

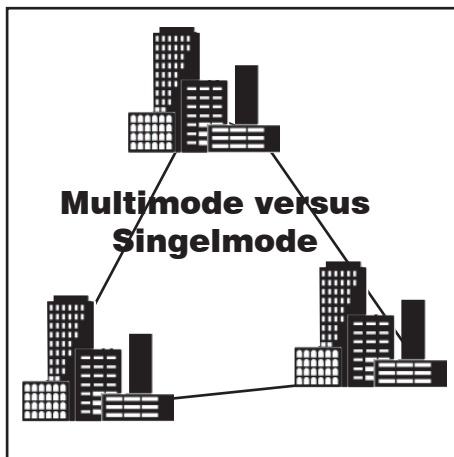
Schwerpunktthema

Singlemode und Multimode in der Backbone-Verkabelung

von Dipl.-Ing. Hartmut Kell

Für den Aufbau einer Verkabelung zur Verbindung von Verteilern, allgemein als Backbone-Verkabelung bezeichnet, hat sich bereits in den 90er-Jahren die Glasfasertechnik als Hauptmedium etabliert. Die Überarbeitung der damit verbundenen Kabelstandards ist im Vergleich zur Kupfertertiärverkabelung eher behäbig vonstatten gegangen.

Der Grund dafür ist, dass die Anforderungen der im Backbone-Bereich benötigten Datenraten an die Verkabelung „bescheiden“ waren und ein permanentes Nacharbeiten nicht notwendig wurde; das galt bis zur Einführung von Gigabit-Ethernet. Mit der Einführung von Gigabit-Ether-



net insbesondere der Einführung von noch höheren Datenraten änderten sich die Anforderungen sprunghaft, so dass sowohl die Nutzbarkeit der vorhandenen Verkabelung als auch die Anforderungen an eine neue Backbone-Verkabelung teilweise völlig neu analysiert werden muss. Häufig werden bei dieser Analyse die physikalischen Anforderungen der Ethernet-Zugangsverfahren zu wenig betrachtet und man wird trotz einer Erfüllung der Kabelnormen bei der Verkabelungsnutzung mit erheblichen Beeinträchtigung konfrontiert.

weiter auf Seite 13

Zweitthema

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

von Dr. Franz-Joachim Kauffels

Die im Consumer-Bereich schon länger verbreiteten SSDs haben eine Reife und ein Produktspektrum erlangt, die sie auch für viele Anwendungen im Enterprise-Umfeld interessant macht. Mittelfristig werden sie mindestens die Hälfte der bisherigen Platten ablösen.

Warum das so ist, kann man nur verstehen, wenn man in die spezielle Technologie dieser Geräte einsteigt. Genau die möchte dieser Artikel mit relativ anschaulichen Erklärungen verdeutlichen. Die weitreichende Einführung von SSDs wird fürchterliche Konsequenzen für bestehende

Netzwerk-Strukturen im RZ haben, denn im Reifegrad und in Folge bei Übertragungsgeschwindigkeit und Latenz liegen zwischen der SSD-Technologie und der aktuellen Netzwerk-Technologie Welten.

weiter auf Seite 23

Geleit

RZ Infrastruktur-Redesign: brauchen wir das überhaupt?

ab Seite 2

Aktuelles Seminar

Sommerschule 2011 - Intensiv-Update auf den neuesten Stand der Netzwerktechnik

ab Seite 9

Sonder-Veranstaltung

Verkabelungssysteme für Lokale Netze

auf Seite 11

Standpunkt

Standards sichern Interoperabilität! Sicher?

auf Seite 22

Zum Geleit

RZ Infrastruktur-Redesign: brauchen wir das überhaupt?

Intel kündigt permanent Chips mit geringerem Stromverbrauch an, der Übergang auf 10 Gigabit ist sowieso schon erfolgt, FCoE dümpelt vor sich hin und wir verlagern unser RZ ohnehin aus Kostengründen in die Cloud. Das ist eine der möglichen Sichtweisen. Niemand braucht neue Infrastrukturen. Aber ist diese Sichtweise wirklich haltbar?

Fangen wir mit der Cloud an und betrachten Zynga. Zynga ist als Anbieter von Webspielen (Farmville etc.) eines der Unternehmen, denen man immer unterstellt, dass die Cloud die perfekte Lösung für deren IT-Bedarf ist. Keine notwendigen Vorinvestitionen und nur die wirklich benötigten Ressourcen müssen bezahlt werden. Nun hat es in der amerikanischen Presse in den letzten Wochen eine interessante Diskussion über Zynga gegeben. Tatsächlich benutzt Zynga Amazon-Dienste nur im Rahmen des Aufbaus eines neuen Produkts. Sobald klar ist, dass das Produkt erfolgreich ist und eine kritische Masse von Benutzern erreicht ist, schaltet Zynga auf die eigenen Rechenzentren um. Diese basieren auf hochgradig optimierten virtuellen Umgebungen (in diesem Fall ZEN ergänzt um selbst entwickelte Software). Klare Aussage: Amazon ist zu teuer für ein etabliertes Produkt. Unter Nutzung derselben Technik, die auch Amazon als Cloud-Anbieter einsetzt, kommt man im eigenen Unternehmen zu einer technisch und wirtschaftlich besseren Lösung. Das deckt sich mit der



Cloud-Bewertung von ComConsult Research (siehe die Studie: "Public und Private Clouds in der Analyse" bei ComConsult Research).

Stellen wir also erst einmal fest: unser Rechenzentrum ist lebendiger denn je. Die Cloud mag interessante Zusatz-Optionen liefern, aber sie stellt die grundsätzliche Bedeutung des Rechenzentrums nicht in Frage.

Welche Auswirkungen werden AMD und Intel auf die Entwicklung im RZ haben? Nun, viele Neuentwicklungen großer Applikationen werden auf virtuellen Architekturen aufsetzen. Sie skalieren besser und harmonisieren somit perfekt mit den Skalierungseigenschaften von Standard-Ser-

ver-Technik. Die Bedeutung von virtueller Server-Technologie wird also weiter zunehmen. Die zentrale Frage ist: was machen wir in Zukunft mit den vielen Transistoren, die uns 22nm und 11nm-Technik in den nächsten 5 Jahren liefern werden? Schon heute überschreitet die Zahl theoretisch möglicher Kerne pro Blade das Leistungsvermögen von mindestens einem wichtigen Server-Betriebssystem (sind wir wirklich überrascht? Was treiben die Microsoft-Entwickler eigentlich mit dem vielen Geld, das sie verdienen?). Aber selbst wenn in Zukunft die Betriebssysteme viel mehr Kerne unterstützen können, macht das überhaupt Sinn, steigt nicht die Verlustleistung in der Kernverwaltung so an, dass der weitere Zugewinn zu klein wird?

Parallel werden wir vermutlich immer mehr RAM zu immer besseren Preisen bekommen. Kombiniert man das mit dem generellen Trend zu 64 Bit-Architekturen, dann ist das eine wichtige Entwicklung. Applikationen können den im Prinzip nun unendlich großen Speicher besser zur Optimierung der Performance nutzen und die Abhängigkeit von der SAN-Performance senken. Das ist vermutlich deutlich wichtiger als die Erhöhung der Kernanzahl.

Aber trotzdem: betrachten wir die Entwicklung der letzten 30 Jahre, dann würden Rechenzentren der 70er Jahre heute in einen kleinen Teil eines 19 Zoll Schrankes passen. Die permanente Verdopplung der Transistorzahl alle 18 Monate (Moore's Law) muss diese Entwicklung beschleunigen (denken sie an die Verdopplung der Reiskörner auf dem Schachbrett, es fängt harmlos an, aber Verdopplung nimmt ab einer gewissen Zahl von Reiskörnern doch dramatische Züge an). Passen unsere heutigen Rechenzentren in 20 Jahren wieder in einen einzigen 19-Zoll-Schrank? Und was bedeutet das für die Übergangszeit, sagen wir für die nächsten 5 Jahre? Nun interpolieren Sie einmal die Leistung der ersten verfügbaren 11nm-Technik in einen heutigen Schrank und nehmen wir mal großzügig an, dass Intel diese Produktionstechnik in 3 Jahren sowohl im Griff hat als auch für Server-Chips benutzt (das ist vom Timing her durchaus unklar, da Intel mit 22nm und 11nm erst einmal andere Zielmärkte im Visier hat). Was bedeutet das für diesen Schrank sagen wir im Vergleich zu unserer Serverlandschaft aus dem Jahr 2000? Nun, wenn ein größeres Unternehmen im Jahr 2000 zwischen 1000 und 5000 physikalischer Server hat-



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

Neue Lösungen verändern den Storage-Markt

te, dann passen die mit der neuen Technik ohne Probleme in einen Schrank (Vorsicht, die Festplatten sind ins SAN ausgelegt und haben immer schon einen erheblichen Teil des benötigten Platzes eingenommen). Was bedeutet dann der Begriff „Virtuelle Infrastruktur“ eigentlich noch? Wandern dann virtuelle Maschinen überhaupt noch außerhalb des Schrankes?

Parallel entwickelt sich SAN-Technik geradezu rasant weiter. Die neue 10000 Euro-Einstiegsklasse erreicht ein Leistungsniveau, für das man noch vor wenigen Jahren den 10-fachen Betrag und mehr gezahlt hätte. Der Wettbewerb konzentriert sich auch immer mehr auf diesen Einstiegsbereich. Anders formuliert: SAN-Technologie wird auch für kleine und mittlere Unternehmen die Normalität werden. Und SAN-Technologie ist schon lange nicht mehr auf große Datenbank-Anwendungen begrenzt. Sie macht auch für normale File-Server-Anwendungen durchaus Sinn. Was bedeutet das? Nun, wir konzentrieren Speicher immer stärker. Parallel erhöht SSD-Technik den Multiusergrad auch der Einstiegs-Klasse. Kombiniert mit Caching und immer größeren Caches (siehe 11nm-Technik) kommen wir in Leistungsklassen, die bisher undenkbar waren. Auch auf der Server-Seite können SSDs als Caches in Zukunft eine Rolle spielen. Dies ergibt ein sehr gemischtes Bild für die notwendige Infrastruktur zwischen Servern und SANs. Die Zahl der SANs steigt, ihre Leistung steigt, also auch ihre Anbindungs-Leistung, aber eventuell sinken Anforderungen auf der Server-Seite. Dafür generiert ein Multi-Tiering zwischen SANs eventuell wieder neue Anforderungen. Zumindest, wenn nicht FC zum Einsatz kommt.

Damit sind wir bei der leidigen FCoE und FC-Diskussion, die durch die Einführung des 16 Gbit FC natürlich unvermeidlich neuen Nährboden erhalten hat. Nun, zum einen ist die Diskussion zunehmend unwichtig, da auf der SAN-Seite der klare Trend zur Unterstützung beider Verfahren existiert. Zum anderen muss aus meiner Sicht eher die Frage gestellt werden, welche Zukunftsbedeutung iSCSI und (p)NFS haben werden. Die Masse unseres Speichers ist File-Speicher und nicht Block-Speicher. Muss die iSCSI und NFS-Diskussion nicht neu aufgelegt werden?

Damit sind wir bei den Netzwerk-Infrastrukturen. Die Teilnehmer mehrerer ComConsult-Kongresse wurden in den letzten 2 Jahren mit der Diskussion der neuen Standards erbarmungslos gequält (geht leider nicht anders, die Tücken liegen in den Details der neuen Standards, sorry :-)). Die Bedarfslage ist weiterhin kon-

fus und die Befürworter und Gegner stehen sich wie immer in der Geschichte des LAN erbittert gegenüber.

Das Kernargument der Gegner der neuen Verfahren ist klar: wir sollten wie immer auf den ganzen Schnickschnack verzichten (das Gelump abschalten wie Walter Röhrl zu sagen pflegte, wenn er über die Entwicklung im Automobilbereich sprach) und einfach genügend Kapazität einplanen. Sprich Multi-10-Gigabit Netzwerke in Parallelschaltung mit TRILL oder SPB bieten so viel Kapazität, dass man eigentlich gar nicht mehr braucht. Na ja, zumindest bräuchte man dann ja schon mal einen der neuen Standards, auch aus der Sicht der Gegner der neuen Standards.

Die Befürworter der neuen Standards sehen naturgemäß einen hohen Bedarf. Ihre Argumente sind:

- FCoE geht nicht ohne DCB mit allen darin enthaltenen Variationen
- Gerade die Einbindung von Speicher-Systemen generiert so hohe Lastspitzen und Summenlasten, dass wir den Verkehr besser als bisher steuern müssen
- Die Konzentration von Servern in immer weniger Schränken erhöht die Punktlasten und führt zu stärkeren Schwankungen als bisher
- Gerade im RZ müssen die Netzwerke Virtualisierungs-bewusst sein, dies geht mit den „alten“ Standards gar nicht. An einer Mindestmenge neuer Standards zumindest zur Integration virtueller Systeme kommen wir gar nicht vorbei
- Latenz wird in Zukunft ein wichtiger Planungsparameter, da immer mehr Applikationen ins RZ kommen, die von sehr kleinen Latenzwerten profitieren (auch außerhalb des Finanzmarktes?), wir brauchen also Verfahren, die uns Maxi-

mal-Latenzwerte garantieren

Damit noch nicht genug, liegt auch im Virtualisierungsmarkt noch weiteres Potenzial für Ärger, zumindest aus der Sicht der Infrastrukturbetreiber. Speziell High Availability und Disaster Recovery sind Themen, die immer wieder Diskussionen über Bandbreiten in Kombination mit ausgedehnten Layer-2-Strukturen aufkommen lassen. Da auch die Virtualisierungshersteller sich weiter entwickeln müssen, ist gerade dieser Markt eine sehr wahrscheinliche Richtung der Entwicklung. Geht HA über den bisherigen einen Kern hinaus, kommen schnell völlig neue Anforderungen ins Spiel.

Betrachten wir das Gesamtbild der Situation im Rechenzentrum, dann ist das eine sehr komplexe Lage. Das Kernproblem ist, dass die Entscheidung nicht in den Einzeltechnologie-Bereichen fallen kann. Wo immer der Weg hingeht, er muss als Gesamtweg im Sinne eines integrierten Applikation-Server-SAN-Netzwerk-Weges gegangen werden. In dieser Gesamtsicht müssen die kritischen Elemente in den einzelnen Technologie-Bereichen identifiziert werden und mögliche Sackgassen aufgezeigt werden.

In diesem Sinne ist die Antwort auf die in der Überschrift gestellte Frage klar: wir brauchen ein Redesign, aber nicht in Form eines Mischmasches aus technischen Einzelprojekten. Wir brauchen ein Redesign der Gesamt-Architektur, die wir in Zukunft im Rechenzentrum umsetzen wollen.

Hier setzen wir mit unserem diesjährigen Rechenzentrums-Infrastruktur-Redesign-Forum im November an. Und die vorangegangenen Ausführungen zeigen, dass dies ein spannendes Forum wird.

Ihr
Dr. Jürgen Suppan

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.

Moderatoren: Dipl.-Inform. Matthias Egerland, Dr. Franz-Joachim Kauffels

Preis: € 2.190,-* zzgl. MwSt. bzw. € 1.790,-* zzgl. MwSt. - *gültig bis zum 31.08.11



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

Rechenzentrum Infrastruktur-Redesign Forum 2011

07. - 10.11.11 in Königswinter

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

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

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

Die Ausgangssituation ist durch folgende Stichworte zu kennzeichnen:

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

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

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



Kommunikation im RZ von drei neuen Verkehrsströmen geprägt wird:

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

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

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

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

dungen erfolgreich sein soll.

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

Das zementiert eine unangenehme Tatsache:

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

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

Was passiert in den nächsten Jahren? Knapp zusammengefasst:

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

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

- Fat Tree-Strukturen für Unified Fabrics
- Virtual Chassis

- ScaleOut
- Single-Hop-Netze

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

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

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

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

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

stäbe anlegen?

Also kann man sich eigentlich auf wenige Kernfragen konzentrieren:

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

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

Frühbucherphase bis zum 31.08.2011

Fax-Antwort an ComConsult 02408/955-399

Anmeldung Rechenzentrum Infrastruktur- Redesign Forum 2011

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

mit Workshop am letzten Tag

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

ohne Workshop am letzten Tag

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

*gültig bis 31.08.2011

☐ Bitte reservieren Sie mir ein Zimmer

vom _____ bis _____ 11

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift



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

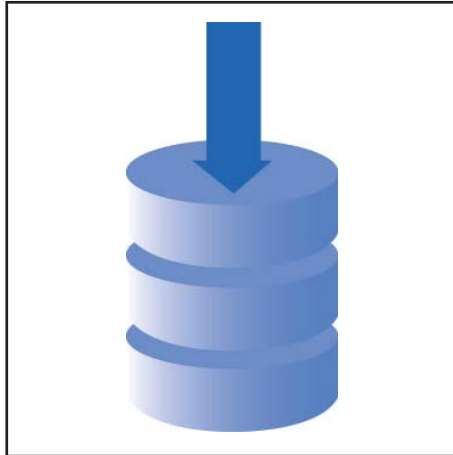
ComConsult Storage-Forum 2011 13.07. - 15.07.11 in Bonn

Die ComConsult Akademie veranstaltet vom 13.07. - 15.07.11 ihren neuen Kongress „ComConsult Storage-Forum 2011“ in Bonn.

Server- und Speicherkonsolidierung führen zu einem Bedarf nach einer neuen Kategorie von Speicher. War früher der Markt zweigeteilt in den Server-Speicher auf der einen Seite (DAS) und das SAN auf der anderen Seite, so ist diese Struktur nicht weiter aufrecht zu erhalten. Direct Attached Storage wird im Rahmen der Konsolidierung durch zentralen Speicher abgelöst, doch das traditionelle SAN kann die Rolle des neuen zentralen Speichers nicht übernehmen. Es ist zu teuer und zu komplex.

In Konsequenz ist der Bedarf nach einer neuen Produktgruppe entstanden mit den Kriterien:

- Preiswert: die Kosten der bisherigen DAS- und SAN-Lösungen müssen deutlich unterboten werden. SAN-Qualität zum Preis von weniger als DAS aber mit der Performance von DAS ist gefragt
- Einfach zu handhaben: Automatisierung der Konfiguration und Zuordnung im Sinne eines Self-Provisioning
- Flexibel: schnelle Reaktion auf einen weiter wachsenden Bedarf nach mehr Speicher
- Intelligente Performance: zentraler Speicher darf den Server und die Applikation nicht bremsen, angesichts der Entwicklung bei Servern und Web-Applikationen ist das eine Herausforderung. Der Bedarf ist klar: mehr Benutzer, mehr Transaktionen, mehr Durchsatz. Aber intelligent: Leistung nur da, wo sie auch benötigt wird
- Gestaltbar: leicht anpassbar an unterschiedliche Lokationen wie Remote Offices und verteilte Rechenzentren sowie gestaltbar in der Form der Nutzung
- Universell: Integration von Block- und File-Zugriff, Angebot aller gängigen Schnittstellen vom Fibre Channel über FCoE bis iSCSI



Die bisherigen Lösungen versagen angesichts dieser Anforderungen:

- Der Betriebsaufwand und die Kosten der Kombination traditioneller SANs mit virtuellen Umgebungen sind deutlich zu hoch
- Einstiegslösungen aus dem unteren iSCSI-Markt erfüllen viele der Anforderungen nicht, zum Beispiel fehlen Thin Provisioning, Deduplizierung, Kompression und Integration in virtuelle Umgebungen
- Viele der bisherigen Lösungen erfüllen die Leistungs-Anforderungen der Zukunft nicht

In Konsequenz entsteht momentan ein völlig neuer Typ von Speicher-Lösungen. Dieser stellt gleichzeitig die Wirtschaftlichkeit der traditionellen SANs in Frage. Diese sehen in der Tat gegenüber den neuen Lösungen „alt“ aus. Das ganze explosive Gemisch wird gewürzt mit neuen Technologien wie SSDs und Cloud Storage. Heraus kommt eine Infrastruktur, die es so bisher nicht gab:

- Hochgradig wirtschaftlich
- Sehr leistungsstark mit extrem hoher Performance
- Extrem flexibel und gestaltbar
- Für alle Unternehmensgrößen nutzbar

Diese Entwicklung hat ihren Preis auf der Seite der Infrastrukturen. Mehr und mehr Speicherverkehr wird über Ethernet transportiert, auch wenn der Fibre Channel-

Zugang weiterhin stark genutzt wird. Wir sind in einem Übergang hin zu 10 Giga-bit und mehr pro Port, sowohl für FC als auch für Ethernet. Und trotzdem stellt sich die Frage, ob die vorhandene Infrastruktur wirklich auf Dauer skaliert.

Hier setzt das ComConsult Storage Forum 2011 an. Wir analysieren mit hochkarätigen Experten:

- Was passiert momentan im Speichermarkt?
- Wohin geht die Entwicklung?
- Wie sieht der Zukunfts-Bedarf aus?
- Welche der neuen Technologien sind kritisch?
- Welche Produkte werden in Zukunft eine große Rolle spielen?
- Was bedeutet das für die Infrastrukturen?
- Welche Handlungsoptionen bestehen für die Unternehmen?

Durch das Forum führen Sie:

Dipl.-Inform. Matthias Egerland hat an der RWTH Aachen Informatik studiert und ist seit 2005 Mitarbeiter der ComConsult Beratung und Planung GmbH. Er ist Leiter des Competence Center Data Center und unterstützt die Competence Center IT-Sicherheit und Netze. Neben den Schwerpunkten Desktop-, Server- und Infrastruktur-Virtualisierung beschäftigt sich Herr Egerland insbesondere mit Speicherlösungen in virtualisierten Umgebungen. In Projekten erstellt er Konzepte und Ausschreibungen von IT-Infrastruktur-Lösungen gemäß UfAB in den Bereichen Lokale Netze (LAN) für mehrere tausend Teilnehmer, aktive Rechenzentrumskomponenten - insbesondere Server, Speicher, Netzwerk und Firewalls - sowie Virtuelle Informationstechnologie.

