

Schwerpunktthema

Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

Netzwerk-Design und Auslegung haben sich in den letzten 12 Monaten rapide verändert. Neue Mechanismen (802.1p, TOS, RSVP, COPS, DEN, VRRP, ...) werden mit traditionellen (was heißt: älter als ein Jahr) Highend-Verfahren kombiniert (OSPF, HSRP, ...) und bieten vielfältige Möglichkeiten ein Netz zu strukturieren und zu steuern. Leider bieten sie auch gute Chancen, sich sehr abhängig von einzelnen Herstellern zu machen und ein völliges Chaos zu produzieren. Der folgende Artikel versucht, eine Linie in das Design moderner Netzwerke zu bringen.

Traditionelle Netzwerke im Bereich Ethernet und Token Ring werden auch als "best effort" Netzwerke bezeichnet. Kennzeichnend für diese Netzwerke ist, dass es kaum Regel- und Steuerungs-Mechanismen gibt, um den Kapazitäts- und Leistungsbedarf bestimmter Anwendungen sicher zu stellen. Alle Netzwerk-Anwendungen werden gleich behandelt und stehen im freien Wettstreit

von Dr. Jürgen Suppan



Marketing-Ballon
»Policy-based-Networking«:
Was steckt dahinter?

um die Ressourcen des Netzwerks. Keine Applikation hat im "best effort" Netzwerk eine Garantie auf eine bestimmte Leistung. Bei der Kombination aus niedrigen Übertragungsraten (10 oder 16 MBit/s) mit hohen Spitzenlasten treten so unvermeidbar Überlasten auf, die zufallsbasiert einigen Applikationen schaden und anderen nicht. Speziell die Leistung Datenbank-basierter Netzwerk-Anwendungen, von Clustern oder auch die Nutzung von Sprache und Video im Netz kann in dieser Umgebung erheblich leiden (starke Performance-Schwankungen). Naturgemäß wird in derartigen Überlast-Situationen auch die Nutzung von Management-Tools eingeschränkt, da diese auch nicht bevorzugt behandelt werden (im Fall SNMP sorgt allerdings die verbindungslose Kommunikation auf der Basis von Einzelpaketen dafür, dass Management besser funktioniert als ein verbindungsorientierter Dienst).

weiter Seite 7

Zum Geleit

Glasfaser am Scheideweg

von Dr. Jürgen Suppan

Betrachtet man die Schwierigkeiten bei der Normung von 1000 Base T und sieht die begleitenden Diskussionen, so muss parallel die Frage nach der Zukunftsfähigkeit von Twisted Pair und der Glasfaser-Alternative entstehen. Unter anderem sind neben den Diskussionen um die Link Klassen E und F offensichtlich weder IEEE noch die ISO in der Lage, das Steckerthema im Twisted Pair Bereich abschließend zu klären.

Mehrere Varianten lagen vor, deren Einsatzspektrum bis 1000 MHz reichte, doch haben sie offensichtlich alle einen Nachteil: sie sehen nicht aus wie RJ45. Dabei kann genau dieser Stecker nur als Ergebnis einer sorgfältigen Suche nach dem ungeeignetsten Stecker für Datenübertragung angesehen werden (parallele und somit nebensprech-

empfindliche Kontakte, ermüdungsgefährdete Feder speziell bei Steckzyklus 1, keine hohen Frequenzen erreichbar usw.).

Unabhängig von der Frage des Steckers entsteht bei Twisted Pair mit Frequenzen über 100 MHz die Frage nach einer funktionsfähigen Schirmung. Diese ist nur erreichbar, wenn der Schirm an beiden Enden aufgelegt wird. Für viele bestehende Verkabelungen ist dies nicht gegeben, für viele neue Verkabelungen kann dies zu einem Problem des Potentialausgleichs führen. Dieses scheinbar harmlose Thema wird in vielen Projekten völlig übersehen. Wer sich jedoch noch an das Problem der effizienten Schirmung von Breitbandnetzen aus den 80er Jahren erinnert, dem wird klar, dass hier ein sehr ernstzunehmendes Problem droht.

