenzentren Kommunikationsvorgänge analysieren und wirklich verstehen wollen, müssen Sie letztlich auf die Datenpakete schauen. Sie müssen die Daten mit den im RZ auftretenden Bitraten – im allgemeinen also mit 10 GBit/s oder Vielfachen davon – aufzeichnen und vor allem auch abspeichern können.

Dazu benötigen Sie Systeme mit entsprechenden Schnittstellen und viel Speicher. Eine Stunde Datenübertragung mit 10 GBit/s entspricht eben 4,5 TByte. Außerdem brauchen Sie eine Software, die Sie in die Lage versetzt, interessante Kommunikationsvorgänge und eventuelle Fehlerereignisse herausfiltern zu können. Und wenn Sie schon Datenpakete aufzeichnen, könnten Sie diese gleich noch auf sicherheitsrelevante Ereignisse überprüfen. Sie wären dann in der Lage, im Zweifel eine Post-Mortem-Analyse durchzuführen und beispielsweise die Wege herauszufinden, über die kritische Daten nach Außen gelangt sind (vgl. hierzu den „Standpunkt“ von Dr. Simon Hoff im Netzwerk Insider vom April 2011).

Solche Systeme, also die Kombination aus High-Speed-Netzwerkkarten, leistungsfähiger Appliance und massenhaft Festplattenspeicher sind sehr teuer. Eine Investition in diese Analysetechnik verschlingt ohne weiteres fünf- bis siebenstelligen Summen. Man geht also eine langfristige Bindung mit dem Hersteller ein, die vorab gut überlegt sein sollte. Zu diesem Zweck haben wir in den vergangenen Wochen ein Proof of Concept mit mehreren solcher Systeme in einem Kundennetz durchgeführt. Dabei sind spannende und teilweise erstaunliche Ergebnisse ans Licht gekommen.



Erstaunt haben uns beispielsweise die Performance-Unterschiede. Das einfache „Wegschreiben“ der Daten gelang allen von uns untersuchten Systemen mit vergleichbarer Performance. Kommunizierten jedoch innerhalb kurzer Zeit sehr viele Stationen miteinander, sank die Performance bei solchen Analyse-Systemen dramatisch ab, die bereits während der Paketaufzeichnung Statistiken erzeugen und als „Metadaten“ abspeichern. Diese Systeme waren dagegen unschlagbar schnell, wenn es darum ging, bestimmte Daten aus einer Terabyte-großen Paketaufzeichnung herauszufiltern.

Andererseits, was nützen einem die tollsten Performance-Werte, wenn man gar nicht weiß, welche Pakete man aus dem Datenstrom heraussuchen soll. Hierfür sind einerseits Statistiken nützlich und andererseits so genannte Experten, die bestimmte

Ereignisse in Datenströmen feststellen und melden. Und auch in dieser Beziehung unterschieden sich die Systeme im Test erheblich. War ein System stark auf das Trouble Shooting zugeschnitten und glänzte mit einem guten Experten, konnte man mit einem anderen System fast beliebige Statistiken in bunten Kurven und Kuchendiagrammen auf benutzerspezifischen „Dashboards“ zusammenstellen. Bot ein System viele nützliche Darstellungen bereits vor-konfiguriert, war dafür auf einem anderen System ein erheblicher Aufwand inklusive Hersteller-Unterstützung vonnöten.

Letztlich kommt es auf die typische Arbeitsweise an, mit der ein bestimmtes Analyssystem zu bedienen ist. Und darauf, wie sich diese Arbeitsweise mit Ihren Ansprüchen und Vorlieben deckt. Unsere Tester beispielsweise wollten alarmiert werden, wenn zu viele TCP-Retransmissions auftraten. Dann sollte im entsprechenden Zeitraum ermittelt werden, welche Clients und Server davon am meisten betroffen waren, um das Problem räumlich eingrenzen zu können. Mittels Paketanalyse besonders betroffener Verbindungen sollten anschließend Ursachen nachgewiesen werden.

In Ihrer Umgebung stellen sich ganz andere Fragen und Sie suchen das passende Werkzeug? Machen Sie die Probe aufs Exempel und testen Sie vor Ort in Ihrer Live-Umgebung. Die Herstellerpräsentation reicht nicht. Beim Autokauf ist die persönliche Probefahrt eine Selbstverständlichkeit. Warum nicht auch bei der Beschaffung von Analysetechnik, deren Preis den eines Autos zuweilen um ein Vielfaches übersteigt?

Seminar

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.

Rferenten: Dipl.-Inform. Oliver Flüs, Dr.-Ing. Joachim Wetzlar
Preis: € 2.290,- zzgl. MwSt.