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.

Gratis bei Buchung dieses Kongresses

Als Teilnehmer dieses Kongresses erhalten Sie vor Ort ein kostenloses Exemplar des neuen Reports "Public und Private Clouds in der Analyse" von Dr. Jürgen Suppan.

Programmübersicht - ComConsult Storage-Forum 2011

Mittwoch, den 13.07.2011

9:30 bis 11:00 Uhr

Keynote: Vom traditionellen RZ zur Private Cloud

- Cloud Computing: Funktion, Chancen, Risiken
- Entwicklung der Web-Architekturen
- Server-, Speicher- und Netzwerkvirtualisierung
- Entwicklungen bei Netzwerken (Leistung, Latenz, Scale-Out)
- Strukturwandel des RZs
- Hochkonzentrierte Blade-Server: der Host-Phönix?

Dr. Franz-Joachim Kauffels, Unternehmensberater

11:00 bis 11:30 Uhr - Kaffeepause

11:30 bis 12:30 Uhr

Erfahrungen von Unternehmen mit ihrer Speicherinfrastruktur

- Datenwachstum: wie schnell wachsen die Datenbestände?
- Wie ist dem Problem des Datenwachstums beizukommen: mit organisatorischen oder technischen Mitteln?
- Was wird mehr benötigt: Volumen oder Performance?
- Welche neuen Verfahren können genutzt werden, um dem Bedarf an mehr Volumen und Performance gerecht zu werden?
- Abgrenzung zwischen verschiedenen Datenklassen und Speichertypen
- Hochverfügbarkeit der Speicherinfrastruktur
- Herausforderung Datensicherung
- Sinn und Zweck der Speichervirtualisierung
- Wie viel Speichervirtualisierung braucht ein Unternehmen?

Dr. Behrooz Moayeri, ComConsult Beratung und Planung GmbH

12:30 bis 14:00 Uhr - Mittagspause

14:00 bis 15:00 Uhr

iSCSI im Rechenzentrum - ein Erfahrungsbericht

- Wie sieht die Grobstruktur des betrachteten RZs aus?
- Welche Anforderungen bestehen seitens der Server beim Zugriff auf den Speicher?
- Welche Aspekte haben die Entscheidung in Richtung iSCSI begünstigt?
- Welche (negativen) Erfahrungen wurden seitens Vattenfall gemacht,

als das iSCSI-Netz noch als VLAN über eine gemeinsame Infrastruktur geführt wurde?

- Welche Verbesserungen haben sich ergeben durch die Separierung des iSCSI-Netzes auf dedizierte Hardware?
- Welche Leistungswerte (Durchsatz, I/Os, Latenzen) liefert das aktuelle iSCSI-Netz?
- Welche Gründe sprechen aus Sicht von Vattenfall für einen Wechsel auf NFS?

Michael Witschke, Vattenfall IT

15:00 bis 15:30 Uhr - Kaffeepause

15:30 bis 16:30 Uhr

Dateibasierter Speicherzugriff -

Erfahrungen aus aktuellen Speicherprojekten

- Wofür eignen sich dateibasierte Verfahren wie NFS, CIFS?
- Was ist neu an NFSv4 und wie wirkt sich das auf die Nutzbarkeit dieses Protokolls als Ersatz von blockbasiertem Speicherzugriff aus?
- Welche Vorteile ergeben sich im Umfeld von NFS bei der Nutzung für Datenbanken oder virtuelle Umgebungen?
- Welche Einschränkungen ergeben sich hierbei?
- In der Praxis erzielbare Leistungswerte: Durchsatz, IOPS, Latenzen
- Wie skaliert eine NFS-Umgebung für Datenbanken und virtuelle Maschinen?
- Gibt es eine Größenordnung (Anzahl Hosts / Anzahl VMs), ab der eine NFS-Umgebung nicht mehr effizient ist?
- Wie gestalten sich Disaster Recovery Szenarien mit NFS?

Michael Albers, Hellweg Data

16:30 bis 17:30 Uhr

Technologische Entwicklungen bei Speichersystemen

- Dauerhaftigkeit von Back-End-Systemen
- Masselose SSDs als neue Speicherkategorie
- Weiterentwicklung 3D-SSDs
- Das Netz als Systembus der Virtualisierten Umgebung
- Enterprise Storage Management

Dr. Franz-Joachim Kauffels, Unternehmensberater

Happy Hour ab 18:00 Uhr

Donnerstag, den 14.07.2011

9:00 bis 10:00 Uhr

Konstruktive Sicherheitsaspekte virtualisierter Storage Networks

- Sicherheit in SAN und NAS
- Welche Gefährdungen sind in SAN und NAS relevant?
- Maßnahmen zur Absicherung von SAN und NAS - Elemente einer SAN-Sicherheitsrichtlinie
- Konzepte für eine Zonierung im SAN
- Wann sollte verschlüsselt werden?

Dr. Simon Hoff, ComConsult Beratung und Planung GmbH

10:00 bis 11:00 Uhr

Speichervirtualisierung: Steigerung der Flexibilität und Funktionalität oder nur der Komplexität?

- Konzept der Speichervirtualisierung
- Implementierungsvarianten: Inband, Split-Path, Netzwerk-basiert
- Leistungsmerkmale aktueller Lösungen
- Nutzung zur Datenspiegelung zwischen entfernten Standorten

Dr. Behrooz Moayeri, ComConsult Beratung und Planung GmbH

11:00 bis 11:30 Uhr - Kaffeepause

11:30 bis 12:30 Uhr

Vereinfachtes Datenmanagement und Kostenreduktion durch Dateivirtualisierung

- Die Herausforderung: exponentielles Wachstum unstrukturierter Daten in traditionellen NAS-Infrastrukturen (NFS, CIFS)
- Die Lösung: Dynamische NAS-Speicherinfrastrukturen durch Dateivirtualisierung • Funktionsprinzip der Dateivirtualisierung
- Unterbrechungsfreie Datenmigrationen
- Automatische Zuordnung von Daten zu Speicherklassen
- Optimierung von Backups und Kapazitätsverteilung

Dr. Shahriar Daneshjoo, F5 Networks GmbH

12:30 bis 14:00 Uhr - Mittagspause

14:00 bis 15:00 Uhr

Moderne Datensicherung mit VTLs und Disk-2-Disk-Verfahren

- Funktionsweise einer VTL
- Wann ist eine VTL einem allgemeinen Disk-2-Disk-Datensicherungsverfahren vorzuziehen? • Wie kann eine VTL dabei helfen, aktuelle SLAs gerade im Bereich der Wiederherstellungszeiten einzuhalten?
- Einsatz von Snapshots im Umfeld der Datensicherung
- Datensicherung in virtuellen Umgebungen: LAN-basiert oder LAN-free?
- Möglichkeiten und Grenzen Image- und Datei-basierter Datensicherungen in virtuellen Umgebungen
- Datensicherung mit Bordmitteln der Virtualisierungshersteller und aktuellen Produkten von Drittherstellern

Dipl.-Inform. Matthias Egerland,

ComConsult Beratung und Planung GmbH

15:00 bis 15:30 Uhr - Kaffeepause

15:30 bis 16:15 Uhr

SAN Performance Monitoring & Troubleshooting

- Erfahrungen eines Troubleshooters - Häufigste Probleme im SAN
- Fehlerfindung an Hand eines Kundenbeispiels
- Wie lassen sich Probleme darstellen - ein Monitoringkonzept
- Welche Herausforderungen erwarten uns in Zukunft - Technologieausblick

Sebastian Schröder, MEN@NET GmbH

16:15 bis 17:00 Uhr

Automatisierte Speicherinstrumente und mobile Daten für wirtschaftliche Unternehmenslösungen

- Mit universellen Speicherinstrumenten durch „dick und dünn“
- Ökonomische Methoden, die Ressourcen und Budgets gleichermaßen schonen und optimieren
- Mittels Virtualisierung die richtigen Informationen, zur richtigen Zeit, am richtigen Ort
- Lösung von globalen Anforderungen mit flexiblen Instrumenten

Dr. Georgios Rimikis, HITACHI DATA SYSTEMS GmbH

Programmübersicht - ComConsult Storage-Forum 2011

Freitag, den 15.07.2011

9:00 bis 15:00 Uhr

Strategien und Produkte der führenden Hersteller

- Wie sieht die Architektur der Speicherlösung aus?
- Wie skaliert die Leistung des Speichersystems?
- Welche Hochverfügbarkeitsmechanismen besitzt das Speichersystem?
- Wie können Multi-Site Disaster Recovery Szenarien mit dem Speichersystem gelöst werden?
- Welche Vorteile bringt das Speichersystem insbesondere in virtualisierten Umgebungen?

9:00 bis 09:45 Uhr

EMC Symmetrix vMax

Carsten Haak, EMC Deutschland GmbH

9:45 bis 10:30 Uhr

Fujitsu Eternus

Wolfgang Schenk, Fujitsu Technology Solutions GmbH

10:30 bis 11:00 Uhr - Kaffeepause

11:00 bis 11:45 Uhr

HP Lefthand

Florian Bettges, Hewlett-Packard Deutschland GmbH

11:45 bis 12:30 Uhr

IBM XIV

Dirk Vogelsang, IBM Deutschland GmbH

12:30 bis 13:30 Uhr - Mittagspause

13:30 bis 14:15 Uhr

NetApp FAS/V-Serie

Roger Zink, NetApp Deutschland GmbH

Anschließende Diskussion mit den Herstellern:

- Wie sieht die Speicherarchitektur der Zukunft aus?
- Welches Speicherzugriffsverfahren wird sich durchsetzen?
- Gehört die Zukunft den SSDs?

15:00 Uhr Ende der Veranstaltung



Inklusive kostenlosem Report "Public und Private Clouds" bei Kongressteilnahme

Sie erhalten den im März 2011 erschienenen Report bei Teilnahme an diesem Forum gratis als Bestandteil der Kongressunterlagen.

Fax-Antwort an ComConsult 02408/955-399

Anmeldung ComConsult Storage-Forum 2011

Ich buche den Kongress
ComConsult Storage-Forum 2011

☐ 13.07. - 15.07.11 in Bonn
zum Preis von € 1.990,-- zzgl. MwSt.

☐ Bitte reservieren Sie mir ein Zimmer

vom _____ bis _____ 11

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort



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

eMail

Unterschrift

Aktuelles Seminar

Sommerschule 2011 - Intensiv-Update auf den neuesten Stand der Netzwerktechnik

18.07. - 22.07.11 in Aachen

Die ComConsult Akademie veranstaltet vom 18.07. - 22.07.11 ihr Intensiv-Seminar „Sommerschule 2011“ in Aachen.

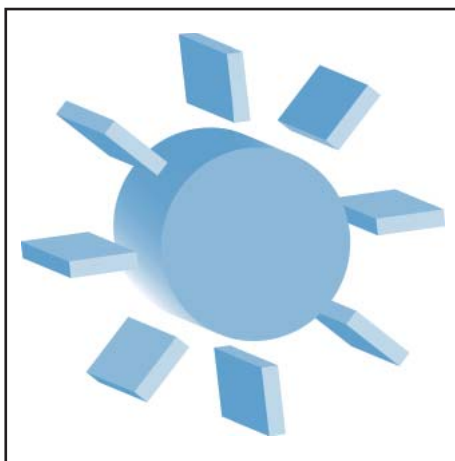
Die Sommerschule gibt innerhalb von 5 Tagen den kompakten und intensiven Überblick über die neusten Entwicklungen im Umfeld der Netzwerk-Technologien: Anforderungen an zukunftssichere Netzwerke: was ändert sich; Neue Technologien und Standards; Design-Verfahren im Vergleich; Ausgewählte Technologien in der Analyse; Umfeld-Analyse: was passiert um Netzwerke herum.

Die Sommerschule wendet sich an Teilnehmer mit bestehenden Grundlagen-Kenntnissen und ist als Weiterbildung für berufserfahrene Netzwerker konzipiert.

Die 5 Intensivtage der Sommerschule 2011 teilen sich wie folgt auf:

Switching 2011: Stand der Technik und Prognosen, Dipl. Inform. Petra Borowka

- Aktuelle Herausforderungen
- Bandbreiten und Architekturen
- Neue Verfahren und ihre Auswirkung auf das Design
- Analyse: TRILL kontra SPB, wo liegt die Zukunft



IPv6: von den Grundlagen zur Migration, Markus Schaub

- IPv6: was bedeutet das?
- IPv6: welche Alternativen der Umsetzung gibt es?
- Wie sieht der Übergang aus?
- Projekterfahrungen und typische Probleme
- Zeitaufwand und Komplexität der Migration

Rechenzentrum: Server und Speicher im Netzwerk

- Herausforderung: Technologien wach-

sen zusammen

- Analyse: wesentliche Entwicklungen und Auswirkungen auf Infrastrukturen

Infrastrukturen und Trouble Shooting

- Verkabelung 2011: aktuelle Entwicklungen, Alternativen und Zukunftsprognosen, Dipl.-Ing. Hartmut Kell

- Trouble Shooting 2011: mehr Bandbreite, stabilere Netzwerke, höhere Protokolle - Herausforderungen und Strategien für eine professionelle Fehlersuche und Beseitigung, Dr. Jochen Wetzlar

Sicherheit 2011, Dr. Simon Hoff

- Aktuelle Herausforderungen
- Öffnung zur Cloud, höherer Schutzbedarf: wie passt das zusammen
- Mobile Teilnehmer: vom Alptraum zur Lösungs-Strategie

Die Sommerschule 2011 greift die aktuellsten Entwicklungen auf, analysiert die Auswirkungen auf Ihre Netzwerke und diskutiert Entscheidungs- und Handlungs-Alternativen. Zögern Sie nicht, sich einen Platz in dieser herausragenden Veranstaltung zu sichern. Die Plätze sind limitiert.

Fax-Antwort an ComConsult 02408/955-399

Anmeldung

Sommerschule 2011

Ich buche das Seminar

Sommerschule 2011 - Intensiv-Update auf den letzten Stand der Netzwerktechnik

☐ **18.07. - 22.07.11 in Aachen**
zum Preis von € 2.490,-- zzgl. MwSt.

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

Vorname

Nachname

Firma

Telefon/Fax

Straße

PLZ, Ort

eMail

Unterschrift

Programmübersicht - ComConsult Sommerschule 2011

Montag, den 18.07.2011 - Switching 2011: Status, Trendanalyse und Prognosen

9:30 - 11:00 Uhr

Was sich im Netz verändert

- IT-Trends und ihre Auswirkungen auf Netze
- Wo werden 40/100 Gigabit benötigt: RZ, Campus, WAN?
- Was passiert im Access-Bereich?
- Auswirkung von Virtualisierung auf Netze
- Konvergenz: Daten, Speicher, Voice, Video in einem Netzwerk?
- Mandantenfähige Netzwerke und warum sie benötigt werden
- IPv6: noch aufzuhalten?

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

11:00 - 11:15 Uhr Kaffeepause

12:30 - 13:30 Uhr Mittagspause

15:00 - 15:15 Uhr Kaffeepause

11:15 - 17:30 Uhr

Switching 2011: Status, Trendanalyse und Prognosen

Layer-2 versus Layer-3: Shortest Path Bridging und was es für das Netzdesign bedeutet

- Nachteile des Spanning Tree
- Spanning Tree Alternative: "Virtuelles Chassis"
- Spanning Tree Alternative: Layer-3
- Was ist Shortest Path Bridging?
- TRILL versus IEEE SPBM und SPBV
- Einsparpotenziale mit Shortest Path Bridging
- Herstellerpositionierung
- Fazit und Empfehlungen zu Shortest Path Bridging

Neue Layer-2 Standards und was sie für Netzdesign und -betrieb bedeuten

- IEEE Energy Efficient Ethernet: Strom und Abwärme sparen mit Twisted Pair für 1 Gbit und 10 Gbit Ethernet

- Data Center Bridging: Anpassung von Ethernet an Data Center Bedarfe
 - Congestion Notification
 - Priority Based Flow Control
 - Enhanced Transmission Selection
 - DCB Management (DCBXP)
 - Audio Video Bridging
 - Echtzeit-Anforderungen von High End Anwendungen im LAN
 - Zeitsynchronisierung
 - Forwarding und Queueing für Echtzeit-Anwendungen (FQTSS)
 - Ressourcen-Reservierung für Echtzeit-Anwendungen (SRP)
 - Herstellerpositionierung
 - Fazit und Empfehlungen zu den neuen Standards
- Dipl.-Inform. Petra Borowka-Gatzweiler,
Unternehmensberatung Netzwerke UBN*

Happy Hour ab 19:00 Uhr

Dienstag, den 19.07.2011 - IPv6

9:00 - 17:00 Uhr

IPv6: von den Grundlagen zur Migration

- Was ändert sich und was bleibt bei IPv6?
- IPv6 Adresskonzept:
 - Adresstypen • Adressaufbau
 - Adressverteilung
- Wie kommt die Konfiguration auf das Endgerät?
 - Auto-Konfiguration
 - DHCPv6: stateful und stateless
 - Vor- und Nachteile der Verfahren

- Neue Aufgaben für den Router:
 - Neighbor Discovery • Forwarding
 - Neues bei OSPFv3 • VRRPv3
- Migrationsverfahren:
 - Tunnelvarianten: ISATP, 6to4, 6in4, 6rd, Teredo
 - Dual-Stack, DS-Lite • NAT
 - Wie sind die Verfahren zu bewerten?
 - Welche Verfahren kommen wann zum Einsatz?
- Kritische Aspekte bei IPv6
 - Sicherheit und IPv6

- Stand der Implementierung
 - Anforderung an Netzwerkkomponenten, Endgeräte und Software
 - Wie und wann migriert man zu IPv6?
- Was gibt es zu beachten?
Markus Schaub, ComConsult-Study.tv*

11:00 - 11:15 Uhr Kaffeepause

12:30 - 13:30 Uhr Mittagspause

15:00 - 15:15 Uhr Kaffeepause

Mittwoch, den 20.07.2011 - Rechenzentrum: Server und Speicher im Netzwerk

9:00 - 17:00 Uhr

Neue Struktur von RZ-Netzen

- Mehrstufige Netzhierarchien in Rechenzentren?
- Routing und Switching im RZ-Netz
- Redundanz, aber wie?

Virtualisierung: was bedeuten wandernde VMs für unsere Ressourcen?

I: Hochverfügbarkeit, Fehlertoleranz, Disaster Recovery

- Welche Verfügbarkeit kann erzielt werden?
- Wie sehen Standort-Redundanzen aus?
- Welche Auswirkungen haben über weite Entfernungen wandernde VMs?

II: Netzanbindung virtueller Server

- Server-Netze in virtuellen Infrastrukturen
- Was leisten neue virtuelle Switches?
- Welche Technologien werden zukünftig zur Netzanbindung virtueller Server genutzt werden?

Speicherlösungen

- Speichervirtualisierung
- ThinProvisioning
- Implementierungsvarianten

Backup & Restore, Disaster Recovery am Beispiel von SAP-Umgebungen

- Block-level Incremental Backup

- Beispielrechnung für eine Datensicherungsinfrastruktur
- Ansätze zur Standort-übergreifenden Datensicherung
- Backup-Szenarien mittels virtuellen Tape Libraries (VTLs)

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

11:00 - 11:15 Uhr Kaffeepause

12:30 - 13:30 Uhr Mittagspause

15:00 - 15:15 Uhr Kaffeepause

Donnerstag, den 21.07.2011 - Infrastrukturen und Trouble Shooting

9:00 - 12:30 Uhr

Rechenzentrum: Physikalische Anforderungen der Gigabit-Übertragungsraten

- Kupfer versus Glasfaser
- Was fordert die IEEE
- Längenrestriktionen verursacht durch Modendispersion und „knappe“ Dämpfungsbudgets
- Kupfer, alles ganz einfach und doch nicht geeignet?
- Die Schwierigkeit von Glasfaserdurchdringungen

Kommunikationsverkabelung im Rechenzentrum

- Vorteile der Anwendung der EN 50173-5
- Strukturvarianten im RZ und deren Einfluss auf

- die Kommunikationsverkabelung
 - Güteklassen bei Kupfer und Glasfaser
 - Sinnvolle Kombinationen von Stecker und Kabel
- Dipl.-Ing. Hartmut Kell,
ComConsult Beratung und Planung GmbH*

12:30 - 13:30 Uhr Mittagspause

13:30 Uhr - 17:00 Uhr

Trouble Shooting 2011

- Von Freeware bis zum Millionengrab - Aktuelle Werkzeuge für das Trouble Shooting
- Masse statt Klasse? Langzeitaufzeichnung als Kö-

- nigdisziplin der Netzanalyse
- Rechtliche Aspekte: Darf man präventiv alle Pakete aufzeichnen?
- Erkenntnisse aus der Performance-Analyse zahlreicher Anwendungen
- IPv6 als neue Herausforderung beim Trouble Shooting
- Und umgekehrt: Protokoll-Know-How als Voraussetzung für erfolgreiche IPv6-Migration

*Dr. Joachim Wetzlar,
ComConsult Beratung und Planung GmbH*

15:00 - 15:15 Uhr Kaffeepause

Freitag, den 22.07.2011 - IT-Sicherheit

9:00 - 15:30 Uhr

Sicherheit 2011

- Aktuelle Herausforderungen
- Informationssicherheit trotz maximaler Konsolidierung bei Virtualisierung
- Server-based Computing und Desktop-Virtualisierung
- Öffnung zur Cloud, höherer Schutzbedarf: wie passt das zusammen

- Integration von Smartphones und Tablets: vom Alptraum zur Lösungs-Strategie
- Consumerization of IT: Fremdgeräte, Consumer-Geräte, Consumer Software und soziale Netze im Unternehmen
- Schutz vor Datendiebstahl und -verlust: Content Security, Endpoint Security, Data Leakage Prevention und sichere Netze
- Seiteneffekte der Konvergenz: Standard-IT-Kompo-

- nenten in Anlagen, Steuerungen und Maschinen
- Notwendigkeit werkzeuggestützter Prozesse in der Informationssicherheit

*Dr. Simon Hoff,
ComConsult Beratung und Planung GmbH*

11:00 - 11:15 Uhr Kaffeepause

12:30 - 13:30 Uhr Mittagspause

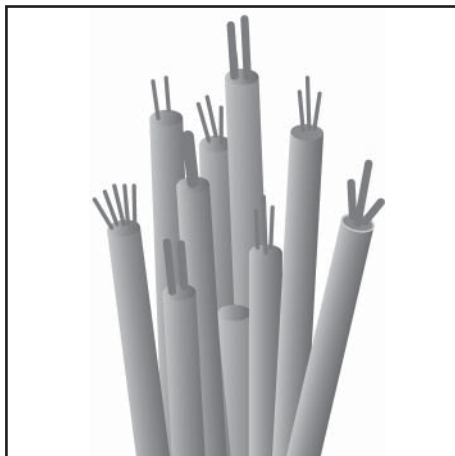
15:30 Uhr Ende der Veranstaltung

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.