Genau in dieser Situation bekommt die Glasfaser durch neue Entwicklungen frischen

Wind in die Segel. Neue Stecksysteme, sogenannte "Small Form Factor Fiber Optic Connectors" oder auch "Small Connector Fiber Cabling" genannt, reduzieren die Konfektionierungskosten um mindestens 20% gegenüber den bisherigen Lösungen. Die ersten Projekterfahrungen lagen sogar über diesem Wert. Im Vergleich zu Twisted Pair Lösungen hängt die Kostensituation allerdings sehr von der Größe der Installation ab. Je mehr Verteillerräume für Twisted Pair notwendig sind, desto besser sieht die Glasfaser aus. Dies kann im Extremfall auch bei Link Class E Verkabelungen zu fast identischen Preisen führen, im Vergleich Link Class F kann eine LWL-Verkabelung sogar Kostenvorteile zwischen 15 und 30% erreichen. Leider werden diese Kostenvorteile durch die LWL-Portkosten der eingesetzten Switch-Systeme wieder in Frage gestellt.

weiter Seite 2

Glasfaser am Scheideweg

Fortsetzung von Seite 1

Erfreulich ist die Entwicklung im NIC-Bereich. Hier sind die ersten MT-RJ-Karten für 300 DM auf dem Markt.

Wer einen Gesamtkostenvergleich zwischen Glasfaser und Twisted Pair im Gebäude anstrebt, sollte folgende Kostenbereiche beachten und kalkulieren (dabei ist zu beachten, dass zur Zeit keine Möglichkeit besteht, Telefon-Endgeräte mit Glasfaser anzuschließen, eine parallele Sprachverkabelung also in der Kalkulation berücksichtigt werden muss, bis dieses Problem gelöst ist):

- Kabel,
- Stecker,
- Verlegekosten,
- Konfektionierungskosten,
- Herstellkosten Technikräume inkl. Strom, Klima, Zugang,
- Baumaßnahmen für Kabelwegebau (VDE, Biegeradien!),
- Komponentenkosten,



- Umrüstkosten für Endgeräte,
- Kosten für Adapterkarten,
- laufende Raumkosten hochgerechnet,
- laufende Personalkosten hochgerechnet,

Die Dämpfungen der neuen optischen Systeme sind für den Endgerätebereich auf jeden Fall ausreichend, im Gelände kann je nach Situation der Wunsch nach besseren Werten und einer Optimierung der Dämpfung herrschen.

Welche neuen "Small Connector Fiber Cabling" Systeme existieren? Leider konkurrieren mindestens 3 verschiedene Systeme miteinander und die TIA (Telecommunications Industry Association) sieht sich nicht in der Lage, sich auf eine Variante festzulegen:

- LC unterstützt durch Lucent
- MT-RJ unterstützt durch Alcoa, AMP und Siecor (sowie HP und IBM)
- VF-45 unterstützt durch 3M

Eine Vorentscheidung in Richtung MT-RJ scheint gefallen, da die wesentlichen Hersteller von Netzwerk-Komponenten diesen Stecker schon anbieten oder in diese Richtung tendieren (Bay/Nortel, CISCO, HP, IBM). Dringend erforderlich wäre eine intensivere Unterstützung durch 3Com und Intel bei den Adapterkarten.

Auf jeden Fall bieten alle 3 Systeme den Vorteil erheblich höherer Packungsdichten auf RJ45-Niveau, was insbesondere auch für die Kosten der Switch-Ports entscheidend ist, da im Vergleich zu Twisted Pair Einschüben inzwischen identische Packungsdichten erreicht werden können (siehe 3Com). So lassen sich bereits bei den ersten Switch-Herstellern Port-Kosten-Verbesserungen um ca. 25% durch Einsatz der neuen Stecksysteme feststellen.

Gleichzeitig steigt die Geschwindigkeit der Installation um bis zu 70%, da die neuen Systeme in wesentlich kürzerer Zeit installiert werden können.

Noch ist es zu früh, um ein abschließendes Urteil zu fällen. Es fehlt der abschließende Standard und auch ein Redesign der Small-Connector-Komponenten ist noch denkbar. Auch müssen wesentlich mehr Erfahrungen in der Praxis sowohl zum Thema Installation als auch zum Betrieb gesammelt werden. Aber eines steht bereits jetzt fest: im Rennen um die Verkabelungszukunft hat LWL speziell im Bereich der Endgeräte-Verkabelung erheblich aufgeholt. War im Gesamtkostenvergleich über alle Projekte das Verhältnis im Schnitt 70% pro Twisted Pair, so kann bereits für die nähere Zukunft ab einer Mindest-Gebäudegröße mit ausgeglichenen Kosten gerechnet werden.