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

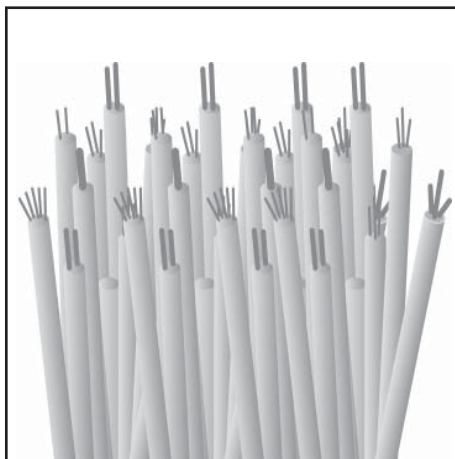
Sonderveranstaltung

Verkabelungssysteme für Lokale Netze, alles standardisiert, alles klar?

Die ComConsult Akademie veranstaltet am 04.10.11 ihre Sonderveranstaltung „Verkabelungssysteme für Lokale Netze, alles standardisiert, alles klar?“ in Köln.

Dieses Seminar erklärt die Zusammenhänge der wichtigsten Standards und Normen, vergleicht diese mit dem aktuellen Stand der Technik und bewertet insbesondere die Praxistauglichkeit der im Normenumfeld getroffenen Empfehlungen. Neben einer Betrachtung des aktuellen Normungsstands aus der Sicht eines Normennutzers, der Bewertung von ausgewählten herstellerspezifischen Lösungen wird auch auf Planungs- und installationsbegleitende Maßnahmen eingegangen, die im Rahmen einer anstehenden Verkabelung zu beachten sind.

Um die Verkabelung im Bereich der Lokalen Netzen ist es ruhig geworden. Die wichtigsten Standards sind seit nun gut zwei Jahren verabschiedet, neue revolutionäre Technologien nicht in Sicht. Doch ist damit alles klar, reichen diese Spezifikationen aus, um eine zukunftssichere und langfristig nutzbare Verkabelung zu planen bzw. zu realisieren. Mitnichten! Der Projektalltag im Umfeld der Planung und Qualitätssicherung bei Verkabelungsprojekten zeigt häufig, wie wenig sowohl die Standards gelesen, richtig bewertet oder auch angewendet werden. Systematische Fehler bei Materialauswahl, Ausschreibung, Messverfahren, Abnahme und Dokumen-



tation reduzieren den Nutzbarkeitszeitraum der Kommunikationskabelanlage, führen zu frühen Neu- oder vermeidbaren Nachverkabelungen oder schließen die Einführung von moderneren Übertragungsverfahren mit hohen Datenraten aus.

Das Seminar geht unter anderem auf folgende Fragen ein:

- Mehr als 15 Jahre strukturierte Kommunikationsverkabelung, was waren die richtigen Entscheidungen, was waren die falschen Entscheidungen, welche Prognosen waren richtig, welche falsch, was ist für die Zukunft zu beachten, wo unterscheiden sich Theorie und Praxis?
- LWL-Messtechnik, welche Unterschiede gibt es, was sind die richtigen Methoden?
- Glasfaser bis zum Arbeitsplatz kontra

Standard-Kupferverkabelung, wann eignet sich welches Medium am besten, wie ist mit vorhandenen Glasfaserverkabelungen umzugehen?

- Warum gewinnt die Dämpfungsproblematik an Bedeutung, gerade in Zusammenhang mit der Nachverkabelung von Backbone- oder Rechenzentrums-Verkabelungen? Wieso beeinflusst die Dämpfung die verschiedenen möglichen Topologien?
- In welchem Zusammenhang stehen die Spezifikationen der Kabelnorm EN 50173 und die Normen der IEEE 802.3, warum ist das Verständnis dieses Zusammenhangs wichtig?
- Ist die Multimode am Ende? Wenn nein, wo macht sie weiterhin Sinn, in welcher Variante: OM2, OM3, OM3+, OM4? Wie ist die Singlemode in Zukunft zu bewerten, gibt es auch hier Unterschiede?
- Welche Steckertechnik bei Glasfaser ist zu empfehlen, wie sind die Herstellerangaben zu den optischen Eigenschaften zu bewerten? Gibt es Standards zur LWL-Steckertechnik?
- Welche Kupferkategorie ist für Kabel und Verbindungstechnik einzusetzen: 6, 6x, 7, 7x etc.? Und warum? Wo sind die Grenzen von Kupfer?
- Erfahrungen aus dem „Ausschreibungsalltag“ eines Fachplaners: Reichen die Vertragsbestandteile der VOB aus, was ist darüber hinaus festzulegen? Die unterschätzte Wichtigkeit der „Technischen Vorbemerkungen“ in Ausschreibungen, was sollte dazu gehören?

Fax-Antwort an ComConsult 02408/955-399

Anmeldung

Verkabelungssysteme für Lokale Netze

Ich buche das Seminar
**Verkabelungssysteme für Lokale Netze,
alles standardisiert, alles klar?**

☐ 04.10.11 in Köln

zum Preis von € 890,- zzgl. MwSt.