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



wählten herstellerspezifischen Lösungen wird auch auf Planungs- und installationsbegleitende Maßnahmen eingegangen, die im Rahmen einer anstehenden Verkabelung zu beachten sind.

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 Dämpfungproblematik 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

Aktuelle Neuerscheinungen bei ComConsult-Study.tv

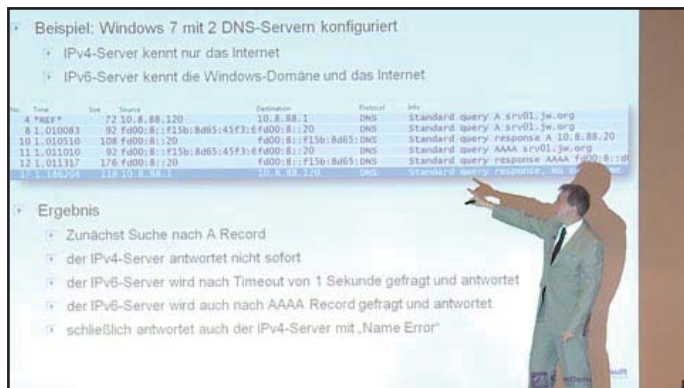
Themenbereich: Betrieb und Architekturen

Trouble Shooting IPv6

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

Zeit: 00:38:54

Preis: Kostenlos mit Abo



Als besonderes Highlight für die Abonnenten von ComConsult-Study.tv den best bewerteten Vortrag des ComConsult IPv6-Forums 2011.

Themenbereich: Analyse und Strategie

Seminar: Cloud Storage in der Analyse

Referent: **Dr. Jürgen Suppan**

Zeit: 00:56:13

Preis: Kostenlos mit Abo



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

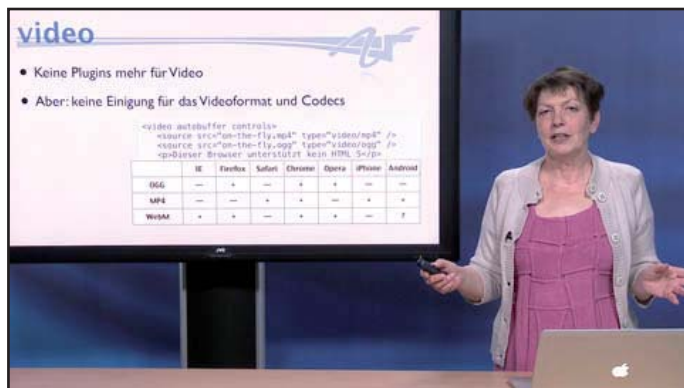
Themenbereich: Web

HTML 5

Referent: **Dipl.-Inform. Ulrike Häßler**

Zeit: 00:31:54

Preis: Kostenlos



Auf vielfachen Wunsch der dritte Teil unseres erfolgreichen Seminars über Webtechnologien hier als kostenfreie Version. HTML5 legt die technische Basis für eine völlig neue Art von Applikationen im Web. Die bisher extrem störenden Unterschiede zwischen den Browsern sollen verschwinden, Plugins sollen überflüssig werden. Die Vision ist eine völlig neue Generation von Client-Server-Technologie basierend auf Web-Clients. Lernen Sie welche technischen Elemente HTML5 liefert, um diese Vision Realität werden zu lassen.

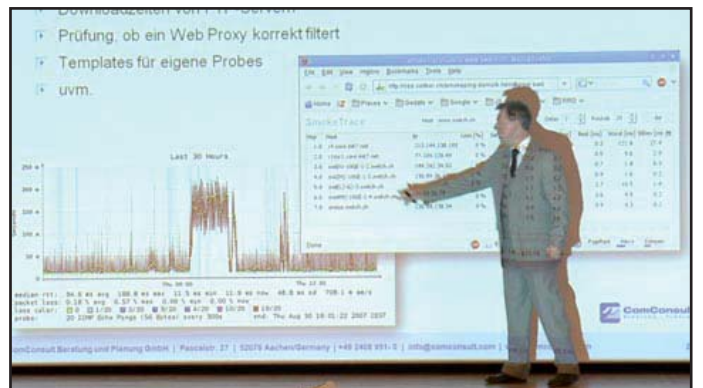
Themenbereich: Netzwerke

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

Singlemode und Multimode in der Backbone- Verkabelung

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.

Der nachfolgende Artikel zieht neben den Empfehlungen der anwendungsneutralen Kabelnormen auch die spezifischen Anforderungen der wichtigsten Ethernet-Zugangungsverfahren heran und leitet daraus Planungsempfehlungen ab für das Design von modernen Primär- und Sekundärverkabelungen.

Elemente einer Backbone-Verkabelung

Am Anfang des Artikels wird zunächst eine Begriffsbestimmung durchgeführt, diese vereinfacht das Verständnis der Erläuterungen. Da in der Normierung DIN/EN 50173 eine Begriffsbestimmung der Grundelemente sehr anschaulich durchgeführt wird, erscheint eine Übernahme der dortigen Definition von Teilelementen der Kommunikationsverkabelung und damit der Backbone-Verkabelung sinnvoll. Eine Verkabelungsanlage, die zur Informationsübertragung genutzt wird, besteht gemäß DIN aus Verteilern und einer Verkabelung zwischen diesen Verteilern. Der Begriff des Verteilers bezeichnet weder den Schrank noch den Raum, sondern er liefert eher eine funktionelle Beschreibung. Obwohl der Begriff des „Verteilerpunktes“ diese Definition besser trifft, soll im Sinne der Norm im weiteren Text von Verteilern gesprochen werden. Folgende Verteiler kennt die Norm:

- Standortverteiler (SV),
- Gebäudeverteiler (GV),
- Etagenverteiler (EV),
- Sammelpunkt (SP).

Das zweite elementare Teilelement zum Aufbau einer Kommunikationskabelanlage besteht aus der Verkabelung zwischen den Verteilern bzw. auch aus der Verkabelung zwischen Teilnehmeranschluss (Endgerätedose) und Verteiler:

- Primärverkabelung,

- Sekundärverkabelung,
- Tertiärverkabelung.

Die DIN/EN 50173 definiert diese Teilelemente wie folgt:

Standortverteiler Verteiler, von dem die Primärverkabelung ausgeht.

Gebäudeverteiler Verteiler, an dem das (die) Sekundärkabel endet(n) und auf dem das (die) Primärkabel aufliegen darf (dürfen).

Etagenverteiler Verteiler, der zur Verbindung von Tertiärkabel, anderen Teilsystemen der Verkabelung und aktiven Geräten benutzt wird.

Sammelpunkt Verbindungspunkt im horizontalen Teilsystem der Verkabelung zwischen dem Etagenverteiler und dem informationstechnischen Anschluss.

Primärkabel Kabel, das den Standortverteiler mit dem (den) Verteiler(n) verbindet. Primärkabel dürfen auch Gebäudeverteiler direkt miteinander verbinden.

Tertiärkabel Kabel, das den Etagenverteiler mit dem (den) informationstechnischen Anschluss (Anschlüssen) oder dem (den) Sammelpunkt(en) verbindet.

Sammelpunkt wie auch die Tertiärverkabelung spielen im Zusammenhang mit der im Artikel betrachteten Backbone-Verkabelung keine Rolle, betrachtet werden nur die Primär- und Sekundärverkabelung.

Im Falle eines anstehenden Redesign dieser Verkabelung sind 2 Hauptfälle zu unterscheiden:

1. Ein neues Gebäude (oder ein neuer Etagenverteiler) muss an eine vorhandene Kommunikationsverkabelung angebunden werden.
2. Die vorhandene Verkabelung zum vorhandenen Gebäude (oder vorhandenen Etagenverteiler) unterstützt die geforderten Zugangsverfahren nicht mehr.

Insbesondere die beabsichtigte Weiternutzung einer vorhandenen Verkabelung, möglicherweise ergänzt um neue Teilstrecken, wird vor allem bei hohen Datenraten von mehr als 1 Gbit/s sehr häufig nicht ohne massive Änderungen im Design oder massiven Beschränkungen bei der Materialauswahl möglich sein. Die Qualität der Verkabelung gibt die maximale Übertragungsrate vor und auf diese Qualität kann man sowohl bei der Planung wie auch bei der Realisierung starken Einfluss nehmen. Um einen solchen Einfluss nehmen zu können, muss im Planungsprozess festgelegt werden, welche Datenraten im Backbone-Bereich im Laufe des Nutzungszeitraumes der Verkabelung potenziell bereitgestellt werden sollen. Daraus leiten sich die „nachrichtentechnischen“ Anforderungen ab. Nicht jede Erhöhung der Datenrate soll einen Wechsel oder eine Veränderung der Verkabelung nach sich ziehen. Eine gute Planung ist bestrebt, einen sehr langen Nutzungszeitraum für möglichst alle in diesem Zeitraum entwickelten bzw. benötigten Datenraten und Übertragungs-

Singlemode und Multimode in der Backbone-Verkabelung

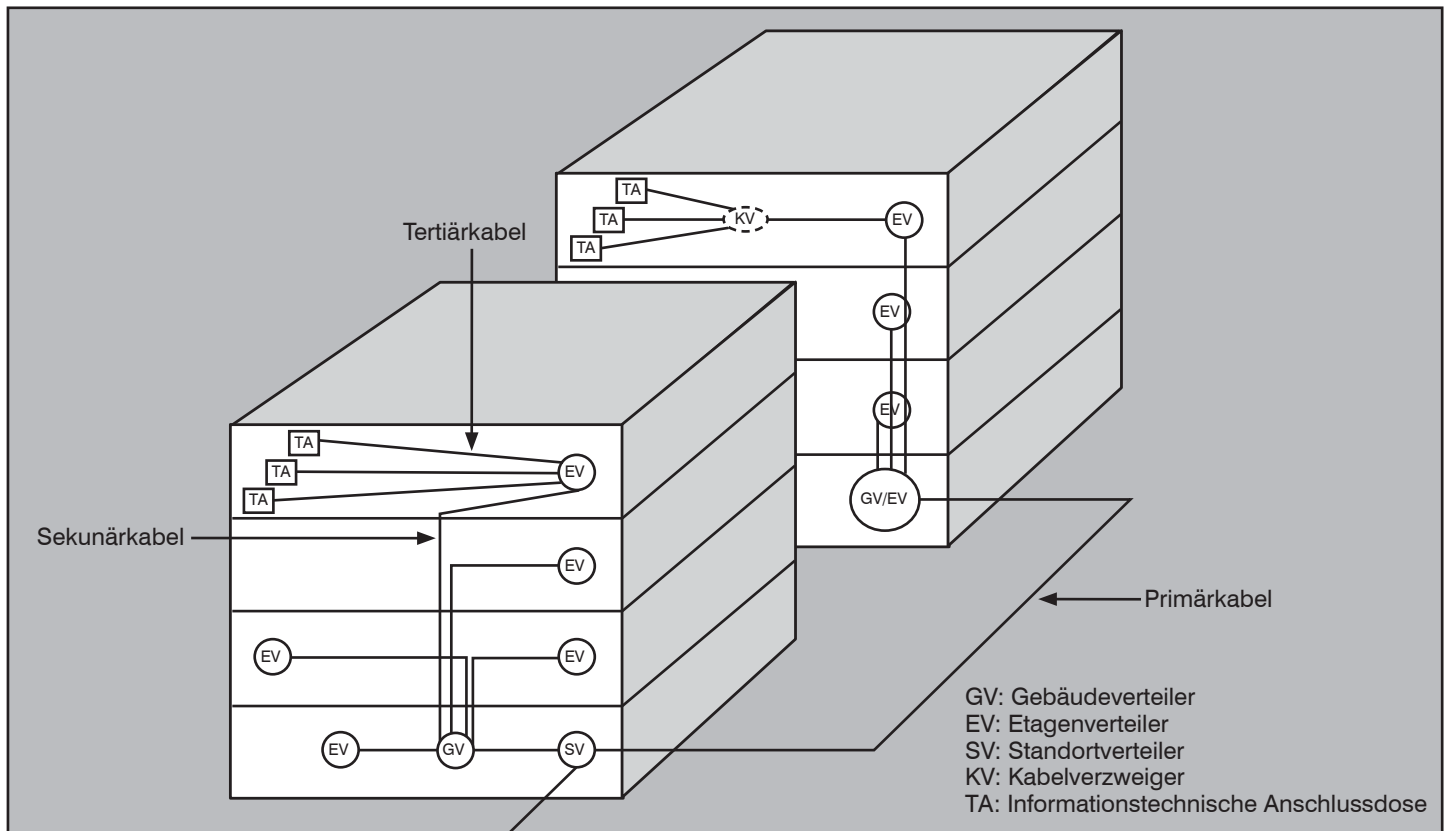


Abbildung 1: Teilelemente nach EN 50173

verfahren sicherzustellen. Eine Hilfestellung bei der Planung leisten dazu zwei wesentliche Elemente:

- Die Anforderungen der Kabelnormen, in Deutschland primär über die DIN EN 50173 spezifiziert.
- Die Anforderungen der Entwickler von Zugangsverfahren, weltweit im Bereich des Ethernets über die IEEE 802.3 spezifiziert.

Nutzbarkeit einer Kommunikationsverkabelung

Die Entwicklung der Verkabelungsnormen erfolgte derart, dass alle verfügbaren oder zum Zeitpunkt der Normverabschiedung „absehbar“ verfügbaren Übertragungsverfahren durch die Verkabelungsspezifikationen berücksichtigt werden, sofern die Kabelnormen eingehalten werden sind weitere Normen nicht zusätzlich zu berücksichtigen. Auf der anderen Seite benötigen die Anbieter von neuen Übertragungstechniken (z.B. Switches) eine möglichst große „passive“ Installationsbasis, um ihre Produkte zu verkaufen und die Hardware-Entwickler werden deshalb bemüht sein, neue aktive Systeme auf möglichst weit verbreiteten Verkabelungsinfrastrukturen einsetzen zu können.

Interessant ist der Vergleich, in welchen Zeitabständen beide „Normierungsinstanzen“ aktualisiert wurden, die nachfolgende Abbildung stellt auf einem Zeitstrahl die Updates der Kabel- und Ethernet-Normen dar. Dargestellt werden bei den Kabelnormen die Meilensteine, bei denen es zu finalen Überarbeitungen der Kabelstandards gekommen ist, nicht dargestellt sind die dazwischen liegenden Entwurfsvarianten. (siehe Abbildung 2)

Einige der dargestellten Ethernet-Updates führten zu einer erhöhten Anforderung an die Verkabelung, Beispiel 10GBaseT oder auch 40GBase. Der Zeitstrahl macht deutlich, dass ein zuvor vorausgegangenes Kabelnormen-Update eine Weiterentwicklung nicht immer berücksichtigen konnte und damit der bei der Planung angesetzte Nutzungszeitraum der Kabelnorm manchmal zwangsläufig zu kurz greifen muss.

Beispiel: Eine Planung einer Verkabelungsanlage zwischen 2003 und 2006 wird sich an der aktuellsten Kabelnorm (2003) orientieren. Der Planer wird den Anspruch haben, Materialien und - für den Backbone-Bereich besonders wichtig - die Längen der Strecken so zu planen, dass für einen Zeitraum von 5 bis 10 (oder gar mehr) Jahren diese Verkabelung für alle Datenraten genutzt werden kann. Wie wir später im Artikel noch genauer sehen wer-

den, hätte er in vielen Fällen diese Anforderung nicht erfüllen können; bereits nach 3 Jahren (!) würden durch einige Ethernet-Standards die Anforderungen so hoch gesetzt sein, dass er für die neuen Ethernet-Standards nachverkabeln muss.

Dieser Sachverhalt hat sich so oft wiederholt, dass im Prinzip jede, bezüglich der Glasfaser getroffene Entscheidung unter der Annahme eines minimalen Nutzungszeitraumes von mindestens 10 Jahren falsch war (eine Ausnahme: Singlemode-faser). Diesen Sachverhalt hat auch die Verkabelungsnorm erkannt und macht dazu folgende Aussage (Kapitel 4.4.3 der EN 50173-1): „Es ist grundsätzlich nicht möglich oder ökonomisch sinnvoll, eine Primär- und Sekundärverkabelung für die gesamte Lebensdauer der anwendungsneutralen Kommunikationskabelanlage zu installieren. Stattdessen darf die Planung auf den gegenwärtigen oder absehbaren Anforderungen der Netzanwendungen beruhen. Derartige kurzfristige Auswahlkriterien sind häufig dann für die Primär- und Sekundärverkabelung angemessen, wenn die Kabelführungssysteme für künftige Änderungen gut zugänglich sind“.

Mit dieser realistischen Einschätzung der Nutzbarkeit einer Backbone-Verkabelung soll im nachfolgenden Teil des Artikels eine Betrachtung der sinnvollen Anforder-

Singlemode und Multimode in der Backbone-Verkabelung

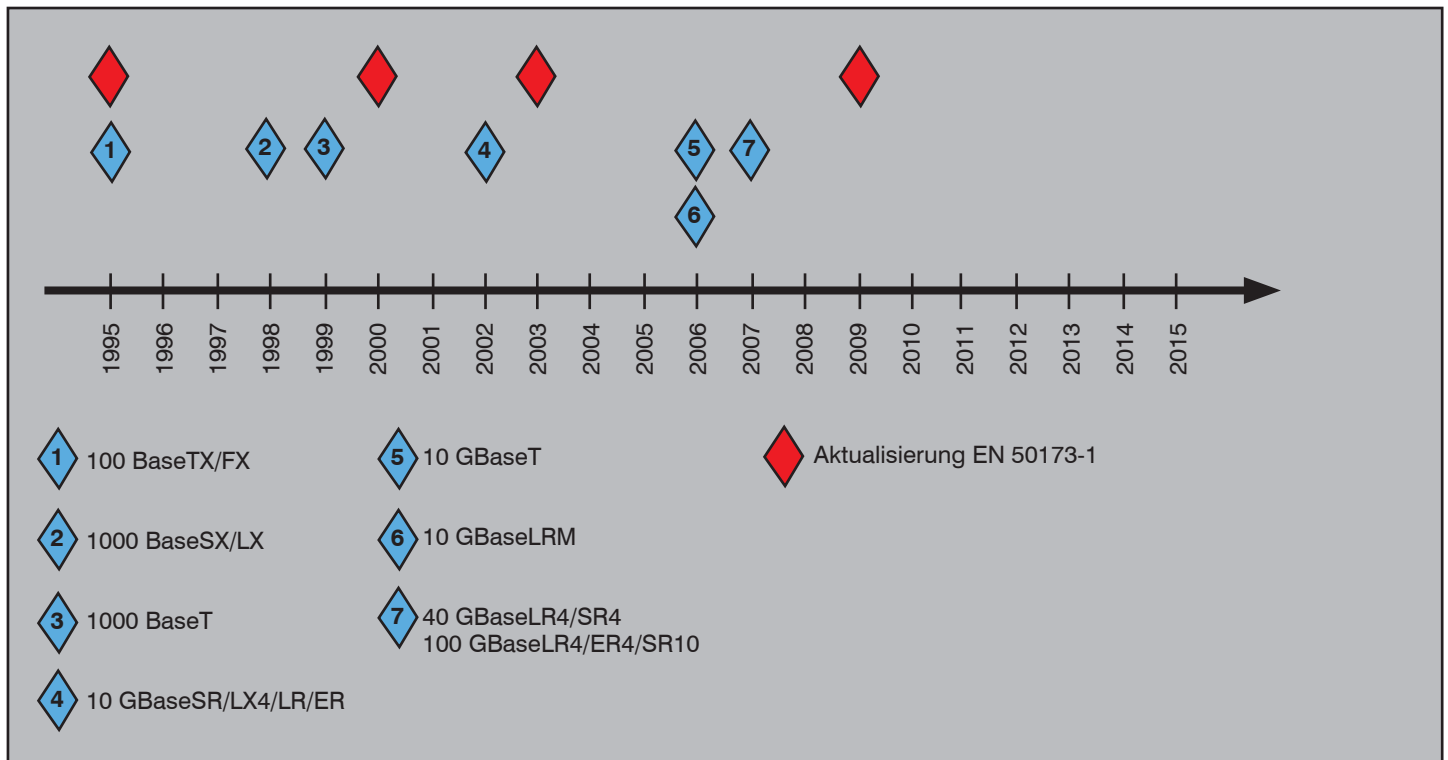


Abbildung 2: Aktualisierungen der Kabel- und Ethernet-Normen

ungen und Lösungsmöglichkeiten für Primär- und Sekundärverkabelungen erfolgen.