Für viele Planer besteht so die Möglichkeit, ein sehr wohl existierendes Unwohlsein zum Thema Twisted Pair und Schirmung wirkungsvoll zu bekämpfen. Die immer wieder in der Presse oder der Werbung erwähnten Bandbreitenvorteile der Glasfaser sind allerdings innerhalb der Gebäude mit 62,5 Mikrometer Multimode nur sehr eingeschränkt gegeben, wie der Differential Mode Delay in der Normung von 1000 Base x gezeigt hat. Ob auf Multimode-Fasern in den gegebenen Längen jemals Übertragungsraten von 10 Gbit/s übertragen werden können, bleibt abzuwarten. Hier könnte im Extremfall sogar ein Vorteil von Twisted Pair gegenüber Multimode-Fasern entstehen. Allerdings wird im Moment eine Datenrate von 10 Gbit/s zum Endgerät nicht angestrebt, sondern in der 10-Jahres-Planung normalerweise als Ziel die 1 Gbit/s-Fähigkeit angesetzt. 10 Gbit/s bieten eine Datenrate und Zugriffsgeschwindigkeit, die auf dem Niveau hochwertiger RAM-Bausteine liegt. Hinzu kommt, dass begingt durch TCP/IP und SMB-Verluste (Netbios über TCP) auf dem typischen Windows 98 Endgerät zur Zeit im Idealfall gerade einmal Datenraten von 30 Mbit/s erreicht werden, eine Projektion auf 10 Gbit/s also sehr viel Vertrauen in die Programmierer von Microsoft beinhalten würde.

Wie man sieht, ist der zukünftige Weg spekulativ. Auf jeden Fall sollte bei der Planung von neuen Gebäude-Verkabelungen ein sehr intensiver Preis und System-Vergleich zwischen Glasfaser und Twisted Pair durchgeführt werden. Vor einer Entscheidung für Glasfaser sollte auf jeden Fall sorgfältig geprüft werden, ob Endsysteme im Haus sind, die nicht über Glasfaser kommunizieren können (Telefone, Terminals, Zugangssicherung, Brandmeldeanlagen).

Auf eine erleuchtete Zukunft

Ihr Dr. Jürgen Suppan

Impressum

Verlag:
ComConsult
Technologie Information GmbH
Pascalstraße 25
D-52076 Aachen
Telefon 02408/14907
Telefax 02408/149233

Herausgeber und verantwortlich
im Sinne des Presserechts:
Dr. Jürgen Suppan

Gestaltung: Zweipfennig.de

Erscheinungsweise:
Monatlich
12 Ausgaben im Jahr

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

Für unverlangt
eingesandte Manuskripte
wird keine Haftung
übernommen

Nachdruck, auch auszugsweise
nur mit Genehmigung des Verlages

© ComConsult Technologie Information

Der Top-Trend: Service-Level- und Availability-Management

Nachdem bei vielen Unternehmen die Grundfunktionen einer Management-Lösung für Netzwerke, Server und Applikationen realisiert worden sind, kristallisieren sich neue Funktionsschwerpunkte heraus. Zu den heißesten Themen der Stunde zählen dabei Service-Level und Availability-Management.