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

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift

Zweitthema

Glasfaserstecker, sinnvolle und weniger sinnvolle Auswahlkriterien

Fortsetzung von Seite 1



Dipl.-Ing. Hartmut Kell kann bis heute auf eine mehr als 20-jährige Berufserfahrung in dem Bereich der Datenkommunikation bei lokalen Netzen verweisen. Als Leiter des Competence Center IT-Infrastrukturen der ComConsult Beratung und Planung GmbH hat er umfangreiche Praxiserfahrungen bei der Planung, Projektüberwachung, Qualitätssicherung und Einmessung von Netzwerken gesammelt und vermittelt sein Fachwissen in Form von Publikationen und Seminaren.

Übertragungstechnische Eigenschaften

Wer könnte sich daran erinnern, dass es jemals „hitzige“ Diskussionen über LWL-Stecker gegeben hat, auch die Spezifikation von Güteklassen in der EN 50173 ist ohne Bedeutung geblieben (es gibt keine!). Der Grund dafür ist wohl darin zu suchen, dass ein LWL-Stecker im Prinzip durch die beiden optischen Kenngrößen Einfügedämpfung und Reflexionsdämpfung ausreichend spezifiziert ist, dabei spielt die Reflexionsdämpfung bei LAN-Zugangsverfahren so gut wie kaum eine Rolle. Der Standard lässt bei der Definition aller Streckenmodelle eine Einfügedämpfung von bis zu 0,75 dB pro Steckverbindung zu, dies ist in Anbetracht der sehr hochwertigen verfügbaren Typen und der hohen handwerklichen Qualitäten eher als „miserabler“ Dämpfungswert zu betrachten. Übliche Forderungen in Ausschreibungen liegen bei 0,4 dB bis 0,5 dB pro Steckverbindung, selbst das ist für kompetente Installationsfirmen keine wirkliche Herausforderung. Für niedrige Datenraten bis maximal 100 Mbit/s ist bei einfachen Verbindungen über ein LWL-Kabel hinweg ein ausreichendes Dämpfungsbudget vorhanden, eine Optimierung des Glasfasersteckers um weitere zehntel dB erscheint wenig sinnvoll. Doch bei höheren Datenraten und insbesondere bei Mehrfachdurchrangierungen spielen solche Differenzen möglicherweise doch eine entscheidende Rolle. Dies haben offensichtlich auch die Hersteller von Glasfaserstecker entdeckt und versuchen erfreulicherweise durch Einführung von normierten Güteklassen bei Glasfaserstecker die Qualität zu erhöhen.

Eine Schwierigkeit dabei ist es, ähnliche Erfahrungen hat die „Kupferfraktion“ schon beim Kupferstecker gemacht, im Datenblatt einen optischen Kennwert

anzugeben, der auf Basis eines **einheitlichen** Messverfahrens entstanden ist und damit die **Vergleichbarkeit** der verschiedenen Typen ermöglicht. Die Normenreihe IEC 61753 (Festlegung der Messverfahren und der optischen Werte) und IEC 61755 (geometrische Parameter wie z.B. die Lage des Faserkerns) versuchen dieses Ziel einer einheitlichen Spezifikation von glasfaserbasierenden passiven Verbindungsstellen zu erreichen. Die IEC 61753 legt dabei Qualitätsklassen für Dämpfung (Insertion Loss) und Rückflussdämpfung (Return Loss) fest, 4 Klassen von A bis D für Singlemode und 1 Klasse M für Multimode werden definiert (letzte noch nicht verabschiedet). Leider sind diese Klassen in den letzten Überarbeitungen der EN 50173-1 noch nicht eingeflossen und auch die Hersteller selber gehen noch sehr zögerlich mit der Kennzeichnung ihrer Produkte damit um. Dennoch lohnt es sich, kurz die aktuellen Tabellen zu betrachten, um die grundsätzliche Vorgehensweise der Normierung

kennenzulernen. (siehe Tabelle 1)

Festgelegt werden maximale und typische Einfügedämpfungen in einer A- bis D-Klassifizierung sowie minimale Reflexionsdämpfungen in einer 1- bis 4-Klassifizierung, je nach Qualität des Steckers lassen sich diese Anforderungen unterschiedlich kombinieren. Beispielsweise hätte ein A1-Stecker (Einfügedämpfung max. 0,15 dB / Reflexionsdämpfung min. 60 dB) eine geringere Qualität als ein C2-Stecker (Einfügedämpfung max. 0,5 dB / Reflexionsdämpfung min. 45 dB); nicht jede Kombination ist technisch machbar und damit sinnvoll. Wichtig dabei ist die Forderung, dass die maximale Dämpfung für 97% aller Kombinationen dieses Steckers mit einem beliebigen anderen Stecker derselben Klasse sichergestellt sein muss. Das bedeutet aber auch, dass durchaus 3 von 100 produzierten Steckern oberhalb der spezifizierten Einfügedämpfung liegen dürfen und der Steckertyp trotzdem die entsprechende

Einfügedämpfungs-Qualitätsklasse	Einfügedämpfung (maximal)	Einfügedämpfung (typisch)
A	max. 0,15 dB	max. 0,07 dB
B	max. 0,25 dB	max. 0,12 dB
C	max. 0,50 dB	max. 0,25 dB
D	max. 1,0 dB	max. 0,5 dB