Die Planung einer Glasfaser-Backbone-Verkabelung zur Bereitstellung von LAN-typischen Übertragungstechniken konzentriert sich im Wesentlichen auf zwei Parameter, der Dämpfung (genauer gesagt der Einfügedämpfung) und der Streckenlänge. Die Rückflussdämpfung spielt derzeit eine untergeordnete Rolle, ihre Bedeutung nimmt zu, wenn auf einer Faser eine Übertragung in beide Richtungen stattfindet. Auch die Anfang der 90er-Jahre stark beachtete Numerische Apertur (= Maß für die Menge des eingekoppelten Lichtes) beeinflusst nur noch sehr gering die Planung. Dämpfung und Länge lassen sich bekanntlich nicht ganz voneinander trennen, denn eine große Länge führt zwangsläufig zu einer hohen Dämpfung. Andererseits kann es aber sein, dass trotz guter Streckendämpfung die Länge zu groß ist, dafür ist z.B. bei der Multimodefaser die Modendispersion verantwortlich. All dieses ist nicht neu und wird bei vollständigen Neuverkabelungen für die reine Betrachtung eines Kabels zwischen zwei Verteilern meistens berücksichtigt. Wurde die Anfang der 90er-Jahre im Vordergrund stehende Leitungsdämpfung mit Einführung von Gigabit etwas in den Hintergrund gestellt durch die modendispersionsbedingte Längeneinschränkung, so spielt die Dämpfung jedes einzelnen Elementes der Strecke heute wieder

eine deutlich größere Rolle. Grund dafür sind die sehr hohen Anforderungen an das Dämpfungsbudget bei Datenraten ab 1 Gbit/s. Dies ist vor allem dann ein Problem, wenn eine bereits vorhandene Verkabelung um weitere zusätzliche Strecken ergänzt werden soll oder wenn ein Übertragungslink aus mehreren Teilstrecken besteht. In diesem Falle hat die vorhandene oder auch neu zu realisierende Topologie einen erheblichen - häufig unterschätzten - Einfluss auf die Nutzbarkeit der Verkabelung.

Topologie-Betrachtung

Die Standard-Topologie, so wie sie innerhalb der EN 50173 seit Jahren als Richtlinie (keine Forderung!) postuliert wird, sieht sternförmig aus. Ausgehend von einem Standortverteiler werden die Ge-

bäudeverteiler mit diesem über einfache Verbindungen angeschlossen (Primärverkabelung), analog erfolgt die Anbindung der Etagenverteiler an den Gebäudeverteiler (Sekundärverkabelung). Die redundante Verbindung wird durch eine Querverkabelung mit dem Verteiler derselben Hierarchiestufe realisiert. Bei einer höheren Anforderung an die Verfügbarkeit wird man zwei Standortverteiler vorsehen und jeden hochverfügbaren Verteiler an zwei unterschiedliche übergeordnete (= höherwertigere) Verteiler anbinden. (siehe Abbildung 3)

Diese Topologie hat sich vor allem deshalb bewährt, weil die sternförmige Architektur von hierarchischen Ethernet-Netzwerken bestehend aus Access-, Distribution- und Core-Layer sehr gut darauf abzubilden ist. In nachfolgendem

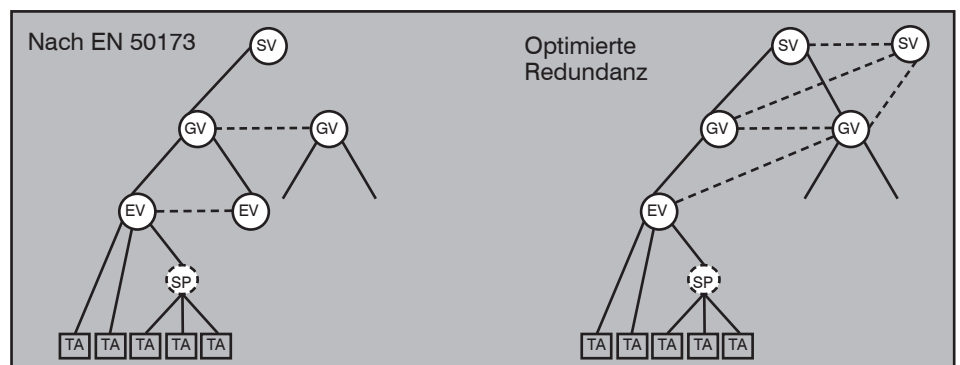


Abbildung 3: Sterntopologie nach Norm und Redundanz-Optimierung

Singlemode und Multimode in der Backbone-Verkabelung

einfachen Beispiel eines Ethernet-Netzwerkes mit nur zwei Hierarchieebenen (Distribution-Layer wird ausgespart) wird ein Access-Switch im Etagenverteiler über eine Durchrangierung in einem Gebäudeverteiler mit einem Core-Switch im Standortverteiler verbunden; jeder Access-Switch in jedem Etagenverteiler kann so an den Core-Switch angeschlossen werden. (siehe Abbildung 4)

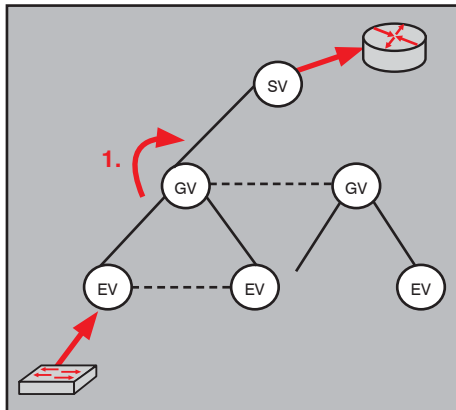


Abbildung 4: Einfaches Ethernet-Netzwerk in einer Sterntopologie

Nicht „normkonform“, aber typisch für „gewachsene“ Verkabelungen die insbesondere älter als 1995 sind, ist die ringförmige Topologie oder gar die vermaschte Topologie. Der Ring war Ende der 80er- und Anfang der 90er-Jahre „die“ Topologie, die bei Hochverfügbarkeit oder z.B. auch bei FDDI oder dem damals stark verbreiteten Tone-Ring verwendet wurde. Bei dieser Topologie gibt es keine Hierarchien in der Verkabelung, jedoch fordert wie bereits beschrieben die Ethernet-Technik eine Hierarchie. Das lässt sich wie im nachfolgenden Beispiel dargestellt im Prinzip auch sehr einfach realisieren. Einer der Verteiler wird als Position für den Core-Switch gewählt und alle Access-Switches in den anderen Verteilern werden mit Hilfe von Durchrangierungen mit diesem verbunden. In Abbildung 5 wurden z.B. 3 Durchrangierungen benötigt um diese Verbindung herzustellen.

Abbildung 6 zeigt eine Struktur, so wie sie (ungünstiger Weise) real geplant worden ist. Im Rahmen der Planung für das physikalische Design des Ethernet-Netzwerkes mussten Standorte für Backbone-Switches festgelegt werden, jeder Etagenverteiler mit Access-Switches musste über zum Teil sehr viele Durchrangierungen zu den Core-Switches hin geschaltet werden. Durch derartige „Durchschaltungen“ wird zum einen die Streckenlänge und zum anderen die Dämpfung der Gesamtstrecke sehr groß. Deshalb ist im Rahmen der Planung zu prüfen, ob diese Erhöhungen normkonform sind oder nicht. Dazu können zwei

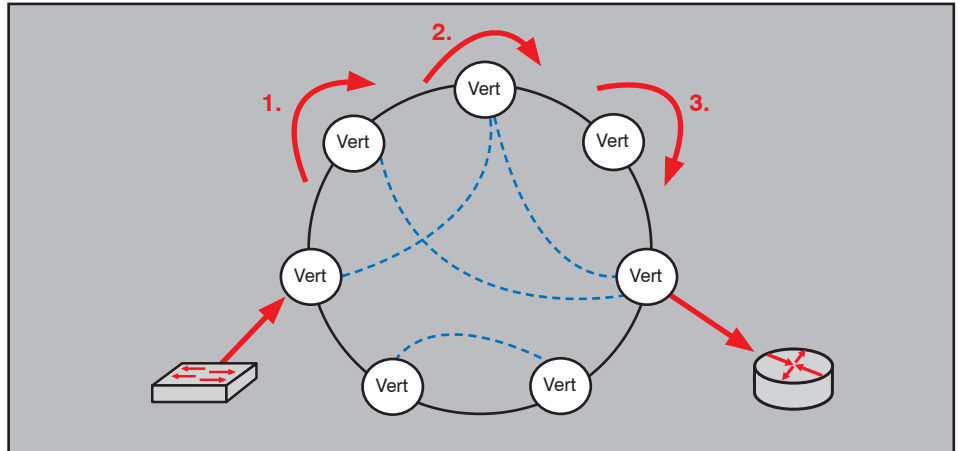


Abbildung 5: Einfaches Ethernet-Netzwerk in einer Ringtopologie

Planungshilfen herangezogen werden:

- Spezifikationen der IEEE802.3
- Spezifikationen der EN 50173-1

Übertragungstechnische Werte nach Norm

Beginnen wir bei der Betrachtung mit der - im Regelfall für den Verkabelungsplaner bekannten - Spezifikation der EN 50173-1. Tabelle 1 stellt sehr übersichtlich dar, welches Übertragungsverfahren mit welcher Glasfaser über welche Entfernung genutzt werden kann und welche optische Übertragungsklasse (OF-Klasse) die Strecke dazu erfüllen muss. Dieser Übertragungsklasse lässt sich mit einer weiteren Tabelle der Norm ein ganz konkreter maximaler Dämpfungswert für die Strecke zuordnen, dieser muss eingehalten werden.

Dazu ein Beispiel:

Das Ethernet-Übertragungsverfahren 1000 BaseSX darf nach Tabelle 1 mit einer OM2-Faser bis mindestens 550 m genutzt werden, wenn die Strecke die optische

Übertragungsklasse OF500 einhält. Gemäß Norm EN 50173-1 (Tabelle 42) wird OF500 eingehalten, wenn die Dämpfung im 2. optischen Fenster 3,25 dB (Channel Insertion Loss) nicht überschreitet.

Obwohl bei der Spezifikation der Verkabelungsnorm sinnvollerweise die Forderungen der Ethernet-Spezifikationen aus der IEEE übernommen wurden, ist festzustellen, dass umgekehrt die Spezifikationen der IEEE über viele Jahre hinweg nie Bezug genommen auf das Modell der EN 50173 mit seinen Faserklassen oder optischen Übertragungsklassen haben. Bis zur Normierung von 10GBaseLRM wurden in den Tabellen der IEEE dämpfungstechnische Anforderungen immer auf Basis von unterschiedlichen Fasern mit unterschiedlichen Bandbreite-Längen-Produkten spezifiziert (unabhängig von OM1, OM2 oder OM3). Dazu zwei Beispiele in Form von Auszügen aus der IEEE 802.3ae (Ethernet mit 10Gbit/s), erkennbar ist, dass nirgendwo ein Hinweis auf Reichweiten mit OM1 oder ähnlichem existiert: (siehe Tabelle 2 und 3)

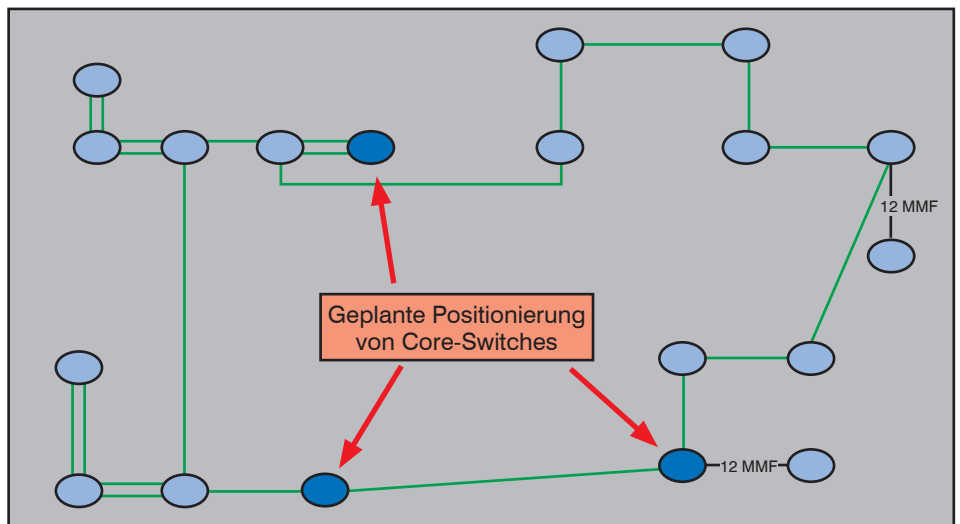


Abbildung 6: Praxisbeispiel Ethernet im Ring

Singlemode und Multimode in der Backbone-Verkabelung

LWL-Anwendung	λ	Fasertyp	Multimodefaser					
			OM1		OM2		OM3 / OM4	
			Reichweite	Klasse	Reichweite	Klasse	Reichweite	Klasse
10BASE-FL	850 nm	G50	1514 m	OF-500	1514 m	OF-500	1514 m	OF-500
		G62,5	2000 m	OF-2000	2000 m	OF-2000	-	-
100BASE-FX	1300 nm	G50	2000 m	OF-2000	2000 m	OF-2000	2000 m	-
		G62,5	2000 m	OF-2000	2000 m	OF-2000	-	-
10000BASE-SX	850 nm	G50	-	-	550 m	OF-500	550 m	OF-500
		G62,5	275 m	OF-100	-	-	-	-
10000BASE-LX	1300 nm	G50	550 m	OF-500	550 m	OF-500	550 m	OF-500
		G62,5	550 m	OF-500	550 m	OF-500	-	-
		E9	-	-	-	-	-	-
10GBASE-SR/SW	850nm	G50	-	-	82 m	-	300 m	OF-300
		G62,5	32 m	-	-	-	-	-
10GBASE-LX4	1300 nm	G50	300 m	OF-300	300 m	OF-300	300 m	OF-300
		G62,5	300 m	OF-300	300 m	OF-300	-	-
		E9	-	-	-	-	-	-
10GBASE-LR/LW	1310 nm	E9	-	-	-	-	-	-
10GBASE-ER/EW	1550 nm	E9	-	-	-	-	-	-
40GBASE-SR4	850 nm	G50	-	-	-	-	100 m	OF-100
100GBASE-SR10	850 nm	G50	-	-	-	-	100 m	OF-100

Tabelle 1: Auszug aus der Tabelle aus der EN 50173-1

Parameter	62.5 μm MMF	50 μm MMF		10 μm SMF	Unit
Modal bandwidth as measured at 1300 nm (minimum, overfilled launch)	500	400	500	n/a	MHz·km
Operating distance	300	240	300	10000	m
Channel insertion loss	2.0	1.9	2.0	6.2	dB

Tabelle 2: Auszug aus der Spezifikation der IEEE 802.3ae 10GBase-LX4

Schwierig wird es für alle, die bei ihrer Altverkabelung nicht wissen, welche Faserqualität bzw. welches Bandbreite-Längenprodukt die Faser hat. Dieses ist im Nachgang mit Feldmessgeräten nicht mehr nachprüfbar und man wird im Zweifel von einer Qualität im unteren Bereich mit allen damit verbundenen Einschränkungen ausgehen müssen.

Was wird in den bisher aufgeführten Tabellen in Zusammenhang mit einer Backbone-

Verkabelung deutlich?

- Ab einer Streckenlänge von 300 m ist 10 Gbit/s mit einer Multimodefaser weder nach Verkabelungsnorm noch nach IEEE nutzbar.
- Eine 62,5 μ m-Faser ist für 10 Gbit/s im Prinzip nur in Kombination mit der teuren LX4-Technik „backbone-tauglich“, die üblichere (und kostengünstigere) SR-Technik ist mit 26 m und 33 m indiskutabel.

- Selbst eine herkömmliche 50 μ m-Faser (nicht OM3) lässt sich im Backbone-Bereich für 10 Gbit/s nur mit LX4-Technik sinnvoll verwenden.

- Das Dämpfungsbudget liegt für 10 Gbit/s und Multimodefaser im Bereich von 1,9 dB bis 2,6 dB und ist damit noch anspruchsvoller als bei 1 Gbit/s.

Die einzig „wahre“ Multimodelösung scheint OM3 in Kombination mit 10GBaseSR zu sein; wenn da nicht 10GBaseLRM wäre. Auch hier lohnt sich ein Blick in die IEEE, da insbesondere diese Technik in der Tabelle der Verkabelungsnorm EN 50173-1 überhaupt nicht aufgeführt wird (die Gründe sind dem Autor nicht bekannt). Zum ersten Mal nimmt eine IEEE-Norm Bezug auf die OM-Faserklassen, von größerer Bedeutung ist aber, dass auch Nutzer von „Altverkabelungen“ mit 62,5 μ m-Fasern oder auch „schlech-

Description	62.5 μm MMF		50 μm MMF			Type B1.1, B1.3 SMF			Unit
Nominal wavelength	850					1310	1550		nm
Modal bandwidth (min)	160	200	400	500	2000	N/A	N/A		MHz·km
Operating distance (max)	26 m	33 m	66 m	82 m	300 m	10 km	30 km	40 km	
Channel insertion loss (max)	2.6	2.5	2.2	2.3	2.6	6.0	11.0		dB

Tabelle 3: Auszug aus der Spezifikation der IEEE 802.3ae 10GBase-SR/LR

Singlemode und Multimode in der Backbone-Verkabelung

Multimode fiber type	ISO/IEC 11801:2002 fiber type	Operating range (m)	Maximum channel insertion loss (dB)
62.5 μ m 160/500		0.5 to 220	1.9
62.5 μ m 200/500	OM1	0.5 to 220	1.9
50 μ m 500/500	OM2	0.5 to 220	1.9
50 μ m 400/400		0.5 to 100	1.7
50 μ m 1500/500	OM3	0.5 to 220	1.9

Tabelle 4: Auszug aus der Spezifikation der IEEE 802.3aq 10GBase-LRM

ten“ 50 μ m-Fasern in den Genuss kommen, bei 10 Gbit/s Reichweiten mit ihren Kabeln von deutlich mehr als 100 m realisieren können. Für alle 3 Fasertypen des Kabelstandards wird eine Reichweite von bis zu 220 m spezifiziert. Damit liegt 10GBase-LRM nur 80 m unterhalb der lange postulierten Idealkombination von OM3/10GBaseSR. Neben der grundsätzlichen Nutzbarkeit von alten Fasern bietet 10GBaseLRM einen weiteren Vorteil: Die Kosten der SFP-Module für die Switches sind nochmals günstiger als die der 10GBaseSR-Technik, Preisunterschiede von bis zu 50% sind erzielbar. Neben den beiden Vorteilen ist allerdings auch ein Nachteil in der Tabelle erkennbar, das maximale Dämpfungsbudget liegt knapp unter 2 dB, also nochmals 0,5 dB (= 1 Steckverbindung) niedriger als bei 10GBaseSR (siehe Tabelle 4).

Die Planungserfahrungen des Autors zeigen, dass aktuelle Backbone-Planungen sowohl für die aktiven wie auch die passiven Netzwerkelemente deutlich im Fokus einer Datenrate von maximal 10 Gbit/s stehen. Die im Rahmen von Rechenzentrumsplanungen bereits berücksichtigte nächste Generation mit 40 Gbit/s oder gar 100 Gbit/s wird nur sehr zögerlich als Maßstab für eine Backbone-Verkabelung angelegt; wer eine absolut zukunftsichere Glasfaserverkabelung bereits heute realisieren möchte, muss sich mit den Anforderungen dieser beiden Datenraten beschäftigen. Zunächst einmal sollen einige, gerade für die Planung von aktuellen Backbone-Verkabelungen wichtige Zwischenergebnisse festgehalten und deren Einfluss anhand von Beispielen veranschaulicht werden.

Projektbeispiele

Die Länge und Dämpfung der Verkabelungsstandards basieren auf den Anforderungen der IEEE, wie oben dargelegt mit übernommenen Dämpfungsmaximalwerten. Die Dämpfung spielt bei direkten Verbindungen zwischen zwei aktiven Knoten über 1 Kabel (mit einer Steckverbindung an jedem Ende) eine geringere Rolle als die bandbreitenbedingte, faserspezifische Längenrestriktion. Wa-

rum ist das so? Nimmt man als Beispiel 10GBaseLRM mit ca. 2 dB Maximaldämpfung und kalkuliert eine Glasfaserstecker-Verbindung in mittlerer Qualität mit ca. 0,5 dB, so bleiben nach Abzug der beiden Steckverbindungen an beiden Enden noch ca. 1 dB reine Faserdämpfung. Dies wäre bei einer mittleren Faserqualität eine Strecke von ca. 300 bis 400 m, der Standard lässt aber „nur“ 220 m zu. Diese Betrachtung ändert sich bei „durchgeschalteten“ Verbindungen zwischen zwei aktiven Knoten über mehrere Kabel hinweg, jetzt spielt die Dämpfung eine - häufig unterschätzte - größere Rolle. Dazu ein ausführliches Projektbeispiel:

Der IT-Netzbetreiber beabsichtigte, auf einer bereits existierenden Backbone-Verkabelung ein Ethernet-Netz zu realisieren und strebte eine Einführung von 10 Gbit/s als Uplink zwischen Access-Switches und Core-Switches an. Eine einfache Doppelauslegung des Core-Bereiches mit „nur“ zwei Core-Switches war aus Gründen der Längenrestriktion nicht möglich, deshalb wurde ein aus 3 Core-Switches bestehender Collapsed-Backbone geplant (siehe Abbildung 7). Der Access-Switch eines jeden Etagenverteilers sollte an mindestens zwei Core-Switches angebunden werden.