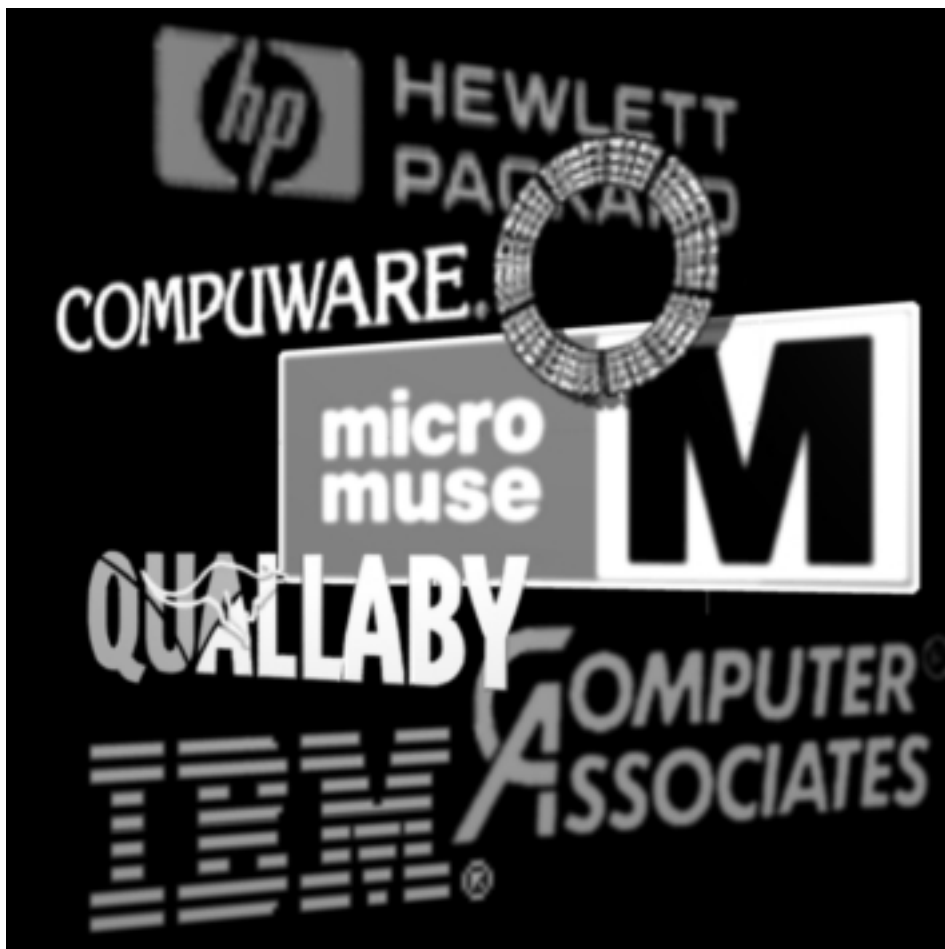
Unter Service-Level-Management wird verstanden, dass nicht die Betriebssituation einer einzelnen Komponente überwacht wird sondern die Einhaltung zuvor dem Anwender gegenüber festgelegter Service-Level. Dies ist wesentlich komplexer, wie an Beispielen für Service-Level-Parameter zu sehen ist:

- Antwortzeiten von Applikationen
- Verfügbarkeit von Applikationen
- Performance-Werte im weitesten Sinne
- Fehlerfreiheit von Applikations-Nutzungen (auch für Sprache und Video)

Typisch für das Management von Service-Leveln ist, dass die Einhaltung der festgelegten Werte nur Technologie-übergreifend sichergestellt werden kann. Bei einem WAN-Zugriff auf ein zentrales Rechenzentrum zur Nutzung von R3 müssen zur Einhaltung des Service-Levels untergeordnete Sub-Service-Level realisiert werden für:

- das Endgerät
- das LAN am Endgerät
- die WAN-Verbindung
- das LAN am Server
- den Server
- die Applikation

Die bestehenden Produktansätze sind sehr unterschiedlich. Im Idealfall messen sie primär die Service-Level-Situation aus der Sicht des Endanwenders. Hier gibt es unterschiedliche technologische Ansätze (Agent auf jedem Endsystem, Messprobe im LAN des Endanwenders, Referenzsystem im LAN des Endanwenders). Ist der definierte Service-Level nicht gegeben, muss die Suche nach der Ursache ge-



Trendsetter: Compuware, Micromuse, Quallaby, Computer Associates, HP und IBM

startet werden. Dazu muss Technologie-übergreifend auf alle Systeme zugegriffen werden, die den Service-Level beeinflussen.

Unter Availability-Management wird das gezielte Überwachen und Ermitteln der Verfügbarkeit einzelner Komponenten oder Dienste verstanden. Einerseits kann Verfügbarkeit in diesem Sinne Bestandteil von Service-Level-Vereinbarungen sein, es kann aber auch für die Vertragsüberwachung eines externen Vertragspartners (Lieferant, Wartung, Outsourcer etc.) wichtig sein. Availability ist komplex, wenn hierunter nicht die generelle Verfügbarkeit (läuft, läuft nicht) verstanden wird sondern ein detailliert beschriebener Leistungsgrad, der verfügbar sein muss. Selbstverständlich wird hierunter auch die Projektion auf die erwartete zukünftige Verfüg-

barkeit verstanden.