Rückflussdämpfungs-Qualitätsklasse	Rückflussdämpfung
1	min. 60 dB
2	min. 45 dB
3	min. 35 dB
4	min. 26 dB

Tabelle 1: Qualitätsklassen von LWL-Singlemodestecker

Glasfasterstecker, sinnvolle und weniger sinnvolle Auswahlkriterien

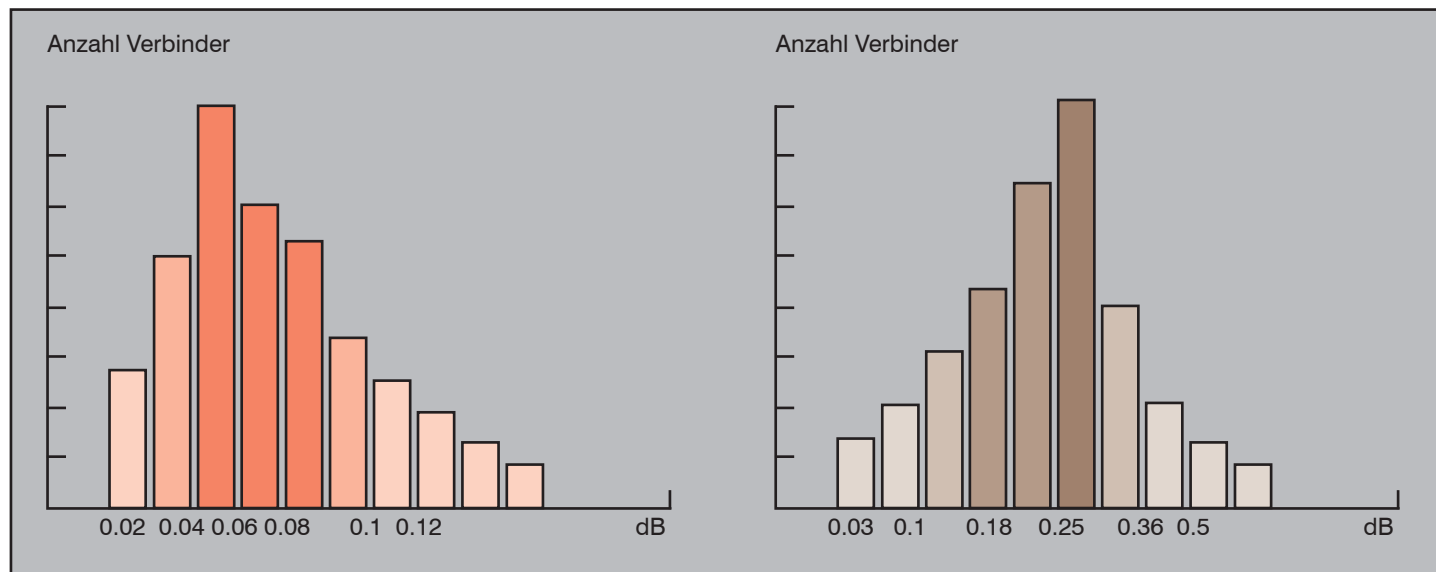


Abbildung 1: Beispiele für Häufigkeitsverteilung unterschiedlicher Qualitäten

Klassifizierung haben darf bzw. nicht mangelhaft ist. In der Praxis hätte dies zur Konsequenz, dass ein einzelner schlechter Stecker nicht reklamiert werden kann, wenn er die Grade-Spezifikationen nicht erfüllen würde und diese die alleinige Anforderung z.B. einer Ausschreibung wäre. Nur Abweichungen der Häufigkeitsverteilung der Einfügedämpfung bedeuten mangelhafte Qualität.

Zusammengefasst bedeutet das also, man entscheidet sich in Zusammenhang mit der Fabrikatsauswahl im Prinzip nicht für einen Stecker, der einen garantierten maximalen Dämpfungswert hat, sondern für einen Stecker mit einer bestimmten Häufigkeitsverteilung der gemessenen Dämpfungswerte. Die Abbildung 1 gibt beispielhaft den Qualitätsunterschied zwischen einem 0,1dB-Stecker (dazu siehe unten mehr) und einem „Standard-Stecker“ wieder.

Die Erfahrung im Bereich der Abnahmemessungen zeigten ohnehin, dass die Erwartungshaltung an die Genauigkeit der gemessenen Werte häufig viel zu hoch ist. Einige Fehlerquellen sind unvermeidbar, wodurch selbst bei sorgfältigster Messung eine Minimalabweichung einzukalkulieren ist. Wer kann z.B. beurteilen, ob die Innenflächen der Kupplung oder die Außenflächen der Ferrulen sauber sind. Empfohlen wird die Akzeptanz einer Abweichung von mindestens $\pm 0,1$ dB.