Im Bild sieht man die beispielhafte Zuordnung eines Access-Switches (ASW), dieser ist an Core-Switch 1 und Core-Switch 2 (CSW1 und CSW2) anzuschließen. Zu berücksichtigen ist, dass für beide Anbindungen bzw. Wege eine unterschiedliche Anzahl von Durchrangierungen notwendig ist. Jede Durchrangierung in einem Verteiler bedeutet, es werden zwei Rangierfelder (bzw. deren Stecker) mit einem Rangierkabel verbunden und dies ergibt bei einer nach Kabelstandard kalkulierten typischen Steckverbinderdämpfung von je 0,75 dB pro Durchrangierung eine Dämpfung von 1,5 dB reine Rangierdämpfung pro Verteiler. Hinzu kommen nochmals 2 x 0,75 dB zur Aufschaltung der beiden Switches. Lässt man im Beispiel die Faserdämpfung völlig aus der Rechnung heraus, so ergeben sich die in der Tabelle 5 dargestellten Dämpfungswerte, die beiden unteren Zeilen beinhalten sogar eine Berechnung mit einem „realistischeren“ Dämpfungswert für eine Steckverbindung mit 0,4 dB. Zum Vergleich nun einige maximale Dämpfungsbudgets der IEEE:

- 1000BaseSX: 3,25 dB
- 10GBaseSR: 2,6 dB
- 10GBaseLRM: 1,9 dB
- 10GBaseLR: 6 dB

Die ersten drei technischen Lösungen basieren auf Multimodefasern, sie wären auch bei vernachlässigter Faserdämpfung (in der Regel ca. 3 dB/km) fast gar nicht nutzbar. Erst die letzte Technik auf Basis von Singlemodefaser würde mit einem ausreichenden Puffer für die Faserdämpfung „brauchbare“ Distanzen zulassen. Genau dies war im Projektbeispiel die unangenehme Konsequenz bei der Planung der Switches bzw. deren SFP-Techniken, es mussten die teuren Techniken einge-

Sonderveranstaltung**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 herstellereinspezifischen Lösungen wird auch auf Planungs- und installationsbegleitende Maßnahmen eingegangen, die im Rahmen einer anstehenden Verkabelung zu beachten sind.

Referent: Dipl.-Ing. Hartmut Kell
Preis: € 890,- zzgl. MwSt.



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

Singlemode und Multimode in der Backbone-Verkabelung

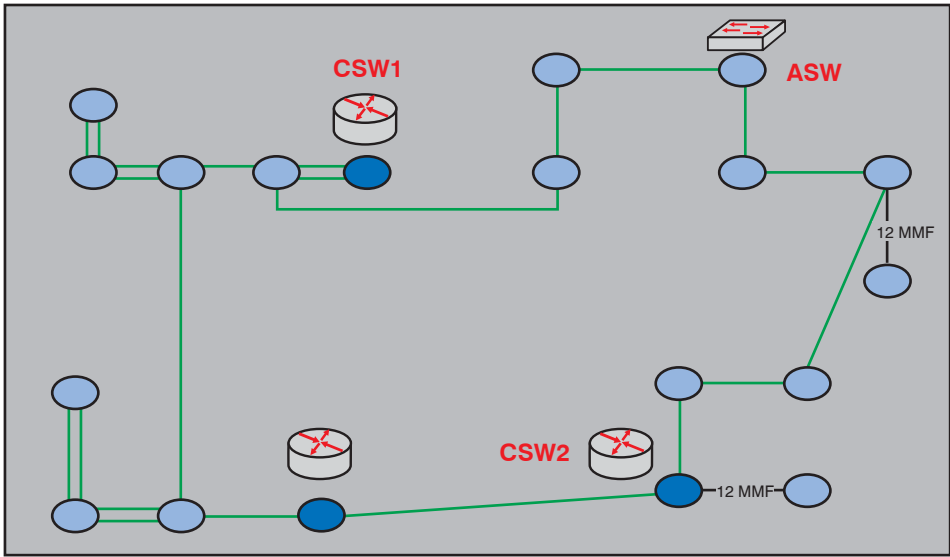


Abbildung 7: Wechselwirkung passive und aktive Netzwerk-Topologie, Projektbeispiel

Verbindung von		Durch-rangierungen	Einaus-kopplungen	Stecker-dämpfung	Gesamtdämpfung ohne Faser/Spleißdämpfung
ASW	CSW1	3	2	0,75 dB	6,00 dB
ASW	CSW2	4	2	0,75 dB	7,50 dB
ASW	CSW1	3	2	0,40 dB	3,20 dB
ASW	CSW2	4	2	0,40 dB	4,00 dB

Tabelle 5: Dämpfungsbudget am Projektbeispiel

setzt werden; glücklicherweise standen Singlemodefasern in ausreichender Anzahl zur Verfügung. (siehe Abbildung 6 und 7)

In einem weiteren Projekt bildete eine sehr stark vermaschte Topologie mit über 50 Verteilern trotz der vielen (Quer-)Verbindungen ein Problem bei der Einführung von neuen Gigabit-Techniken. Dieses, für eine Industrieumgebung sehr typisches Netz war dadurch geprägt, dass es eine große Anzahl von weit gestreuten kleineren Etagenverteilern (genau genommen sind es „Kleingebäudeverteiler“) mit wenigen Endgeräteanschlüssen bzw. Access-Switches gab. Der Anschluss der Access-Switches mit Hilfe von einfachen 100Mbit/s-Uplinks an zwei Core-Switches wäre kein Problem gewesen (Dämpfungsbudget von über 10 dB und Glasfaservlänge von 2.000 m), doch im Sinne einer zukunfts-sicheren Planung sollte eine Uplink-Qualität von mindestens 1 Gbit/s bereitgestellt werden. Somit standen nur 3,25 dB als Dämpfungsbudget zur Verfügung und wie oben bereits dargestellt, war die Anzahl von möglichen Durchrangierungen über mehrere Verteiler hinweg sehr eingeschränkt. Eine Nachverkabelung wäre aufgrund der ungünstigen Geländesituation nur mit sehr hohen Kosten möglich gewesen. Das Ergebnis

war, dass man zum Aufbau eines neuen Netzes in 10(!) sich im Gelände befindenden Verteilern Core-Switches positionieren musste. Dabei sorgten einzelne Core-Switches nur für die Anbindung von 2 bis 3 Access-Switches, auch die Realisierung des Redundanzkonzeptes wurde

trotz der hohen Anzahl von Core-Switches sehr schwierig.

Neben den beschriebenen planerischen Aspekten des sehr engen Dämpfungsbudgets bei hohen Datenraten ist im Zusammenhang mit der Errichtung des Netzes bzw. der Realisierung der Uplinks ein weiterer Aspekt von zunehmender Bedeutung: Die messtechnische Überprüfung des Dämpfungsbudgets. Wurde die Einmessung einer komplett beschalteten Glasfaserstrecke bisher nur sehr selten bei einer Ausschreibung zum Aufbau des aktiven Netzwerkes vorgesehen (eher im Rahmen der Abnahmemessung der Kabelinstallation), so kann nur dringend geraten werden, dies zur Sicherstellung eines nicht ausgereizten Dämpfungsbudgets grundsätzlich vor der Aktivierung der Links vorzunehmen. Verschmutzte oder zerkratzte Steckerstirnflächen wirken sich bei knappem Dämpfungsbudget möglicherweise gravierend aus und sollten messtechnisch erfasst bzw. ausgeschlossen werden.

Ausblick: 40 Gbit/s und 100 Gbit/s im Backbone

Zurück zur Analyse des Einflusses von 40 und 100 Gbit/s. Die Tabelle 6 verdeutlicht sehr einfach, mit welchen Fasertypen welche Reichweiten zu erzielen sind (Informationen der IEEE, nicht der Kabelstandards).

Bereits die Längenkategorien 40km/10km/100m machen deutlich, dass 40/100-Gigabit-Ethernet zum jetzigen Zeitpunkt nicht den Schwerpunkt im Backbone-Bereich setzt. Entweder gibt es Län-

Anwendung	Glasfaser				Bezeichnung und Spezifikation
	Reichweite	40 km	10 km	100 m	
	Technologie	Singlemode	Singlemode	Multimode	
	40-GbE		40-GBASE-LR4 4 x 10-GbE, 4 Pfade/Link = 2 Fasern/Link, OS1	40-GBASE-SR4 4 x 10-GbE, 4 Ports/Link = 8 Fasern/Link, OM3 (TypA1a nach IEC 60793-2-10)	
	100-GbE	100-GBASE-ER4 4 x 25-GbE, 4 Pfade/Link = 2 Fasern/Link, OS1 (Typ B1.1 & B1.3)	100-GBASE-LR4 4 x 25-GbE, 4 Pfade/Link = 2 Fasern/Link, OS1 (Typ B1.1 & B1.3)	100-GBASE-SR4 10 x 10-GbE, 10 Ports/Link = 20 Fasern/Link, OM3 (TypA1a nach IEC 60793-2-10)	

Tabelle 6: Faserauswahl bei 40/100 Gbit/s

Singlemode und Multimode in der Backbone-Verkabelung

genkategorien mit mindestens 10 km (Zielgruppe: Provider oder MAN-Verbindungen) oder Kategorien bis 100 m (Zielgruppe Rechenzentrum). Beschränken wir uns bei der Analyse auf 10 km und 100 m. Folgt man dem bisher üblichen Ansatz, Multimodetechnik als Basistechnik zu verwenden, so würden beide Datenraten nur Backbone-Verbindungen von maximal 100 m zulassen, das reicht knapp für typische Sekundärverkabelungen. Diese Möglichkeit würde aber auch nur OM3-Fasern zur Verfügung stehen, schlechtere Faserqualitäten erlauben gar keine 40/100Gbit/s. Von extrem großer Bedeutung ist aber, dass man bei 40 Gbit/s 8 Multimodefaser für einen Link und bei 100 Gbit/s bereits 20 Multimodefasern vorsehen muss. Diese müssten sich alle in einem Kabel befinden und mit einem sogenannten MPO-Stecker abschließen (= Multi Push On). Entscheidende Frage: Wer wird heute im Backbone-Bereich bei seiner Planung Glasfaserverbindungen vorsehen, bei denen vielfache von 20 Fasern verlegt werden und diese nicht in Rangierfelder mit 20-Einzelports enden, sondern in 20-Port-Steckern? Dies kann fast ausgeschlossen werden und damit ist 40/100-Gigabit-Ethernet in Multimodetechnik eine kaum vorstellbare Option für den Backbone-Bereich.

Anders dagegen Singlemode, hier lässt sich weiterhin die 2-Fasertechnik verwenden und diese bietet mit Längen von 10.000 m weit mehr Reichweite, als in den meisten Geländenetzen notwendig ist.

Bewertung Singlemode versus Multimode

Folgende wichtigen Zwischenerkenntnisse lassen sich zusammenfassen:

- Singlemode ist bei hohen Übertragungsraten (ab 1 Gbit/s) wichtiger denn je.
- Singlemode erleichtert die Planung von

redundanten Switching-Topologien, besonders in stark vermaschten Topologien (oder Ringtopologien).

- Aber: Singlemode erhöht die Kosten bei den aktiven Komponenten.
- Bei Einführung von 40/100 Gbit/s im Backbone-Bereich ist die Multimodefaser definitiv „am Ende“, dies ist nur sinnvoll mit Singlemodetechnik zu realisieren.

Also gar keine Multimodefaser mehr in der Primär- und Sekundärverkabelung? Diese extreme Abkehr wird vom Autor nicht als sinnvoll erachtet, zum einen wären auch bei kurzen Strecken der Einsatz der teureren GBIC/SFP-Techniken notwendig. Zum anderen ist derzeit ein Trend festzustellen, dass neue Anwendungen aus dem Bereich der Gebäudemess- und regelungstechniken (noch) keine Gigabit- oder Mehr-als-Gigabit-Übertragungsraten benötigen und es deshalb dazu keine entsprechenden elektronischen Komponenten gibt, hier wird weiterhin sehr stark auf Multimode gesetzt. Deshalb scheint eine Reduzierung der Anzahl von Multimodefasern in einer Backbone-Verkabelung sinnvoller zu sein als der völlige Wegfall.

Multimodeauswahl

Lassen die Distanzen zwischen den Verteilern Multimodetechnik zu, so wird sich natürlich die Frage stellen, welche der derzeit kategorisierten 4 Fasertypen OM1 bis OM4 zu verwenden sind. Dargelegt wurde, dass die „älteren“ Fasern OM1 bis OM2 bis zu einer Reichweite von 220 m durchaus für 10 Gbit/s genutzt werden können. Über diese Distanz hinaus geht wohl kein Weg an OM3 oder OM4 vorbei. OM4 bringt normativ keinen Vorteil in diesem Einsatzumfeld, deshalb sind aus Sicht des Autors keine Notwendigkeiten erkennbar, von der OM3-Faser abzuwei-

chen. Diese wird also bei absoluten Neuverkabelungen ohne Altbestand die Nummer-1-Wahl sein für Strecken bis 300 m. In Umgebungen mit vorhandenen OM1/OM2-Fasern kann natürlich auch OM3 nachinstalliert werden, in diesem Falle gibt es jedoch keine nennenswerten Erfahrungen bei einem Mix von OM1/OM2/OM3-Fasern innerhalb eines Links, hier wird man wohl eigene Tests durchführen müssen. Die Standards sehen diesen Mix natürlich nicht vor.

Neben den Reichweiten, welche die standardisierten Fasern nutzbar machen, muss auch darauf hingewiesen werden, dass viele Kabel- bzw. Faserhersteller mit eigenen, optimierten Fasern auf dem Markt auftreten und für diese Fasern bei Einsatz von Standardelektronik deutliche höhere Reichweiten versprechen. Marktinggerecht werden solche Fasern dann auch als OM3+ oder OM3e o.ä. gekennzeichnet. Corning sagt z.B. mit seiner OM3+-Faser Infinicor eSX+ eine Erhöhung der Reichweite für 10GBaseSR um 250 m im Vergleich zum Standard zu. Unabhängig von der Zuverlässigkeit der Herstellerzusagen ist festzuhalten, dass es nach Erfahrungen des Autors seitens der Switch-Hersteller hier keinerlei Garantie zu den höheren Reichweiten gibt, hier beschränkt man sich weiterhin auf die Längen der IEEE-Spezifikationen. Deshalb kommt ein Nutzer dieser besseren Fasern an eigenen Tests nicht vorbei.

Singlemodeauswahl

Weit weniger bekannt ist die Vielfalt der zur Verfügung stehenden Singlemodefasern. Planer, die nur mit dem Verkabelungsstandard vertraut sind, kennen dazu die Kategorie OS1, neuerdings auch OS2. Diese Unterteilung ist jedoch sehr grob, wie die Tabelle 7 zeigt.

Zunächst einmal ist auffällig, dass es drei

Bezeichnung ITU-T ...	Bezeichnung IEC (60793-2)	Bezeichnung EN 50173	Typ	Optimiert für (nm)
G.652 a	B1.1		Non Dispersion Shifted Fiber	1310, 1550, 1625
G.652 b	B1.1		Non Dispersion Shifted Fiber	1310, 1550, 1625
G.652 c	B1.3	OS1 & OS2	Low Water Peak Single Mode	1310 bis 1550
G.652 d	B1.3	OS2 & OS2	Low Water Peak Single Mode	1310 bis 1550
G.653	B2		Dispersion Shifted Fiber	1310, 1550
G.654	B1.2		Cutoff Shifted Fiber	1550
G.655	B4		Non Zero Dispersion Shifted Fiber	1550, 1625
G.656	B5		Non Zero Dispersion Shifted Fiber	1460, 1625
G.657 A-Typen	B6	OS1	bending loss insensitive single mode	1310 bis 1625
G.657 B-Typen	B6		bending loss insensitive single mode	1310, 1550, 1625

Tabelle 7: Vielfalt der Singlemodefasern

Singlemode und Multimode in der Backbone-Verkabelung

unterschiedliche Normierungsinstitutionen gibt, die sich mit der Kennzeichnung von Singlemodefasern beschäftigen:

Die IEC (60793-2) verwendet die im Zusammenhang mit Glasfaserausschreibungen häufig verwendete Bezeichnung mit einem voranstehenden Buchstaben „G“. In der ITU dagegen wird eine Bezeichnung mit dem Buchstaben „B“ vorneweg verwendet und die EN 51073, die nur einen Teil der in der IEC und ITU spezifizierten Fasern übernimmt, verwendet letztendlich die im IT-Bereich eher bekannten OS1/OS2-Abkürzungen. Die verschiedenen Typen unterscheiden sich in erster Linie in dem nutzbaren optischen Bereich (=Wellenlängenbereich), für den sie ausgelegt wurden. Um LAN-spezifische CWDM-Techniken im Umfeld einer IT-Verkabelung nutzen zu können, benötigt man relativ breitbandige Fasern, die im Bereich zwischen 1300 nm und 1650 nm eingesetzt werden können. Einige der aufgeführten Fasern liegen aber entweder völlig außerhalb dieses Bereiches oder besitzen nur punktuell ein nutzbares optisches Fenster (z.B. die Faser G.625b deckt genau die drei Wellenlängen 1310

nm, 1550 nm und 1625 nm ab). Mit einer Anforderung nach Nutzbarkeit von CWDM reduziert sich die große Anzahl von Fasern auf G.652c-, G.652d- und G.657A-Typen, dies sind die entsprechenden OS-Typen der EN50173-1.

Viele Nutzer vorhandener Verkabelungen werden sich noch nie Gedanken gemacht haben, ob der bereits verlegte Typ denn passend ist, hier kann zur Beruhigung davon ausgegangen werden, dass in den letzten Jahren zumeist eine G.652c-Faser verlegt worden ist und man damit eine normkonforme Singlemodefaser ohne Einschränkungen im Campus-Bereich nutzen kann.

Eine Besonderheit stellen die G.657A-Typen dar, ihre Fähigkeit, extrem enge Biegeradien zuzulassen wird im Bereich der Glasfaserverkabelungen in Wohnhäusern bevorzugt genutzt (die Faser wird teilweise mit Hilfe von „Tacker-ähnlichen“ Befestigungswerkzeugen z.B. direkt auf Holzverkleidungen festgemacht), diese Eigenschaft ist im Bereich der Backbone-Verkabelung meistens nicht gefordert.

Zusammenfassung

Bei der Backbone-Verkabelung in Geländen mit großer Ausdehnung ist die Multimodefaser ganz klar auf dem Rückweg, gerade Nutzer von sehr hohen Datenraten werden mit diesem Fasertyp permanent mit Planungsschwierigkeiten bzw. -einschränkungen rechnen müssen. Die Multimodefaser kann deshalb, trotz der geringeren Kosten bei den aktiven Komponenten, nicht mehr den Hauptanteil der Fasern bei Neuverkabelungen übernehmen. Die stärkere Berücksichtigung der Singlemodefaser bringt insbesondere Vorteile durch eine fast unbegrenzt nutzbare Datenrate, keinerlei Längenrestriktionen im LAN-Bereich, eine einfache Typauswahl der Faser, Beibehaltung der Zwei-Faserigkeit und eine sich abzeichnende praxisfreundliche Standardisierung der Glasfaserstecker. Trotzdem wird empfohlen, die Multimodefaser nicht völlig bei der Planung von neuen Glasfaserstrecken rauszunehmen, es ist lediglich mit einem deutlich geringeren Anteil zu planen, dieser ist aber abhängig von der Nutzungsform der Backbone-Verkabelung.

Seminar

Sommerschule 2011 - Intensiv-Update auf den letzten Stand der Netzwerktechnik 18.07. - 22.07.11 in Aachen

Die Sommerschule gibt innerhalb von 5 Tagen den kompakten und intensiven Überblick über die neusten Entwicklungen im Umfeld der Netzwerk-Technologien: Anforderungen an zukunftssichere Netzwerke: was ändert sich; Neue Technologien und Standards; Design-Verfahren im Vergleich; Ausgewählte Technologien in der Analyse; Umfeld-Analyse: was passiert um Netzwerke herum

Die Sommerschule wendet sich an Teilnehmer mit bestehenden Grundlagen-Kenntnissen und ist als Weiterbildung für berufserfahrene Netzwerker konzipiert.

Die 5 Intensivtage der Sommerschule 2011 teilen sich wie folgt auf:

Switching 2011: Stand der Technik und Prognosen
IPv6: von den Grundlagen zur Migration
Rechenzentrum: Server und Speicher im Netzwerk
Infrastrukturen und Trouble Shooting
Verkabelung 2011: aktuelle Entwicklungen, Alternativen und Zukunftsprognosen
Trouble Shooting 2011: mehr Bandbreite, stabilere Netzwerke, höhere Protokolle -
Herausforderungen und Strategien für eine professionelle Fehlersuche und Beseitigung
Sicherheit 2011

Die Sommerschule 2011 greift die aktuellsten Entwicklungen auf, analysiert die Auswirkungen auf Ihre Netzwerke und diskutiert Entscheidungs- und Handlungs-Alternativen. Zögern Sie nicht, sich einen Platz in dieser herausragenden Veranstaltung zu sichern. Die Plätze sind limitiert.

Referenten: Dipl.-Inform. Petra Borowka-Gatzweiler, Dr. Simon Hoff, Dipl.-Ing. Hartmut Kell, Dr. Behrooz Moayeri, Markus Schaub, Dr.-Ing. Joachim Wetzlar
Preis: € 2.490,- zzgl. MwSt.



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

Standards sichern Interoperabilität! Sicher?

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.

Daten- und Kommunikationsnetze sind das Paradebeispiel schlechthin, wenn es um den Nachweis geht, dass Standards die Interoperabilität von Systemen und Anwendungen sicherstellen und überhaupt erst ermöglichen. Wer kennt es nicht, das OSI-Referenzmodell, das mit seinen 7 Schichten die Kommunikation von Systemen beschreibt, angefangen bei der Repräsentation von Bits und Bytes als Stromstärken oder Funkwellen bis hin zu den verschiedenen Aspekten der Kommunikation von Anwendungs-Programmen. Verschiedene Gremien präsentieren den Herstellern tausende Seiten Standards, die bei der Implementierung berücksichtigt sein wollen. Und die Hersteller schließen sich zu „Kompatibilitäts-Allianzen“ zusammen, die ihre Produkte auf Interoperabilität prüfen.