Alle Produktansätze in diesem Bereich sind noch im Fluss. Sie reichen von in Deutschland noch relativ unbekannten Herstellern wie Compuware, Micromuse, Quallaby bis hin zu den etablierten Anbietern wie Computer Associates, HP und IBM. Einige liefern komplette Lösungen, andere setzen strategisch auf Kombinationen mit bestehenden Produkten wie Microsoft Excel und SAS.

Einen Überblick über die Technologie und Produktsituation wird das Netzwerk- und System-Management-Forum im November in Neuss liefern. Da dort bereits über 80% der Plätze ausgebucht sind, empfiehlt sich für den Interessenten/in eine schnelle Buchung.

Aktuelles zur Technologie- und Produkt-Situation zu Service-Level und Availability-Management auf dem

Netzwerk- und System-Management Forum 1999
08.11. - 11.11.99 in Neuss

► 4 Tage (mit Tutorium) zum Preis von DM 3.270,-- (EUR 1.671,92) zzgl. MwSt.

► 3 Tage (ohne Tutorium) zum Preis von DM 2.890,-- (EUR 1.477,63) zzgl. MwSt.

Die Buchung ist verbindlich, kann aber jederzeit auf andere Mitarbeiter Ihres Unternehmens übertragen werden.

Bitte buchen Sie per

eMail: akademie@comconsult.de oder via Internet www.comconsult-akademie.de

In der Hoffnung, Sie in Neuss begrüßen zu können,
mit freundlichen Grüßen

Stephanie Schroeder
- Organisationsleitung -
ComConsult Akademie

Fax-Antwort an ComConsult 02408/149233

Anmeldung Management-Forum

Ich melde mich an für das
**Netzwerk- und System-Management
Forum 1999**

vom 08.11. - 11.11.99 in Neuss
zu den angekreuzten Konditionen an

☐ 4 Tage (mit Tutorium)
zum Preis von DM 3.270,--
(EUR 1.671,92) zzgl. MwSt.

☐ 3 Tage (ohne Tutorium)
zum Preis von DM 2.890,--
(EUR 1.477,63) zzgl. MwSt.

Faxen Sie uns einfach nebenstehenden Ab-
schnitt an **02408/149233** oder schicken Sie
eine eMail an akademie@comconsult.de
oder buchen Sie über unsere Web-Seite
<http://www.comconsult-akademie.de>
Ihren Platz auf unserer Veranstaltung.

Name

Vorname

Firma

Abteilung

Position

Funktion

Straße

PLZ, Ort

Telefon

Fax

eMail

Unterschrift

Neues Seminar

Live-Evaluierung von Microsoft SMS 2.0

Ist Microsoft SMS 2.0, das augenscheinlich das komplette Management-Spektrum von der Inventarisierung, Remote Support, Software Verteilung, Lizenz-Überwachung, Netzwerk-Management, PC-Management, Server-Management abdeckt, der neue Standard?

Die Vorgänger-Version 1.2 hatte einige Defizite. Naturgemäß stellt sich jetzt die Frage, ob diese mit der SMS 2.0 - Version beseitigt worden sind. Dazu gehört auch die Frage nach der Komplexität und der Praxistauglichkeit. Dieses Seminar prüft unter praxisnahen Gesichtspunkten die technische Nutzbarkeit von SMS 2.0 und führt wichtige Funktionsbereiche live vor. Dabei wendet sich dieses Seminar sowohl an potentielle Neueinsteiger in SMS als auch an die Betreiber der Versionen 1.x.