Die Idee der Norm ist es weiterhin, die typische Einfügedämpfung (= Mittelwert der produzierten Stecker) als „Rechengröße“ bei der planerischen Ermittlung des Dämpfungsbudgets heranzuziehen, z.B. 4 Steckverbindungen mit Klasse A dürfen nicht mehr als 0,28 dB Dämpfung ergeben.

Damit verlieren die bisherigen Dämpfungangaben der Hersteller, die sich sehr oft auf eine Kombination mit einem extrem hochwertigen eigenen Referenzstecker bezogen, an Aussagekraft bzw. an Nutzbarkeit für die Planung. Das Ganze erinnert etwas an die Verfahrensweise im Bereich der RJ45-Normierung, auch hier war es ein langer Weg, bis man Einigung gefunden hat, mit welcher oder wie vielen Bezugsbuchsen der im Datenblatt angegebene Wert gemessen wurde (De-Embedded-Diskussion).

Schaut man sich die aktuellen Datenblätter der verschiedenen Hersteller an, so stellt man fest, dass die Qualifizierung der Stecker auf Basis dieser IEC-Normen nur sehr zögerlich bzw. noch nicht weit verbreitet ist. Hersteller, welche diese hohen

Qualitätsstandards massiv propagieren sind z.B. Reichle de Masari und Huber & Suhner.

Einige Hersteller definieren für eine Auswahl ihrer Produkte eine sehr hochwertige „0,1dB-Klasse“, bei der über eine sehr große Anzahl von Stichproben mit zufälligen Kombinationen der eigenen Stecker der Mittelwert der Einfügedämpfungen unter 0,1dB liegt (siehe Beispiel im Bild oben). Immer häufiger werden diese High-Tec-Stecker auch im LAN-Umfeld gefordert, doch ist das sinnvoll? Beim Einsatz dieser Systeme muss immer klar bleiben, dass eine nachträgliche messtechnische Beurteilung dieser Qualität alleine schon wegen der oben beschriebenen messtechnisch bedingten 0,1dB-Toleranz kaum möglich ist; der akzeptierte Messfehler ist

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: € 2.390,- zzgl. MwSt. bzw. € 1.990,- zzgl. MwSt. ohn Workshop



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

Glasfasterstecker, sinnvolle und weniger sinnvolle Auswahlkriterien

genauso groß wie die Messgröße. Darüber hinaus sollte abgewogen werden, welchen Mehrwert eine solche Steckerklasse im Vergleich z.B. zur „geringwertigeren“ Singlemode LAN-Eco Klasse des Herstellers Reichle de Masari mit 0,2 dB bringen soll; bringt dieses 10tel dB z.B. bei einem Dämpfungsbudget von 6 dB bei 10GBa-se-LR tatsächlich Planungsvorteile?

Es kann also festgehalten werden, dass die Transparenz der unterschiedlichen Qualitäten bei Singlemodesteckern mit Hilfe dieser IEC-Normen zunehmen könnte, falls sie auf der einen Seite von den Herstellern angenommen würde und auf der anderen Seite von den Planern z.B. im Rahmen von Ausschreibungstexten eingefordert werden würde. Leider ist diese Kategorisierung bei den Multimodefaseren noch sehr am Anfang der Standardisierung, auch hier wäre ein ähnliches Modell wünschenswert, insbesondere deshalb, weil das Dämpfungsbudget viel kleiner ist und damit die Bedeutung der Steckerqualität viel höher.

Weitere Eigenschaften

Ein Stecker muss nicht nur übertragungstechnischen Anforderungen genügen, sondern er muss auch dafür sorgen, dass die mechanischen Belastungen bei einem Steckvorgang nicht zu einer Zerstörung der Faser führen. Prinzipiell ergeben sich neben den optischen Forderungen folgende weitere Forderungen an einen Stecker:

- universelle Einsetzbarkeit,
- niedrige Kosten,
- hohe Steckzyklenanzahl mit geringer Dämpfungsänderung,
- gute Handhabbarkeit,
- geringer Platzbedarf.

Die Universalität eines Steckverbinders leitet sich daraus ab, dass er in möglichst vielen Umgebungen unter Sicherstellung von vielen Übertragungstechniken eingesetzt werden kann. Viele Nutzer definieren daraus die Anforderung, dass der Steckverbinder passend zu den Anschlüssen der aktiven Komponenten sein muss. Aufgrund der Erfahrungen der letzten zwanzig Jahre bezweifelt der Autor, dass dies der richtige Ansatz ist.