Das Ergebnis kann sich sehen lassen. So können wir uns heute beispielsweise beim Ethernet und bei der IPv4-Kommunikation darauf verlassen, dass beliebige Endgeräte miteinander kommunizieren können. Im Bereich des Wireless LAN allerdings beobachten wir immer wieder Probleme, obwohl WLAN mit einer Vielzahl von Standards der IEEE 802.11 besonders gut bestückt ist und darüber hinaus die Wi-Fi Alliance Zertifikate für erfolgreiche Interoperabilitäts-Tests vergibt. Hierüber möchte ich Ihnen heute berichten und Ihnen zwei Beispiele geben.

Beispiel „Wireless Bridges“: Ein führerloses Transportsystem (FTS) ist mit einer speicherprogrammierbaren Steuerung (SPS) bestückt. Die SPS verfügt über eine Ethernet-Schnittstelle. Man hat sie mit einem WLAN Access Point verbunden, der als Wireless Bridge betrieben wird. SPS und Wireless Bridge stellen somit aus der Sicht des WLAN ein mobiles Endgerät dar. Der Betreiber nun rüstet sein WLAN auf eine Controller-basierte Technik. Leider vermehren die FTS in der Folge immer wieder Störungen.

Die Untersuchung vor Ort offenbart das Problem. Der WLAN Controller enthält ei-



nen ARP Proxy. ARP Requests aus dem Ethernet-LAN werden nicht als Broadcasts auf der WLAN-Luftschnittstelle ausgesendet, sondern lokal vom WLAN Controller beantwortet. Das ist an sich eine gute Idee, schützt sie doch das - verglichen mit dem LAN - schmalbandige WLAN vor allzu vielen Broadcasts. Der WLAN Controller benötigt für diese Funktion die IP-Adresse des mobilen Endgeräts und legt diese in seiner so genannten Association Table ab. Dort ist Platz für genau eine IP-Adresse pro Endgerät - das FTS hat derer aber zwei! Und zwar die IP-Adresse der SPS und die des Access Point, der hier als Wireless Bridge mitfährt. Die Lösung bestand hier in der Umrüstung der FTS auf Wireless Bridges vom selben Hersteller wie die WLAN-Controller - leider ein teures Unterfangen.

Zweites Beispiel „WLAN-Knöpfe“: In einem Produktionswerk werden fehlende Teile auf Knopfdruck nachgeliefert. Der Knopf be-

findet sich an einem batteriebetriebenen Kästchen, das bei Knopfdruck über WLAN mit einer Logistik-Anwendung Verbindung aufnimmt. Hier führt die Umstellung auf ein modernes WLAN gemäß IEEE 802.11n zu Problemen. Regelmäßig melden die Knöpfe „Netzwerkfehler“.

Die Untersuchung mit dem Protokoll-analysator zeigt, dass sich in den von den Knöpfen ausgesandten Paketen immer wieder veränderte Zeichen im WLAN-Namen (SSID) einschleichen. Wir vermuten ein Speicherproblem der nun schon viele Jahre alten WLAN-Knöpfe. Die Lösung bestand letztlich darin, die Access Points derart zu konfigurieren, dass in den so genannten Management Frames weniger Informationen enthalten waren. Hersteller-support ist wegen des Alters der Knöpfe nicht mehr zu erwarten...

Alle verwendeten Komponenten trugen ein Wi-Fi-Zertifikat. Und dennoch traten Fehler auf. Zum einen, weil Funktionen verwendet wurden, die gar nicht Gegenstand der Wi-Fi-Zertifizierung sind (der ARP-Proxy des ersten Beispiels) und zum anderen, weil offensichtlich ein Programmierer die von zukünftigen WLAN-Standards erzeugte Informationsflut falsch eingeschätzt und in der Software nicht entsprechend behandelt hatte (im zweiten Beispiel). Gerade der Einsatz sich so dynamisch entwickelnder Technologien wie WLAN in vergleichsweise statischen Umgebungen wie beispielsweise in der industriellen Fertigung birgt also gewisse Risiken. Seien sie auf der Hut - auch und gerade hier gilt der alte Grundsatz: „Never change a winning team“.

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

Zweitthema

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

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.

SSD (Solid State Drive) bedeutet übersetzt Halbleiterlaufwerk. Das ist grob irreführend, denn die wesentliche kennzeichnende Eigenschaft ist ja, dass grade nichts mehr „läuft“. Der Begriff hat sich vor allem dadurch eingebürgert, dass SSDs primär für den Consumer Markt angeboten wurden, die in Bauform und Schnittstellen den Plattenlaufwerken, die durch sie ersetzt werden sollen, genau entsprechen. Dadurch kann z.B. der PC-Kunde wählen, ob er zur Datenspeicherung ein Laufwerk konventioneller Bauart mit beweglichen Teilen haben möchte oder eben einen Speicher, der komplett auf integrierter Halbleitertechnologie basiert.

Mittlerweile sind die SSDs aber den Kinderschuhen entwachsen und haben sowohl von Kapazität als auch von ihrer Stabilität (nicht nur im mechanischen Sinne!) eine hohe Quantität hinsichtlich der auf ihnen zu lagernden Daten als auch eine Qualität erreicht, die sie für professionelle Anwendungen mit höchsten Anforderungen im RZ nutzbar macht.

Anwendungsbeispiele

SSDs beginnen damit, Plattenlaufwerke bei Schlüsselanwendungen zu ersetzen. Würden sie sich auch genauso verhalten wie konventionelle Platten, könnte das dem Netzwerker eigentlich völlig gleichgültig sein. Es ist aber so, dass der Reiz der neuen Technologie primär in der teilweise enormen Geschwindigkeitssteigerung bei Schreib- und Lesezugriffen liegt. Sind die SSDs direkt in dem Server, wo die Speicherleistung benötigt wird, montiert, kommunizieren sie mit den Prozessoren über den internen Rechnerbus.

Bei der Virtualisierung ist es aber so, dass die Menge der gegebenenfalls wandernden virtuellen Maschinen freizügig auf die

Menge der Speicherressourcen zugreifen können sollen. Das bedeutet, dass eine Bindung eines Speichersystems an einen oder mehrere Multicore-Prozessoren nicht mehr hinreicht, um diese Anforderung zu bedienen. Vielmehr muss es ein zentralisiertes Speichersystem geben, auf das die Menge der Prozessoren und somit mittelbar die Menge der virtuellen Maschinen systematisch zugreifen kann. Dabei darf es keine Unterschiede geben, die auf die Bindung einer VM an einen spezifischen Prozessor zurückzuführen wäre. (siehe Abbildung 1)

Somit hat das Netz die Rolle des Systembusses für die verteilte virtualisierte Umgebung. Das kann nur dann sinnvoll funktionieren, wenn die Leistung des Netzes harmonisch zu der Leistung der angeschlossenen Komponenten, zu denen auch die neuen SSD-Speicher gehören, ist, und zwar nicht nur hinsichtlich der real erreichbaren Übertragungsrate, sondern vor allem hinsichtlich der Latenz.

Ein völlig anderes Beispiel: wie Anbieter von Broadcast-Diensten in zukünftigen Marktumfeldern überleben können, hängt nicht nur von ihrer Fähigkeit ab, einzigartige digitale Inhalte zu erstellen, zu aggregieren und eine bekannte Marke daraus zu machen, sondern auch davon, welche strategischen Entscheidungen sie darüber treffen, wie und wo ihre Inhalte gespeichert werden sollen.

Die geforderten geschäftlichen und kreativen Entscheidungen müssen von den Application Service Providern durch eine ganze Anzahl integrierter Funktionen unterstützt werden, die unter anderen die Bezahlung von Dienstleistungen, Anzeigengeschäft und Abo-Modelle entsprechend problemlos unterstützen. Kann der Besitzer von Inhalten diese Funktionen nicht ökonomisch und schnell umsetzen,

macht das eben früher oder später der Mitbewerber.

Dies kann man auch auf die Anbieter von Cloud-Dienstleistungen in den bekannten unterschiedlichen Formen (IaaS, PaaS, SaaS) oder auf Finanzdienstleister übertragen.

In all diesen Fällen ist die Schnelligkeit von Antwortzeiten existentiell. Unabhängig von den Optimierungs- und Zugriffsmöglichkeiten eines Datenbanksystems bedeutet das vor allem, dass die Speicher als solche schnell sein müssen.

Physikalische Begrenzungen traditioneller Speichersysteme

Viele etablierte Hersteller von Großrechnern haben anfangs die neue Mikroprozessortechnologie ignoriert und angenommen, dass das Ganze nur eine kleinere, bescheidenere Ausführung von Dingen sei, die sie bis jetzt auch gemacht haben. Die gesamte PC Revolution ist aber dadurch entstanden, dass Tausende in der Schaltkreistechnologie bewanderte Ingenieure, die vorher nicht bei den Rechnerherstellern gearbeitet haben, das völlig anders gesehen und eine ganz neue Art von Produkt kreiert haben.

Vor 1990 gab es nur im wissenschaftlichen Bereich oder bei der Digitalen Signalverarbeitung Prozessor-Arrays. Die Einführung von Multiprozessor-basierten Workstations durch die Firma Sun Microsystems in 1992 erlaubte im Zusammenwirken mit dafür eigens entwickelten Unix-Versionen den Einsatz des Multiprocessings für allgemeine Geschäftszwecke.

Der nächste Schritt war die Möglichkeit, in einem Mikroprozessor multiple Prozessoren unterzubringen, also integrierte Multiprozessoren einzuführen. Dies wurde vor

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

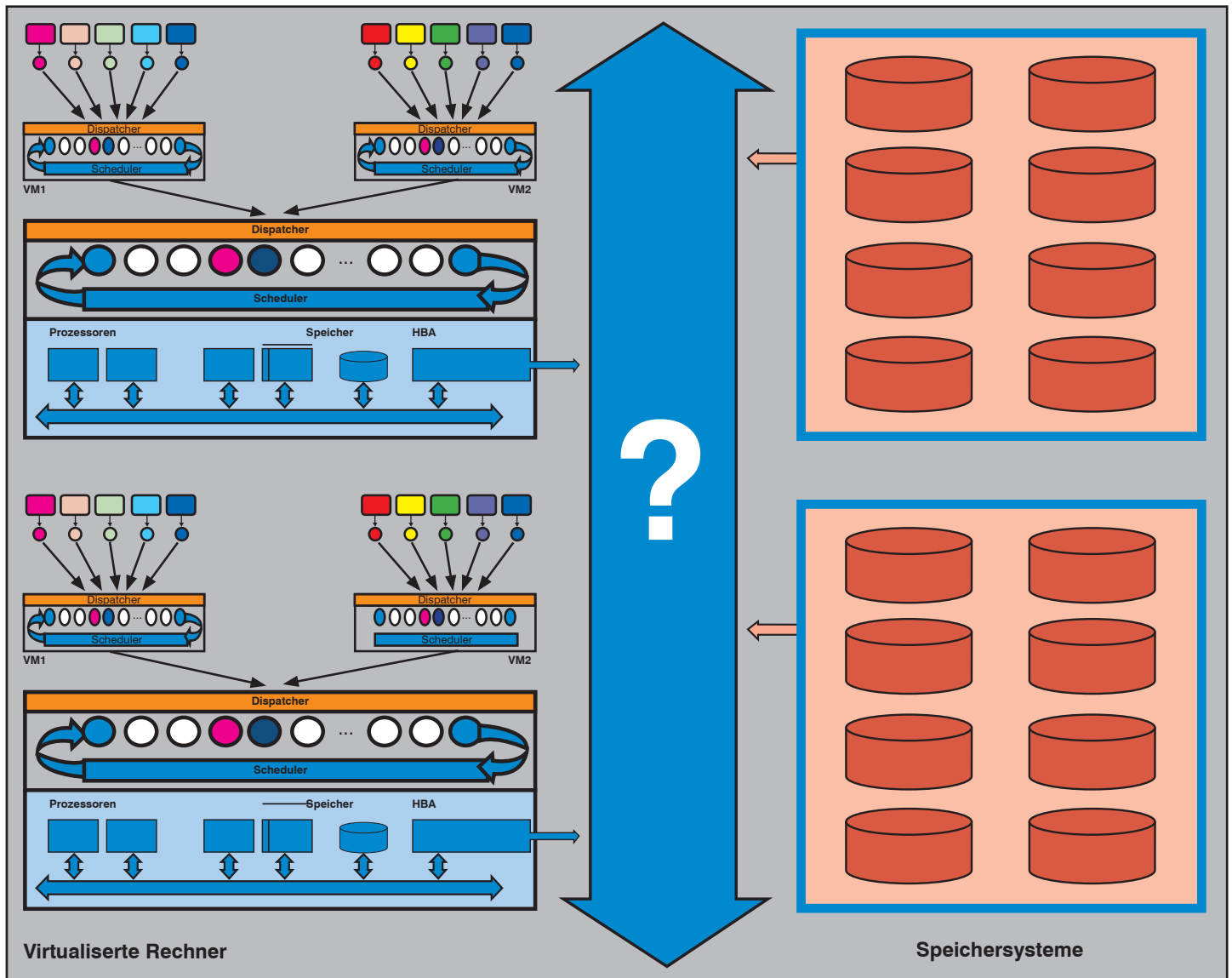


Abbildung 1: Virtuelle Server: Netz wird zum Systembus

allein von Intel und AMD vorangetrieben. Im privaten Bereich hat dies die Wandlung des Internets von einem relativ primitiven Kommunikationsmedium zu einem integrierten Träger multimedialer, breitbandiger Unterhaltungs- und Informationsangebote begünstigt, im professionellen Umfeld zur Virtualisierungswelle geführt: Hunderte bisherige kleine Server mit all ihren täglichen Betriebsproblemen konnten auf Virtuelle Maschinen konsolidiert werden, die ihrerseits auf wenigen leistungsfähigen Multicore-Servern laufen, was bei richtiger Anwendung zu enormen Kostenvorteilen führen kann.

Leider können die magnetenaufzeichnungsorientierten Datenspeicher nicht auf eine derartig fulminante Entwicklung zurückblicken. Plattensysteme sind letztlich in den 90er Jahren stehen geblieben. Man konnte billige und teure bauen, relativ lang-

sam und relativ schnelle, sie einfach betreiben oder zu sehr komfortablen Speichersystemen zusammenfassen, aber insgesamt blieben sie wegen der Eigenheiten der in ihnen enthaltenen beweglichen Teile hinter der Gesamtentwicklung wesentlich zurück.

Im Jahrzehnt zwischen 2000 und 2010 konnte die Kapazität einer Hard Disk von 180 GB auf 3 TB erhöht werden, also um den Faktor 17. In 1999 kostete 1 TB ca. 80.000 US\$, Ende 2010 kostete ein 3 TB Plattenlaufwerk um 250 US\$. Das ist eine Preisreduktion um den Faktor 1000. Plattenlaufwerke sind zwar immer billiger geworden, aber weder zuverlässiger noch schneller. Zugriffszeiten oder Ersetzungszyklen sind über die Jahrzehnte relativ gleich geblieben. Bei sehr großen Installationen kann man konsequenterweise beobachten, dass die MTBF mit zunehmen-

der Kapazität immer weiter gesunken ist, weil die Zuverlässigkeit der Plattenspeicher nicht mit ihrer Größe gewachsen ist. Natürlich haben sich die Hersteller großer Speicherlösungen wie IBM, HP, EMC oder Hitachi extrem viele Zusatzfunktionen einfallen lassen, um die Daten möglichst oft zu duplizieren und nach komplizierten Mustern zu verteilen. Die bekannten RAID-Level sind hier nur die Basis wesentlich komplexeren Verfahren. Langfristig werden sie sich aber nicht wirklich erfolgreich gegen das Problem stemmen können, dass bei weiterem erheblichen Wachstum der Kapazität mit zufällig verteilten, nicht mehr zu korrigierenden Datenverlusten zu rechnen ist.

Alle Begrenzungen resultieren letztlich primär aus den beweglichen Teilen in Hard Disk Laufwerken. Die beweglichen Teile, insbesondere die Köpfe, unterliegen

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

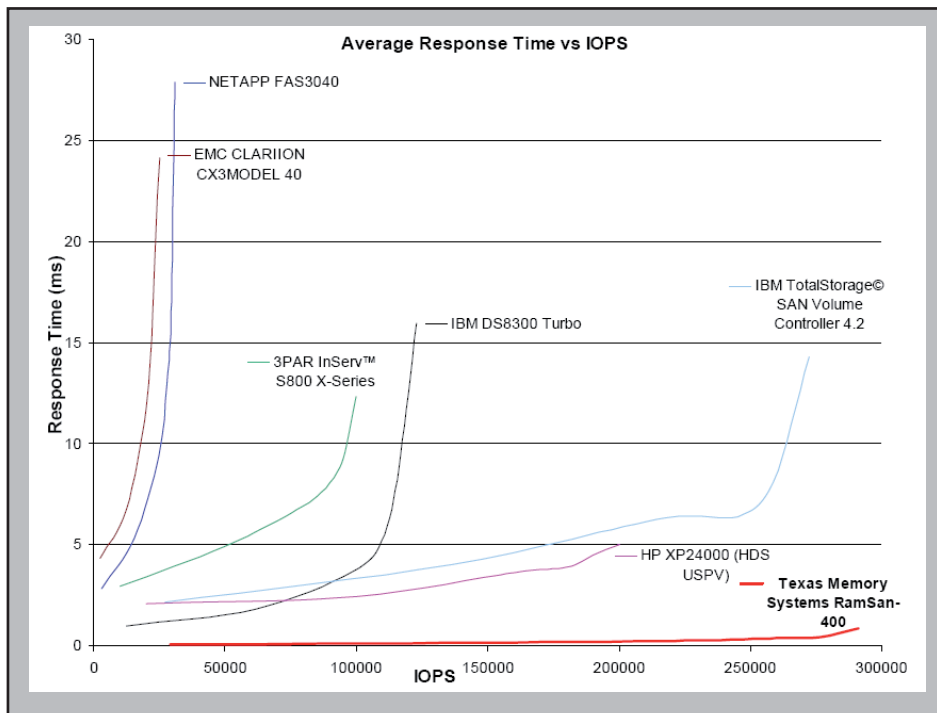


Abbildung 2: Grenzen von Plattenspeichersystemen

physikalischen Randbedingungen. Da dies nun sehr wichtig für das Gesamtverständnis ist, möchte ich es anschaulich erklären. Um die technische Zugriffsgeschwindigkeit einer Festplatte wirklich zu erhöhen, muss man die Köpfe beschleunigen. Jeder weiß aber, dass der Aufwand für die Beschleunigung eines Körpers mit einer definierten Masse exponentiell steigt. Ein normales Auto mit 100 PS kann vielleicht auf 200 km/h beschleunigt werden. Möchte man das gleiche Auto auf 400 km/h beschleunigen, werden 200 PS längst nicht ausreichen, sondern man würde ca. 800 PS benötigen. Eine andere Möglichkeit wäre es, das Auto leichter zu machen, wie es in der F1 geschieht. Aber Motor und Energieaufwand werden ja immer größer. Ein Auto mit einem Motor von 2000 PS würde so schwer, dass auch diese Leistung nicht mehr ausreichen würde, es noch wesentlich voranzubringen. Möchte man schneller reisen, muss man offensichtlich die physikalischen Randbedingungen erheblich ändern, wie es beim Flugzeug geschieht. In dem Fall wird mit steigender Höhe der Luftwiderstand, der das Auto primär aufhält, immer weiter gesenkt. Die optimale Leistungsausbeute erhält man bei einem ballistischen Flug. Ganz dumm ist dann aber, dass der Kopf, den wir aufwändig beschleunigt haben, vor der Spur, die er bearbeiten oder lesen soll, wieder bremsen muss. Natürlich steht auch die notwendige Bremsenergie wieder in exponentiellem Verhältnis zur Geschwindigkeit, von der aus gebremst werden muss und zum Bremsweg. Letztlich ist die Arbeitsweise einer Festplatte

te vergleichbar einem Ampelrennen (was der gelbe Ferrari aus der Nachbarschaft meist gewinnt).

Ein weiteres Problem von Speichersystemen, die auf magnetisierbaren Medien basieren, ist die für eine haltbare Magnetisierung notwendige Zeitspanne. Neben weiteren, eher organisatorischen Problemen, trägt sie zur Differenz zwischen Schreib- und Lesezeiten bei.

Noch eine Nebenbemerkung zu optischen Speichern. Sie fallen aus dieser Betrachtung heraus, weil sie von ihren Reaktionszeiten eher als Langzeitspeicher gedacht sind. Aber auch dort versagen sie, wenn es um wirklich langfristige Speicherung geht. Das primäre Funktionsprinzip optischer Plattenspeicher wie CD oder DVD ist das Verbiegen einer Folie durch einen Laserstrahl. Diese Folie hat aber mit den Jahrzehnten die Tendenz, sich bei sinkenden Temperaturen wieder „glattziehen“ und bei hohen zusätzlichen Verwerfungen zu bilden. Beides ist ungünstig für den gespeicherten Informationsinhalt.

In Abbildung 2 sieht man die Grenzen moderner Plattenspeichersysteme. Ab einer gewissen Zugriffsrates steigt die Reaktionszeit produktbedingt rapide. Die Abbildung 2 zeigt aber auch direkt, wie sich ein SSD-Speicher modernster Bauart vergleichsweise verhält (rote Linie).

Möchte man also wirklich vorankommen, benötigt man eine völlig andersartige Art

der Datenspeicherung. SSD-Speicher sind einfach gesagt genau das Medium, welches zu dichten Multicore-Prozessoren (und damit zur Virtualisierung) wirklich passt. Um den Unterschied direkt zu klären: RAM, wie er immer nahe eines Prozessors vorkommt und SSD-Speicher basieren prinzipiell auf den gleichen Technologien, haben aber andere Aufgaben. Der RAM stellt eine Menge von Registern dar, die mittels begleitender abstrahierender Techniken den Virtuellen Speicher eines Prozessors implementieren. Der Prozessor greift mittels entsprechender Befehle auf diese Register unmittelbar zu. Alle Daten sind im Wortformat des Prozessors organisiert. Das bedeutet, dass der RAM genau zu dem Prozessor passen muss, der ihn benutzt. Prozessor und RAM sind die wesentlichen Teile der CPU. Ein SSD-Speicher hat die Aufgabe, Daten in entsprechenden allgemeinen Formaten so aufzunehmen und bereitzustellen, dass sie auch von unterschiedlichen CPUs benutzt werden können. Der SSD-Speicher ist daher ein peripheres Gerät, welches mit einem entsprechenden Treiber benutzt werden soll.