Kurz vor der Einführung der Standards EN 50173 galt der ST-Steckverbinder als Standardstecker im LAN-Bereich, viele Hubs und ältere Switches boten diese Schnittstelle. Die einfache Handhabung des SC-Duplex-Steckverbinders führte auf der einen Seite zu einer Präferenzierung in der EN 50173 und auf der anderen Seite zu einem Wechsel der Anschlussschnittstellen bei

den neuen aktiven Systemen; GBIC-Technologie setzt auf SC-Duplex-Anschlüsse. Die Ende der 90er-Jahre erfolgte Einführung der Small-Form-Factor-Technologie reduzierte die geometrischen Maße eines LWL-Duplex-Anschlusses auf RJ45-Maße, öffnete damit den Weg zu hohen Packungsdichten im Rangierfeldbereich und insbesondere bei den aktiven Komponenten. Moderne Switches sind mit SFPs (Small Formfactor Pluggable) ausgestattet, deren Basis der LC-Steckverbinder ist. Folgte man also dem Trend, seine Rangierfelder immer dem aktuellen Stand der Anschlusstechnik auf den aktiven Komponenten anzupassen, hätte man über einen Zeitraum von ca. 15 Jahren 3 verschiedenen Techniken eingesetzt, ohne dass vor allem die letzten Wechsel einen technischen Vorteil gebracht hätten (Ausnahme: Packungsdichte). Aus Sicht des Autors ist es viel sinnvoller, ein System mit guten mechanischen und optischen Eigenschaften auszuwählen, diesem möglichst lange „treu zu bleiben“ und die Adapterkabel zu akzeptieren. Auch ein heute eingeführter hochmoderner LC wird wohl kaum eine Garantie abgeben, dass nicht in den nächsten Jahren wieder ein neuer Steckertyp bei den aktiven Systemen eingeführt wird.

Mit Ausnahme des Umfelds von Arbeitsplatzverkabelungen mit einer sehr hohen Anzahl von Anschlüssen sollte bei der Bewertung der Kosten im Backbone-Bereich wohl eher von verhältnismäßig geringer Steckeranzahl ausgegangen werden. Da aber gerade diese Verbindungen oftmals eine hohe Zuverlässigkeit erfordern, ist eine zu starke Kostenbewertung aus Sicht des Autors nicht zu empfehlen.

Auch die Steckzyklenanzahl ist für den Backbone-Bereich nicht zu überbewerten, die herstellerteitigen Zusagen zu mehr als 500 Steckzyklen (teilweise bis 1000) werden wohl in den wenigsten Umgebungen benötigt.

Eine Beurteilung der guten Handbarkeit setzt sich aus vielen Faktoren zusammen.

- Wie lässt sich ein Stecker in einem sehr

dicht bestückten Rangierfeld stecken und insbesondere „ent-stecken“?

- Wie kann verhindert werden, dass falsche Fasern miteinander verbunden werden, z.B. Singlemode und Multimode oder OM1 und OM2?
- Wie erfolgt der Staubschutz bzw. allgemein der Schutz der Ferrulen, gibt es automatischen Staubschutz, müssen lose Staubschutzkappen aufgesteckt werden?

Steckerübersicht

Neben der Betrachtung der optischen und normativen Übertragungsqualität von Glasfastersteckern bleibt in vielen Projekten die Frage offen, welcher Typ ist denn der richtige Stecker bzw. lohnt sich ein Wechsel des bisherigen Steckertyps. Die EN 50173 gibt dazu keine wirkliche Hilfe, Zitat: „Für die Verbindungstechnik können alle von IEC oder CENELEC genormten Lichtwellenleiter-Steckverbinderarten ausgewählt werden“. In der Zwischenzeit reduziert sich die Auswahl der am weitesten verbreiteten Stecker mit Zukunftspotenzial aber auf 3 Typen:

Der **E2000-Stecker** (in der Duplex-Variante standardisiert unter IEC 61754-15) bietet viele technische Vorteile, u.a. eine mechanische Kodierbarkeit, einen automatisch schließenden Ferrulenschutz (Staubschutz) und in der Regel extrem gute optische Kennwerte. Er ist sowohl in einer breiteren Version wie auch in einer Small-Form-Factor-Version erhältlich (SFF). Nachteilig sind

- der hohe Preis,
- die Situation, dass nur drei Hersteller ihn produzieren,
- die hohe Anzahl von nicht grundsätzlich zueinander kompatiblen Untervarianten,
- und der Einsatzfokus Europa.

Der **SC-Stecker** (standardisiert in IEC 60874-19) ist sehr weit verbreitet und darf als sehr zuverlässig charakterisiert wer-

	SC-Stecker	E2000	LC-Stecker
Normierung	++	+	+
Maße	-	++	++
Ferrulenschutz	-	++	-
Verbreitung (international)	++	-	++
Herstelleranzahl	++	+	++
Kodierungsmöglichkeiten	-	++	+

+*: herstellereabhängig möglich

Tabelle 2: Bewertung der Stecker

Glasfaserstecker, sinnvolle und weniger sinnvolle Auswahlkriterien

den. Als Nachteil sind die sehr großen geometrischen Maße zu nennen, dies ist mit Sicherheit einer der Gründe, warum die Verkaufszahlen rückgängig sind. Nebenbei ist zu beachten, dass viele Hersteller nach Drehung der beiden Stecker ein Zusammenclipsen nicht mehr erlauben, der Grund hierfür liegt darin, dass der Clipsvorgang mit Hilfe einer Nut am Stecker erfolgt.

Der **LC-Stecker** (standardisiert in der IEC 61754-20) ist deutlich kleiner als der über viele Jahre im Standard bevorzugte SC-Stecker. Man darf annehmen, dass entsprechende herstellerpolitische Einflüsse dazu geführt haben, dass der nicht unbedingt beste Stecker auf dem Markt sich zu einem Standardstecker bei den aktiven Komponenten entwickelt (SFPs haben LC-Buchsen). Deshalb wird jeder Nutzer von neueren Switches wohl kaum den LC-Stecker vermeiden können. Die Beurteilung des LCs durch die Nutzer selber ist sehr kontrovers, gerade die mechanische Qualität wird häufig bemängelt. Im Gegensatz zum E2000

sind in der Regel keine Kodierungsmöglichkeiten und keine automatische Staubschutzklappen verfügbar. Eine Ausnahme stellt hier der F3000 von Diamond dar, er besitzt eine automatisch schließende Staubschutzklappe.

Alle, die der Multimodefaser auch in Zukunft bei Datenraten ab 40 Gbit/s treu bleiben wollen, müssen sich darüber hinaus mit dem mehrfaserigen **Multipath Push-On-Stecker** vertraut machen. Dieser ist standardisiert in der IEC61754-7 und kann typischerweise 12 oder 24 Fasern aufnehmen, bis zu 72 Fasern sind möglich. Er ist ebenfalls wie die meisten anderen Typen sowohl als PC- (Geradschliff) als auch als APC-Variante (Schrägschliff) verfügbar. Eine Notwendigkeit zum Einsatz dieses Steckers wird in erster Linie im Rechenzentrum zu erwarten sein. Diese Einführung verlangt aber ein völliges Umdenken bei der Planung einer solchen Infrastruktur, nur zu nennen ist z.B. die Einführung von gedrehten Kabeln („Entclipsen“ zum Drehen ist nicht mehr möglich) oder die Fra-

ge, wie solche Strecken einzumessen sind.

Empfehlungen

Eine eindeutige Empfehlung zu einem der genannten Typen im Einsatzbereich Rangierfeld kann der Autor nicht machen, tendenziell zeigt der E2000 aber langfristig nutzbare Vorteile, diese müssen den Mehrkosten gegenübergestellt werden. In einem Punkt kann es aber eine klare Empfehlung geben: Falls ein etabliertes Steckersystem im Rangierfeldbereich vorhanden ist, man mit diesem zufrieden ist und dieses den übertragungstechnischen Anforderungen genügt, so sollte man nicht den „Modetrends“ der Switch-Hersteller oder auch der Verkabelungsnormen folgen und dieses System ersetzen oder um ein neueres System ergänzen. Dieses führt nur zu einer unnötigen hohen Anzahl von verschiedenen Adapterkabeln und die Vergangenheit zeigt, dass selbst die Verkabelungsnormen hier keine ausreichende Nachhaltigkeit bieten um dieses zu vermeiden.

Sonderveranstaltung**Verkabelungssysteme für Lokale Netze, alles standardisiert, alles klar? 04.10.11 in Köln**

Das Seminar geht unter anderem auf folgende Fragen ein:

- Mehr als 15 Jahre strukturierte Kommunikationsverkabelung, was waren die richtigen Entscheidungen, was waren die falschen Entscheidungen, welche Prognosen waren richtig, welche falsch, was ist für die Zukunft zu beachten, wo unterscheiden sich Theorie und Praxis?
- LWL-Messtechnik, welche Unterschiede gibt es, was sind die richtigen Methoden?
- Glasfaser bis zum Arbeitsplatz kontra Standard-Kupferverkabelung, wann eignet sich welches Medium am besten, wie ist mit vorhandenen Glasfaserverkabelungen umzugehen?
- Warum gewinnt die Dämpfungsproblematik an Bedeutung, gerade in Zusammenhang mit der Nachverkabelung von Backbone- oder Rechenzentrums-Verkabelungen? Wieso beeinflusst die Dämpfung die verschiedenen möglichen Topologien?
- In welchem Zusammenhang stehen die Spezifikationen der Kabelnorm EN 50173 und die Normen der IEEE 802.3, warum ist das Verständnis dieses Zusammenhangs wichtig?
- Ist die Multimode am Ende? Wenn nein, wo macht sie weiterhin Sinn, in welcher Variante: OM2, OM3, OM3+, OM4? Wie ist die Singlemode in Zukunft zu bewerten, gibt es auch hier Unterschiede?
- Welche Steckertechnik bei Glasfaser ist zu empfehlen, wie sind die Herstellerangaben zu den optischen Eigenschaften zu bewerten? Gibt es Standards zur LWL-Steckertechnik?
- Welche Kupferkategorie ist für Kabel und Verbindungstechnik einzusetzen: 6, 6x, 7, 7x etc.? Und warum? Wo sind die Grenzen von Kupfer?
- Erfahrungen aus dem „Ausschreibungsalltag“ eines Fachplaners: Reichen die Vertragsbestandteile der VOB aus, was ist darüber hinaus festzulegen? Die unterschätzte Wichtigkeit der „Technischen Vorbemerkungen“ in Ausschreibungen, was sollte dazu gehören?