Der Hersteller Texas Memory Systems nennt seine Systeme RamSan®. Damit weist er darauf hin, dass die verwendete Technologie eigentlich einem RAM entspricht und auch dessen Grundqualitäten, wie eine enorme Reaktionsgeschwindigkeit, besitzt. Andererseits besitzt das Gerät aber zusätzliche organisatorische Funktionen, so dass es zum Bestandteil eines SANs werden kann.

Marktforschungsunternehmen rechnen damit, dass sich der Anteil von Hard Disk Speichern an der Datenhaltung in Unternehmen in den nächsten fünf Jahren halbieren wird und diese in 2019 völlig verschwunden sein werden. Danach wird es sie nur noch in Consumer-Umgebungen mit niedrigen Ansprüchen geben.

Sogar im Consumer Markt haben wir bei den Disketten ja schon gesehen, dass beliebte und extrem verbreitete Speichermedien und –Systeme wirklich vollständig verschwinden können.

Vom Flip Flop zum SSD-Speicher

Sie können den Begriff „SSD“ natürlich in Wikipedia nachschauen. Dort erhalten Sie vor allem über mögliche Organisationsformen mehr Informationen, als Ihnen lieb ist, der eigentliche Kern der Sache ist aber viel einfacher.

SSD-Speicher gibt es schon seit 40 Jahren, nur hatten sie anfangs nicht diesen Namen. In den 70er Jahren kamen die ersten integrierten Schaltungen für den

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

breiten Markt auf. Besonders berühmt war damals die „SN“-Serie in TTL-Technik (Transistor-zu-Transistor-Logik). Ein SN 7400 N enthält vier NAND-Gatter. Je zwei NAND-Gatter kann man zu einem Flip Flop zusammenschalten. Das ist ein bistabiler Multivibrator. Zu Beginn habe er einen Zustand, z.B. repräsentiert durch eine „Null“ an seinem oberen, eine „Eins“ an seinem unteren Ausgang. Schickt man auf seinen (in der Zeichnung oberen) Eingang eine „Eins“ auf den oberen und eine „Null“ auf den unteren Eingang, wechselt er nicht nur seinen Zustand, sondern behält ihn auch solange bei, bis man das invertierte Signal an seine Eingänge schickt und vice versa. Daher bistabil, die Schaltung hat zwei stabile Zustände. Den Zusatz „Multi“ verdient sich die Schaltung dadurch, dass man den Vorgang beliebig oft wiederholen kann. Führt man eine geeignete Rückkopplung ein, kann man die Schaltung auch zum Schwingen bringen, daher „Vibrator“. Obwohl man damit viele wunderschöne Sachen anstellen kann, z.B. einen Peilsender mit nur einem SN7400 bauen, ist dieser Teil hier uninteressant.

Vielmehr kann man den Flip Flop zur Datenspeicherung nutzen. Ein Flip Flop speichert nur ein einziges Bit, aber das macht nichts, wie wir gleich sehen. Schaltet man Flip Flops hintereinander, ergeben sich Schieberegister, die man normalerweise taktet. Baut man jetzt so viele Flip Flops zu einem Schieberegister, wie man braucht, um ein Wort zu speichern, kann man genau das tun. Der primäre Unterschied zur Speicherung auf einer Platte oder auf einem Band ist aber, dass sich das Wort bewegt, wenn man die Taktung nicht wegnimmt. Koppelt man vom letzten auf den ersten Flip Flop zurück, läuft es immer im Kreis.

In Abbildung 3 sehen wir die Entwicklung der SSD-Speicher in den letzten 40 Jahren. Ausgehend von einem Flip Flop hat sich durch die Steigerung der Integrationsdichte mittlerweile die Möglichkeit ergeben, 64 Gbyte in einem Chip der Größe einer Briefmarke zu speichern. Wir erklären jetzt diese Entwicklung.

Höhepunkte der frühen 70er Jahre war das 8 Bit SN 74164 N Schieberegister, das ging immerhin schon ein ganzes Byte herein, und mein Vortrag über die SN-Serie im Physikunterricht der Mittelstufe.

Tatsächlich kann man aber die wesentlichen Vor- und Nachteile der SSD-Technik schon an den geschilderten Schaltungen ausmachen:

- **Schnelligkeit.** Sie ist einzig von der

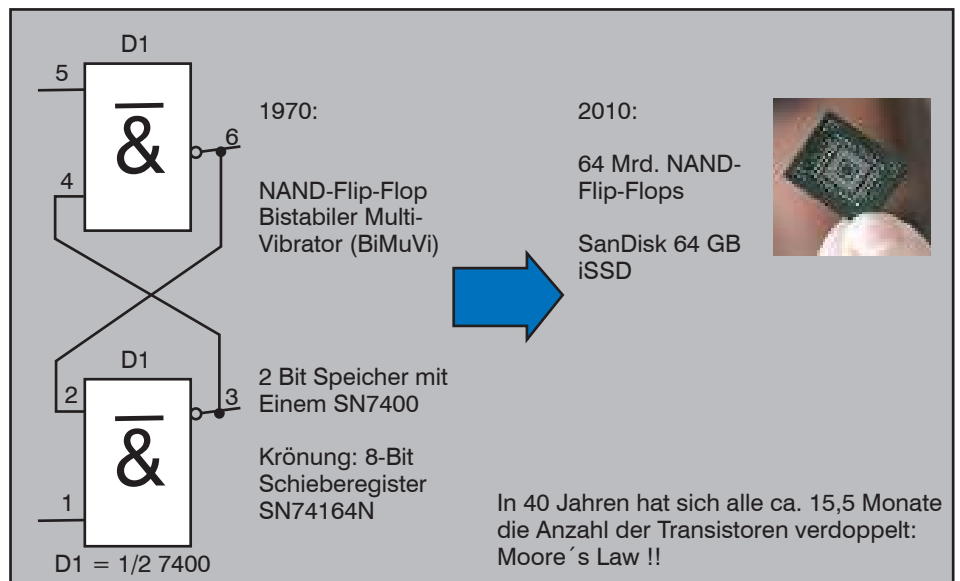


Abbildung 3: SSD-Entwicklung

Taktrate der Logikfamilie limitiert. Bei CMOS-VLSI bedeutet das, dass ein SSD-Speicher, der mit den gleichen Randbedingungen aufgebaut ist wie ein Prozessor aus der gleichen Logikfamilie, auch genauso schnell arbeiten kann. Macht man also bei der Verbindung zwischen Prozessor und SSD-Speicher keine großen Fehler, arbeiten diese Systeme sehr harmonisch. Dieser Vorzug lässt sich mit Platten nicht erreichen.

- **Skalierbarkeit.** Die Größe eines SSD-Speichers ist letztlich primär von der Anzahl der verfügbaren Flip Flops limitiert. Eine weitere Rolle spielen natürlich die Wege zwischen den Komponenten des SSD-Speichers. Sind diese zu lang, ergeben sich zu große Laufzeitverzögerungen. Theoretisch könnte man einen SSD-Speicher aus beliebig vielen SN 7400 N aufbauen. Würde man eines der modernen Fußballstadien, wie die Allianz Arena oder die Arena auf Schalke komplett bis zum oberen Rand mit SN 7400 füllen, ergäbe das einen SSD-Speicher von 4-5 GB Kapazität, wie ich extra für diesen Artikel ausgerechnet habe. Abgesehen von Leitungslängen, die ein Vielfaches an Verzögerung gegenüber den Schaltmöglichkeiten der Flip Flops aufweisen, wäre das ziemlich unpraktisch. Wir kriegen den SSD-Speicher gleich aber schon noch klein.

- **Fällt der Strom aus,** sind die Daten ohne Wenn und Aber weg.

Durch die Fortschritte der Technologie und die damit verbundene Miniaturisierung kamen zunächst noch zwei Probleme hinzu.

- **Empfindlichkeit gegenüber Schwankungen** im Herstellungsprozess. Je dichter eine integrierte Schaltung werden soll, desto höhere Anforderungen muss man an die Herstellungstechnik stellen. Die TTL-SN-Serie war sehr grob, ein Transistor war nicht wesentlich kleiner als das, was man auch in ein Gehäuse gesetzt hätte, um einen individuellen Transistor zu fertigen. Heute sprechen wir von SSD-Chips in 65, 40 oder sogar neuerdings 25 (!!!) nm-Technik. Das bedeutet, dass die Leitungen wirklich nur diese Dimension annehmen und auch die Elemente selbst bzw. ihre Repräsentierungen durch hauchfeine dotierte Schichten entsprechend winzig geworden sind. Es dauert ein paar Jahre zwischen dem Zeitpunkt, wo ein Hersteller Chips in einer Technologie ankündigt und herstellt und dem Zeitpunkt, wo auch wirklich fast alle Chips genau gleich sind und höchste Qualitätskriterien erfüllen. In den letzten Jahren konnten aus diesem Grund SSD-Speicher für professionelle Anwendungen wirklich nur aus (ohne Quatsch) „handverlesenen“ Chips hergestellt werden, was die Sache natürlich erheblich verteuert hat.
- **Leckströme.** Damit hat jede wirklich hochdichte integrierte Schaltung zu kämpfen. Mit der Zeit können sich, anschaulich gesprochen, durch parallel laufende Leitungen, die als Kondensator wirken, Ladungen aufbauen. Sind sie groß genug, neigen sie zum gegenseitigen Potentialausgleich und erzeugen einen Stromfluss, der im Falle eines SSDs kontraproduktiv zu den Strömen stehen kann, die die zu repräsentierenden Daten beinhalten können.

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

Ein weiterer wichtiger Begriff im Kontext der SSDs ist der FLASH-Speicher. Dabei handelt es sich ebenfalls um einen Halbleiterspeicher, der aber auf einem etwas anderen Arbeitsprinzip als der Flip Flop basiert. Die genaue Bezeichnung wäre FLASH-EEPROM (Electrical Erasable Programmable Read Only Memory). Er besteht aus seiner Matrix von Feldeffekt-Transistoren. Jeder Transistor repräsentiert ein Bit. Beim Programmierungsvorgang wird durch ein so genanntes Floating Gate dem Transistor eine elektrische Ladung geschickt, die ihn sperrt. Die Änderung des Ladezustandes ist bei dieser speziellen Bauform nur durch die Nutzung des quantenmechanischen Tunnel-effektes möglich, da eine Ladung eine isolierte Zone durchqueren muss. Das hat den Vorteil, dass der angenommene Zustand im Gegensatz zum Flip Flop nicht mehr davon abhängig ist, ob eine externe Betriebsspannung anliegt. Durch Konstruktion und Betriebsweise ergeben sich folgende Vor- und Nachteile:

- + Informationsspeicherung ohne externe Energiezufuhr möglich
- + extrem kompakte Bauweise
- + sehr preiswert

- Schreib- und Lesevorgänge nur in größeren Blöcken möglich
- Schreibvorgang dauert länger als Lesevorgang
- Anzahl der möglichen Schreibvorgänge begrenzt (heute ca. 2 Mio. für NAND-Flash)

Jeder Leser hat bestimmt eine ganze Anzahl von Flash-Speichern direkt um sich herum, mindestens USB-Sticks und Speicherkarten für Digitalkameras und Smart Phones, ohne jemals darüber nachgedacht zu haben, welche kleinen Wunderwerke sie darstellen. Wir können sie nutzen, ohne es zu merken. Genau das ist der Trick.

Enterprise SSDs heute

Durch die lange Entwicklungszeit haben sich viele unterschiedliche Erscheinungsformen der SSDs gebildet. Der Speichertyp kann hinsichtlich seiner Grundorganisation als Flash, RAM oder gemischt gekennzeichnet werden. Die Darstellungen im letzten Absatz sollten dazu gedient haben, klar zu machen, dass eine Mischung der Technologien RAM und Flash genau zu dem wünschenswerten Ergebnis führt, die jeder Technologie an sich haftenden Nachteile weitestgehend zu kompensieren.

Ergänzt man Flip Flop-Felder um Flash-EEPROMs, können die Daten auf letzte-

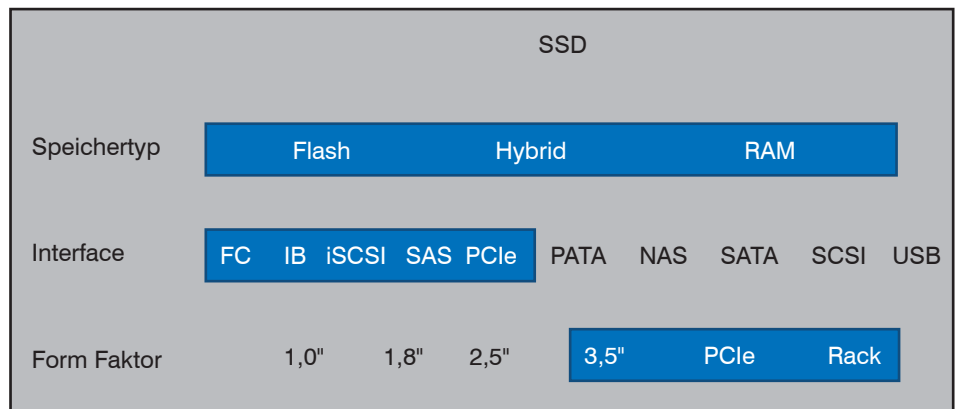


Abbildung 4: SSD Enterprise Formate

ren auch ohne externe Energieversorgung gespeichert werden. RAMs können die Daten schnell und flexibel von außen annehmen oder abgeben und auf die beim Flash benötigten Blöcke konvertieren. Letztlich bilden RAM und Flash eine interne zweistufige Speicherhierarchie innerhalb eines SSD-Speichers. Es wird also eine zusätzliche Kontroll-Logik in einem speziellen Steuerprozessor benötigt.

Und genau dieser Schritt ebnet den Weg zu SSD-Speichern für professionelle Anwendungen im Corporate Umfeld. Dafür ist es aber von untergeordnetem Belang, wie genau diese unterschiedlichen Organisationsformen von den Herstellern dazu benutzt werden, Optimierungen in der Arbeitsweise durchzuführen. Letztlich zählen nur die Ergebnisse.

Bei den Schnittstellen gibt es eine breite Vielfalt, die sich auch daraus erklärt, dass vor allem zu Beginn der Entwicklung eher der Consumer Markt im Vordergrund und hier der Ersatz von Festplattenlaufwerken durch bau- und kapazitätsgleiche SSDs mit genau den Schnittstellen, die man in einem PC findet, stand. Wo wir einmal diesen Bereich ansprechen, kann man sagen, dass viele Anwender, die sich ein solches Laufwerk zugelegt haben, bitter enttäuscht waren, weil sich die erhoffte Geschwindigkeitssteigerung nicht eingestellt hat. Das liegt aber genau betrachtet nicht an den SSDs, den Schnittstellen, der Bauform oder der Datenübertragung, sondern schlicht daran, dass die in PCs üblicherweise verbauten sonstigen Elemente, besonders die Prozessoren, von der möglichen Steigerung der Geschwindigkeit nicht profitieren konnten. Neben den Prozessoren gab es auch Probleme mit den Treibern. Lediglich Apple hat es zeitnah geschafft, in den besonders schlanken Notebook-Bauformen nicht nur von dem geringeren möglichen Formfaktor, sondern auch aufgrund des anders konstruierten Betriebssystems von der Verbesserung der Zugriffsgeschwindigkeit

zu profitieren. Bei PCs konventioneller Bauart müssen erst noch ca. 3-5 Produktionszyklen vergehen, bis auch hier die gewünschten und möglichen Vorzüge der SSDs für den Benutzer spürbar werden. Die eben schon zitierten Marktforscher haben sich hier gründlich vertan. Für 2010 hatten sie erhebliche Steigerungen der SSD-Verwendung bei PCs und kaum Steigerungen bei der Verwendung im Enterprise-Umfeld vorausgesagt. Es kam genau anders. Wegen der genannten Probleme blieben die SSDs bei PCs weit hinter den Erwartungen zurück und im Enterprise Umfeld zogen sie massiv an, weil es die genannten Probleme dort kaum gibt und praktisch jede moderne Serverbauart sofort vom Einsatz der SSDs massiv profitieren kann.

Schließlich gibt es SSD-Speicher in den unterschiedlichsten Formfaktoren. Die Abbildung 4 unterlegt, welche Schnittstellen und Formfaktoren im RZ in Frage kommen.

Was können Enterprise-SSDs Mitte 2011? Sehen wir uns Beispiele an: SanDisk hat vor kurzem eine Familie von SSD-Chips vorgestellt, deren Spitzenmodell eine Speicherkapazität von 64 GB hat. Vergleichsweise müssten wir ca. 14 große Fußballstadien mit SN 7400 füllen, um die gleiche Kapazität zu erhalten. Dieser Chip ist schon in Abbildung 3 zu sehen. SanDisk ist nur einer der Hersteller, die die Probleme der Leckströme und der Produktvarianz gelöst haben. Ein weiterer Spezialhersteller wäre SandForce. Die Abbildung 5 zeigt ein Diagramm des Kontrollprozessors, der eigentliche Speicher wird über die Schnittstellen rechts erreicht.

Intel stellt zwar auch Chips her, verkauft aber lieber fertige SSDs. Die Ende März angekündigten Modelle der neuesten Generation (Elmcrest) zeigen in beeindruckender Weise, dass SSDs zu höchster Leistung bei geringem Preis befähigt sind.

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

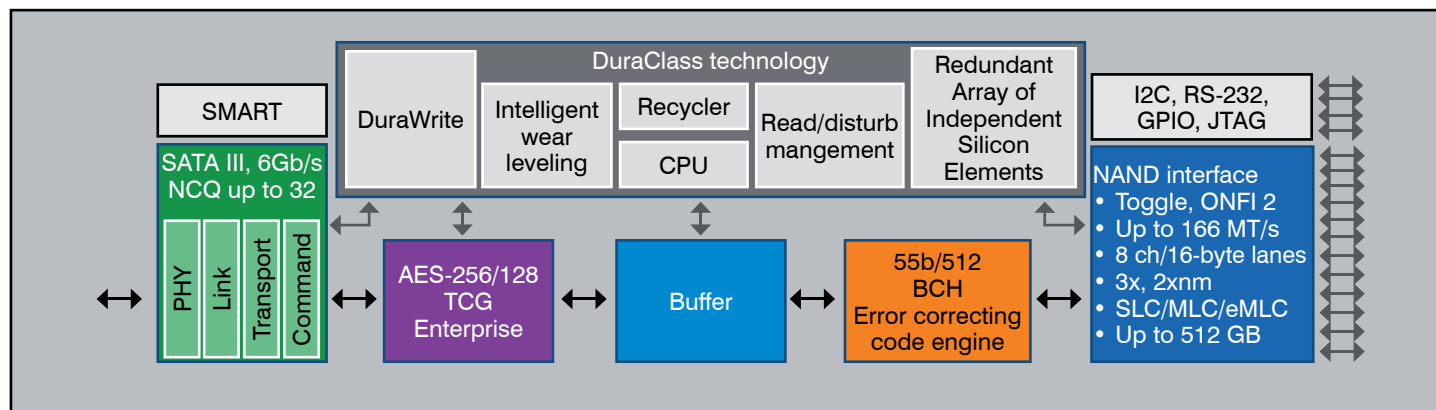


Abbildung 5: SandForce SSD-Prozessor

The World's Fastest Storage®

512 Gbyte

RAM SSD

Flash-Module für Persistenz

Der Alptraum für Netze: 24-36 Gbps, Latenz min. 30 µsek (!), operationell 100 - 300 µsek.

RamSan-440
Texas Memory Systems, Inc.

4 x 4/8 GbFC

oder

4 x 10 G IB

Abbildung 6: Texas Memory Systems RamSan 440

Die 2,5" 80 Gramm 250 GB-SSD erreicht im Test von CHIPonline.de bei mittlerer Datenrate fast 480 MByte/s für das Lesen und 320 MByte/s für das Schreiben. Die Latenz liegt dabei immer unter 100 µsek. Um diese Geschwindigkeit an den Prozessor zu bringen, benötigt man einen sechsfachen SATA-Anschluss auf dem Motherboard. Als eine der schnellsten herkömmlichen Festplatten erreicht die 10.000 U/min. Velociraptor von Western Digital grade mal 121/142 MByte/s. Lesen/Schreiben mit einer Latenz von 3 – 20 MILLISekunden. Die Intel SSD kostet 481,90 €, die 300 GB Velociraptor kostet ausstattungsabhängig 99,- € bis 163,- €. Nimmt man jetzt also bei gleichbleibender Kapazität Formfaktor und Geschwindigkeit als Maßstab für die Kosten, liegt die SSD vorne. Es kommt also ganz darauf an, was man möchte.

Nun, das war jetzt eine kleine SSD für die gehobene Privatanwendung. Sie wird aber nachher noch wichtig für die Argumentation.

Ein bekanntes großes System ist der RamSan von Texas Memory Systems. Das Modell 440 hat 512 GB-Kapazität und wohnt in einem Rack-Gehäuse. Die Latenz liegt minimal bei 30 µsek. bei norma-

ler Operation zwischen 100 und 300 µsek. Man kann sie mit 40 G InfiniBand oder 4x8 GbFC an ein Netz anschließen. Das Modell 12000 in Rackgröße mit 100 TB Kapazität für bescheidene 4 Mio. US\$ hat vergleichbare Leistungsdaten. In Gewicht und Formfaktor liegt bei gleicher Kapazität zwischen der RamSan 440 und der genannten Intel-SSD ein Faktor 100. In Reaktionsgeschwindigkeit sowie beim Lesen und Schreiben steht die kleine billige Intel SSD dem wesentlich größeren Gerät nur

bei extremen Anforderungen nach, die viele Unternehmen gar nicht haben.