Referent: Dipl.-Ing. Hartmut Kell

Preis: € 890,- zzgl. MwSt.



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

Aktuelle Veranstaltungen

UC - Cisco versus Microsoft - Wer hat die bessere Unified-Communications-Lösung?, 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.

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 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.

Preis: € 890,- - 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.

Verkabelungssysteme für Lokale Netze, alles standardisiert, alles klar?, 04.10.11 in Köln

Dieses Seminar erklärt die Zusammenhänge der wichtigsten Standards und Normen, vergleicht diese mit dem aktuellen Stand der Technik und bewertet insbesondere die Praxistauglichkeit der im Normenumfeld getroffenen Empfehlungen. Neben einer Betrachtung des aktuellen Normungsstands aus der Sicht eines Normennutzers, der Bewertung von ausgewählten herstellerspezifischen Lösungen wird auch auf Planungs- und installationsbegleitende Maßnahmen eingegangen, die im Rahmen einer anstehenden Verkabelung zu beachten sind.

Preis: € 890,- - zzgl. MwSt.

Recht und Datenschutz bei Einführung von Voice over IP, 04.10. - 05.10.11 in Berlin

Dieses Seminar erklärt die Zusammenhänge der wichtigsten Standards und Normen, vergleicht diese mit dem aktuellen Stand der Technik und bewertet insbesondere die Praxistauglichkeit der im Normenumfeld getroffenen Empfehlungen. Neben einer Betrachtung des aktuellen Normungsstands aus der Sicht eines Normennutzers, der Bewertung von ausgewählten herstellerspezifischen Lösungen wird auch auf Planungs- und installationsbegleitende Maßnahmen eingegangen, die im Rahmen einer anstehenden Verkabelung zu beachten sind.

Preis: € 1.490,- - zzgl. MwSt.

Interne Absicherung der IT-Infrastruktur, 04.10. - 06.10.11 in Berlin

Bedingt durch Netzkonvergenz, Mobilität und Virtualisierung hat die interne Absicherung der IT-Infrastruktur in den letzten Jahren enorm an Bedeutung gewonnen. Heterogene Nutzergruppen mit unterschiedlichstem Sicherheitsniveau teilen sich eine gemeinsame IP-basierte Infrastruktur und in vielen Fällen ist der Aufbau sicherer, mandantenfähiger Netze notwendig. Dieses Seminar identifiziert die wesentlichen Gefahrenbereiche und zeigt effiziente und wirtschaftliche Maßnahmen zur Umsetzung einer erfolgreichen Lösung auf. Alle wichtigen Bausteine zur Absicherung von LAN, WAN, Endgeräten, RZ-Bereichen, Servern und SAN werden detailliert erklärt und anhand konkreter Projektbeispiele wird der Weg zu einer erfolgreichen Sicherheits-Lösung aufgezeigt.

Preis: € 1.890,- - zzgl. MwSt.

Virtualisierungstechnologien in der Analyse, 04.10. - 06.10.11 in Berlin

Dieses Seminar analysiert die verfügbaren Virtualisierungstechnologien der führenden Anbieter. Sie lernen, welche Gestaltungselemente virtuelle Umgebungen haben, angefangen von einfachen und überschaubaren Lösungen bis hin zu komplexen und umfassenden Rechenzentrums-Gesamt-Architekturen. Dabei wird auch der Bedarf an Infrastruktur-Leistung insbesondere auf der Netzwerk- und Speicherseite untersucht.

Preis: € 1.890,- - zzgl. MwSt.

WAN: Aktuelle Technologie und Erfahrungen aus Ausschreibungen, 04.10. - 05.10.11 in Berlin

Das Programm des Seminars „WAN: Neue Verfahren und Erfahrungen aus Ausschreibungen“ bietet wertvolle Tipps und Empfehlungen sowohl zu technischen als auch zu organisatorischen Aspekten der Konzeption, der Planung, der Ausschreibung und des Betriebs von Wide Area Networks. Die Referenten des Seminars blicken auf langjährige Erfahrungen im WAN-Bereich zurück und vermitteln im Seminar Erkenntnisse aus Dutzenden von Projekten, in denen Wide Area Networks entworfen, ausgeschrieben und optimiert wurden. Der große Erfahrungsschatz von ComConsult bei der Lösung von Problemen und der Lokalisierung von Fehlern in standortübergreifenden Netzen fließt ebenso in das Seminarprogramm ein wie die Expertise der Referenten bei der Gestaltung sinnvoller Service Level Agreements (SLA) im WAN-Betrieb.

Preis: € 1.490,- - zzgl. MwSt.

Zertifizierungen

ComConsult Certified Network Engineer

Lokale Netze

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

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