Deshalb hat Fusion I/O einen anderen Weg gewählt. Der Hersteller bietet eine Reihe von SSD-Karten an, die zur Bildung eines größeren Systems einfach nebeneinander in ein passendes Chassis gesteckt werden. Ein Beispielsystem schafft dann z.B. 5 TB Gesamtkapazität bei 6 Gbyte/s (48 Gbps) Schreib- und Lesegeschwindigkeit und bis zu 800.000 I/O-Operationen pro Sekunde. (siehe Abbildung 7)

Es ist natürlich nur eine Frage der Zeit, bis ein Hersteller auf den Gedanken kommt, einfach einen Stoß der oben genannten Intel SSDs zusammen mit einem Steuer- und I/O-Prozessor in ein Gehäuse zu packen und zum Kampfpfeis anzubieten.

IBM geht für die Server der Serien x3850 und x3950 bereits einen vergleichbaren Weg. Im Angebot sind unterschiedliche PCIe-SSD-Karten (Adapter genannt). Kombiniert man z.B. 14 von diesen Adaptern, erreicht das Gesamtsystem eine I/O-Leistung von über 10 Gbyte/s (über 80 Gbit/s!). Mit nur 4 Adaptern erreicht man schon 677.000 IOPS, alles gemessen mit dem IOMETER nach Standard. Bis

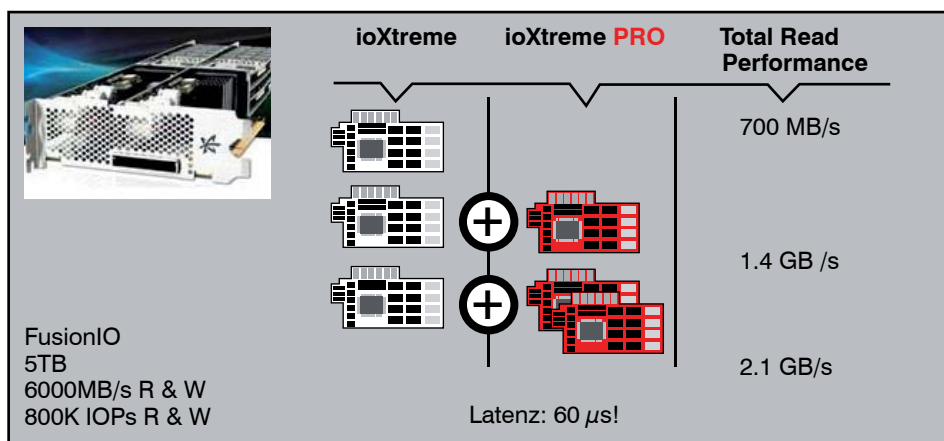


Abbildung 7: Fusion I/O

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

mich IBM korrigiert, bin ich aber der Ansicht, dass das auch die Fusion i/O-Karten sind. Ein Alleinstellungsmerkmal liegt hier eher in der Verwaltung der Platten. Es ist von einem Anwender in der Enterprise Umgebung nicht zu erwarten, dass er die Zuordnung zwischen den Daten und den unterschiedlichen Speichersystemen sozusagen von Hand vornimmt. Wesentlich günstiger ist da doch schon eine Software, die die Zuordnung entlang von Regeln, die im Rahmen von QoS-Anforderungen definiert wurden, vornimmt und dabei auch die für die jeweiligen Speichersysteme passenden zusätzlichen Sicherungsfunktionen vornimmt. Genau das macht der IBM SAN Volume Controller.

EMC ist einer der Pioniere auf dem Bereich SSD. Sie sind fester Bestandteil des Symmetrix Vmax-Systems. Nach Herstellerangaben kann man mit einem SSD-Device ca. 30 FC-Platten ersetzen und 98% Strom sparen. Der Kunde braucht sich eigentlich um technische Eigenschaften nicht zu kümmern, alle Elemente der Storage-Lösung stehen unter einer gemeinschaftlichen Verwaltung.

Um diesen Abschnitt abzuschließen, kommen wir noch zu einem wirklich installierten Beispiel. Die Deutsche Bank hat mit HP, Mellanox/Voltaire und Fusion I/O ein System mit folgenden Leistungsdaten aufgebaut:

- über 6 TB Kapazität pro Storage Server mit Fusion I/O
- über 23 GB/s (184 Gbps) Speicher-durchsatz
- 2,5 M Random IOPS
- Zugriffszeit über 100-mal schneller als bei herkömmlichen Systemen

Der Anschluss zum Server erfolgt mit nur 8(!) 40Gb(!) Infiniband-Adaptoren über entsprechende Voltaire-Switches. Neben der bescheidenen und zurückhaltenden Technologie kam noch zusätzlich eine spezielle Storage Accelerator Software von Voltaire zum Einsatz.

Dieses Beispiel zeigt, was Ende 2010 möglich war. Ende 2012 wird in jeder Dimension mindestens das Doppelte möglich sein.

Konsequenzen

Die Konsequenzen sind erfreulich und erschreckend zugleich, je nachdem, wo man steht.

„Speicher“ ist ein sehr individuelles The-

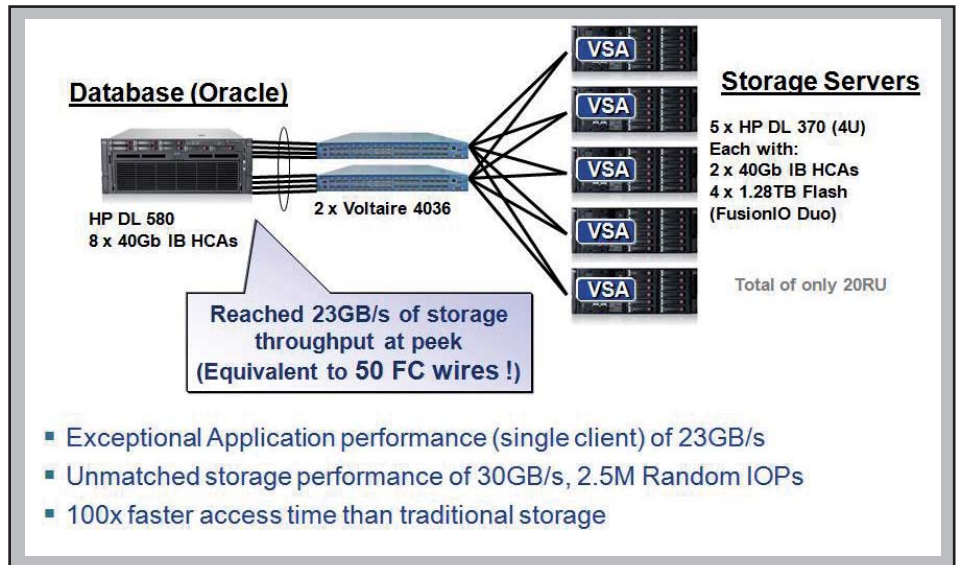


Abbildung 8: Hier hängt der Hammer aktuell

Quelle: Mellanox

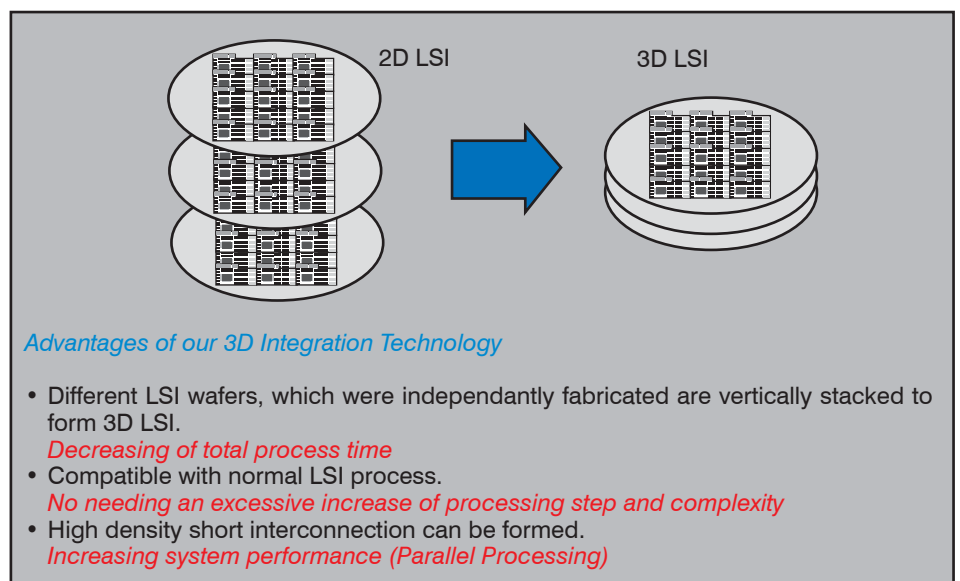


Abbildung 9: 3D-SSD

Quelle: Koyanagi

ma für jedes Unternehmen. Man sollte hier nichts über einen Kamm scheren. Tatsache ist aber, dass jeder Bedarf an SSD-Storage vom Privatanwender mit einem 200 Euro Budget bis zum internationalen Großkonzern abgedeckt werden kann. Letzterer wird natürlich eine abgestufte Speicherhierarchie anwenden, die im Idealfall einheitlich verwaltet und benutzt werden kann. Die SSDs werden zunächst die heute teureren Platten verdrängen.

Die SSDs haben sich nach Moore's Law weiterentwickelt. Für die nahe Zukunft steht jedoch noch ein Sprung in der Evolution bevor. Es ist gelungen, einen stabilen Herstellungsprozess für dreidimensionale integrierte VLSI-Strukturen zu schaffen. Das bedeutet einfach gesprochen das Übereinanderlegen und Verbin-

den von hauchdünnen Wafern. Dadurch können neben anderen Schaltungen auch 3D-SSDs entstehen. Ihre Verfügbarkeit wird die Marktbedingungen nochmals dramatisch verändern.

Die Tabellen 10 und 11 fassen die Spitze des aktuellen Angebots an Enterprise-SSDs nochmal zusammen.

Erschreckend ist aber, dass es für die Anbindung der SSDs streng genommen gar keine wirklich passende Netzwerktechnologie gibt. Ein heutiger Großanwender muss wohl oder übel verschiedene Technologien mischen:

- Ethernet für normalen Datenverkehr und iSCSI-SATAs

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

manufacturer	model	technology	interface	performance metrics and notes
Kove	Xpress Disk	RAM SSD	InfiniBand	20GB/s throughput via 6x InfiniBand ports
Texas Memory Systems	RamSan-440	RAM SSD	Fibre Channel	4GB/s random sustained external throughput, 600,00 random IOPS
NextIO	vSTOR S100	SLC flash SSD	PCIe	2.2 million IOPS, 5.5GB/s sequential read, 6GB/s sequential write
Texas Memory Systems	RamSan-630	SLC flash SSD	InfiniBand	1 million IOPS, 10GB/s bandwidth, 80 microsecond latency (write)
Violin Memory	Violin 1010	RAM SSD	PCIe	1 million random IOPS on a sight port. 1,700MB/s read, 1,00MB/s write with x4 PCIe, 3 microseconds latency
Solid Access Technologies	USSD 320	RAM SSD	Fibre Channel SAS SCSI	100,000 random IOPS on a sight FC port. Upto 320,000 IOPS in a single 2U chassis. 4,000MB/s aggregated bandwidth via 6 FC links.
Solid Access Technologies	USSD 310	RAM SSD	Fibre Channel SAS SCSI	Up to 130,000 IOPS in a single 1U chassis. 1,500MB/s aggregated bandwidth via 2 full duplex FC links. 10 microseconds latency

Abbildung 10: Enterprise SSDs Rack

Quelle: StorageSearch.com

manufacturer	model	technology	interface	performance metrics and notes
Fusion-io	ioDrice Octal double-width card	Flash SSD	PCIe	1 million IOPS 6.2 GB/s of bandwidth
Virident Systems	tachIO on single slot card	Flash SSD	PCIe	300K random IOPS (75% R, 25% W). Peak R/W performance 1.44 GB/s and 1.2GB/s respectively
STEC	ZeusRAM SSD	RAM SSD	SAS	Under 23 microseconds average latency
STEC	ZeusIOPS	regular flash SSD	SAS	80,000 IOPS random read, 40,000 IOPS random write with transfer speeds of 550MB/s read and 300MB/s write
Pliant Technology	Lightning EFD	skinny flash SSD	SAS	160,000 IOPS. R/W throughput 525/340MB/s.
Density Dynamics	DDR2 jet.io	RAM SSD	Fibre Channel	400MB/s R/W, 10 μ S latency

Abbildung 11: Enterprise SSDs PCIe und 3,5"

Quelle: StorageSearch.com

- Fibre Channel für alte Plattenstapel in FC-SANs
- Infiniband für SSDs

Das ist wirklich eine ganz tolle Realität der Konvergenz: drei Sorten Netze, drei Sorten Kästen, drei Sorten Anschlusstechnik und mit ganz viel Glück eine gemeinsame Steuerung für die Speicher.

Der einzige Hersteller, der sich bereits dieses Problems wirklich angenommen hat, ist Mellanox/Voltaire. Hier wird es bald wenigstens einen 320 G IB-Adapter geben.

Alle anderen Hersteller waren in den letzten Jahren wohl derart mit anderen Diskussionen beschäftigt, dass die SSD-Entwicklung ungehindert an ihnen vorüber gehen konnte, ohne irgendeine Wahrnehmung zu erzeugen. Auch die Standardisierung hat kläglich versagt. Die höchste Ethernet-Ausbaustufe mit 100 G ist zwar definiert und es gibt Produkte, aber für

eine SSD-Anbindung führt auch das wieder zu Bündeln. Eine 400 GbE oder TbE-Standardisierung steht noch aus und wird wohl auch noch mindestens zwei Jahre benötigen.

Letztlich liegt es an den Chips. Mit dem Trident von Broadcom feiern wir grade den Einzug der 40 nm-Technologie bei Switch-ASICs. Die SSDs sind allerdings schon längst bei der 25 nm-Technologie angelangt. Zwischen dem modernsten Switch-ASIC und dem modernsten SSD-Chip, wie dem SanDisk 64 GB-Modell liegen technologisch WELTEN. Die einzige technologische Antwort wären CWD-M-Ethernets. Die werden auch wegen der Integration optischer Komponenten in einigen Jahren kommen. Heute wäre eine solche Lösung zu teuer.

Wie kann man angesichts dieser vorliegenden Schiefelage dennoch ein RZ-Netz zukunftsfähig planen?

Nicht ohne Absicht habe ich das eigentlich für den Consumer Markt gedachte Intel-SSD vorgestellt. Es zeichnet sich nämlich wie auch bei Plattenspeichern eine Klassenbildung bei SSDs ab:

- High End SSDs mit ca. 1-10 Mio. IOPS und multipler 100 GB I/O-Fähigkeit in Nutzung und Weiterentwicklung der Technologien, wie sie u.a. Texas Memory Systems liefern.
- Low End SSCs mit ca. 0,1 - 1 Mio. IOPS und 10 - 100 GB I/O-Fähigkeit in Nutzung und Weiterentwicklung der Produkte, die aktuell eher für den Consumer-Markt gedacht sind.

Diese unterscheiden sich massiv im Preis. Ein Unternehmen, welches tatsächlich einen Bedarf an High End SSDs hat, wird genau ausrechnen, welcher Etat dafür zur Verfügung steht. Dieser umfasst dann auch bündelweise Anbindungen.

SSD: Zukunft der professionellen Datenspeicherung und Ende konventioneller RZ-Netze

Die große Mehrheit der Unternehmen wird sich aber mit den Low End SSDs begnügen können. Diese ersetzen die bisherigen Plattenstapel und benötigen weniger Platz und deutlich weniger Strom. Außerdem können sie mit 40 oder 100 GbE an ein Netz angeschlossen werden.

Der aufmerksame Leser wird aber festgestellt haben, dass ich in der obigen Differenzierung die LATENZ nicht aufgeführt habe. Und es ist tatsächlich so, dass sich High End und Low End in IOPS und Da-

tenrate erheblich unterscheiden, aber NICHT WIRKLICH in der möglichen Gesamtkapazität oder hinsichtlich der Latenz. Das sollte sowohl an den einzelnen Produktvorstellungen als auch an den Tabellen klar geworden sein.

Und für den Aufbau latenzarmer Netze mit 10/40/100-Anschlüssen gibt es ja genügend Auswahl. Auch wenn es den Leser langweilt, möchte ich abschließend nochmals darauf hinweisen, dass eine Implementierung der DCB-Funktionen existentiell ist. Aktuell bewegt sich

der breite Markt in Richtung der Anbindung auch von SSDs mit iSCSI. Führend ist hier die Variante iSCSI über IB, aber auch iSCSI über Converged Ethernet ist durchaus ohne wesentliche Qualitätsverluste ein gangbarer Weg. Aber eben NUR über Converged Ethernet.

Der Siegeszug der SSDs ist nicht mehr aufzuhalten. Also müssen wir versuchen, mit den bestehenden Netzwerk-Produkten und -Strategien einen Weg zu ihrer sinnvollen Einbindung zu finden.

Kongress

Rechenzentrum Infrastruktur-Redesign Forum 2011 07.11. - 10.11.11 in Königswinter

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

Die Ausgangssituation ist durch folgende Stichworte zu kennzeichnen:

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

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

Aus der Perspektive der Systemarchitektur kann man vereinfachend sagen, dass die Kommunikation im RZ von drei neuen Verkehrsströmen geprägt wird:

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

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.

Moderatoren: Dr. Franz-Joachim Kauffels, Dr.-Ing. Behrooz Moayeri

Preis: € 2.190,-* zzgl. MwSt. mit Workshop am letzten Tag

€ 1.790,-* zzgl. MwSt. ohne Workshop am letzten Tag

*gültig bis zum 31.08.2011 - dann regulär € 2.390,- zzgl. MwSt. bzw. € 1.990,- zzgl. MwSt.



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

Aktuelle Veranstaltungen

Ausschreibungen im Informations- und Kommunikationsbereich, 11.07. - 11.07.11 in Bonn

Diese Veranstaltung ist als Leitfaden und Praxisseminar für öffentliche Auftraggeber gedacht, die in ihren ITK-Vergabeverfahren unter Einhaltung aller gesetzlichen Auflagen und Vermeidung aller rechtlichen Risiken für ihre Verwaltung das optimale Ausschreibungsergebnis erreichen wollen. Auch die Vertreter von ITK-Unternehmen sollen durch die Veranstaltung angesprochen werden, um die Rechts- und Investitionssicherheit bei der Teilnahme auf Bieterseite zu verbessern und optimales Angebotsverhalten zu erreichen.

Preis: € 890,- zzgl. MwSt.

Aktuelle VPN-Technik, 11.07. - 13.07.11 in Bonn

Die Nutzung von VPN-Technik hat sich in der jüngeren Vergangenheit insbesondere im Bereich des Remote Zugriffs mobiler oder auch stationärer Anwender (Stichwort: Telearbeit) auf zentrale Ressourcen als mehr oder weniger Standard-Lösungsansatz etabliert. Aber auch zur kostenoptimierten Anbindung von (typischerweise kleineren) Remote-Standorten an Corporate WAN-Strukturen bewährt sich dieser Ansatz. Dieses Seminar vermittelt die für einen erfolgreichen VPN-Einsatz notwendigen Kenntnisse der aktuell relevanten Technologien. Alle wesentlichen Bausteine typischer Lösungen werden detailliert erklärt und anhand praktischer Projektbeispiele und Übungen wird der Weg zu einer erfolgreichen VPN-Lösung aufgezeigt

Preis: € 1.890,- zzgl. MwSt.

ComConsult Storage-Forum 2011, 13.07. - 15.07.11 in Bonn

Zentraler, im Netzwerk zugreifbarer Speicher steht im Mittelpunkt aller zukünftigen IT-Architekturen. Die technologische Spannweite ist riesig, permanent kommen neue Entwicklungen und Produkte hinzu. Das ComConsult Storage-Forum 2011 analysiert die aktuellen Entwicklungen im Speichermarkt, vergleicht Alternativen und zeigt auf, wo der Weg hingeht.

Preis: € 1.990,- zzgl. MwSt.

Sommerschule 2011 - Intensiv-Update auf den letzten Stand der Netzwerktechnik, 18.07. - 22.07.11 in Aachen

Die Sommerschule gibt innerhalb von 5 Tagen den kompakten und intensiven Überblick über die neusten Entwicklungen im Umfeld der Netzwerk-Technologien: Anforderungen an zukunftssichere Netzwerke: was ändert sich; Neue Technologien und Standards; Design-Verfahren im Vergleich; Ausgewählte Technologien in der Analyse; Umfeld-Analyse: was passiert um Netzwerke herum.

Preis: € 2.490,- zzgl. MwSt.

Sicherer Internetzugang, 07.09. - 09.09.11 in Aachen

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

Preis: € 1.890,- zzgl. MwSt.

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

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

Preis: € 1.890,- zzgl. MwSt.

IT-Projektmanagement Kompaktseminar, 12.09. - 14.09.11 in Aachen

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

Preis: € 1.890,- zzgl. MwSt.

Lokale Netze für Einsteiger, 12.09. - 16.09.11 in Aachen

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

Preis: € 2.490,- zzgl. MwSt.

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

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

Preis: € 890,- zzgl. MwSt.

Zertifizierungen

ComConsult Certified Network Engineer**Lokale Netze**

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

TCP/IP intensiv und kompakt

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

Internetworking

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

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

ComConsult Certified Trouble Shooter**Trouble Shooting in vernetzten Infrastrukturen**

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

Trouble Shooting für Netzwerk-Anwendungen

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

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

ComConsult Certified Voice Engineer**Session Initiation Protocol Basis-Technologie der IP-Telefonie**

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

Umfassende Absicherung von Voice over IP und Unified Communications

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

IP-Telefonie und Unified Communications erfolgreich planen und umsetzen

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

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

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

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

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

Impressum

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

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

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

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

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