Das herausragende Seminar zeigt die Stärken und Schwächen mit vielen praxisnahen Live-Demos, und setzt die Teilnehmer in die Lage, zu entscheiden, ob der Einsatz von SMS 2.0 für Ihr Unternehmen sinnvoll ist. Der Seminaraufbau gestaltet sich so, dass ca. 60 % der Zeit



Nabelschau: SMS 2.0

die Eigenschaften von SMS vorgeführt und diskutiert werden und die Teilnehmer zu 40% der Zeit die erörterten Funktionen an bereit stehenden Rechnern ausprobieren können. Dazu teilen sich jeweils 2 Teilnehmer einen Rechner. Diese herstellernerneutral moderierte Evaluierung dient zur objektiven Entscheidungsfindung und we-

niger zur technischen Ausbildung hinsichtlich einer selbständigen Implementierung des Systems Management Servers. Zudem beinhaltet das Seminar keinen 1:1 Vergleich zu Mitbewerbsprodukten, da dies im Rahmen der zur Verfügung stehenden Zeit nicht möglich ist. Die Anzahl der Plätze ist auf 16 begrenzt, um die interaktive Kommunikation zu den Referenten zu gewährleisten.

Die Veranstaltung wird unter der Leitung von Dipl.-Ing./MCSE Martin Barbian von der ITC GmbH aus Detmold durchgeführt. Martin Barbian zählt mit seinem Team zu den führenden deutschen SMS-Experten. Das im IT Management Umfeld bundesweit etablierte Consulting Unternehmen agiert in diversen Gremien und Partnerprogrammen wie im Microsoft SMS Specialist Partner Program, MS SMS 2.0 Beta Test Program, MS SMS InnerCircle u.v.m.

Darüber hinaus wurden unter der Leitung von Martin Barbian diverse SMS-Projekte durchgeführt und mittels Tools anderer Hersteller integrierte IT Management Lösungen realisiert.

Fax-Antwort an ComConsult 02408/149233

Anmeldung SMS 2.0-Seminar

Ich melde mich an für das Seminar
Live Evaluierung von Microsoft SMS 2.0
zum Preis von DM 2.990,- (EUR 1.528,76)
zzgl. MwSt. zu dem angekreuzten Termin an

☐ **18.10. - 20.10.99** Hamburg

☐ **13.12. - 15.12.99** München

☐ **21.02. - 23.02.00** Bonn

Faxen Sie uns einfach nebenstehenden Abschnitt an **02408/149233** oder schicken Sie eine eMail an **akademie@comconsult.de** oder buchen Sie über unsere Web-Seite **<http://www.comconsult-akademie.de>** Ihren Platz auf unserer Veranstaltung.

Name

Firma

Position

Straße

Telefon

eMail

Vorname

Abteilung

Funktion

PLZ, Ort

Fax

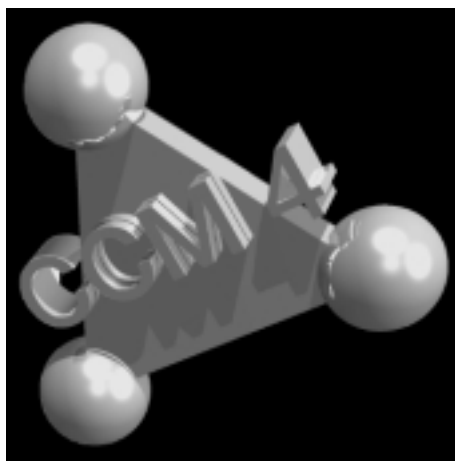
Unterschrift

Neues vom CCM: Version 4.2 ist da!

Der ComConsult Communication Manager, durchgängige Integrationslösung für die Bereiche Dokumentation, Netzwerk-Management, System-Management und User Help Desk auf der Basis marktführender Systeme, wird heute europaweit von über 15.000 CCM-Anwendern für über 3 Milliarden Anschlüsse zur effizienten Netzwerk-Verwaltung eingesetzt.

Die CCM-Arbeitsplätze sind Windows NT-Clients - der CCM-Server ist für die Plattformen Unix und Windows NT verfügbar. Durch sein strukturiertes Datenmodell und seine offene Architektur unterstützt der CCM 4 den wirtschaftlichen, kostenoptimierten Betrieb von Netzen und Systemen in Unternehmen.

Die auffälligste Neuerung der Version CCM 4.2 ist die durchgängige graphische Benutzeroberfläche mit über 200 Masken. Alle Bildschirmmasken sind mit Hilfe eines De-



Verkürzt die Planungszeit
für DV-Arbeitsplätze auf 15 Minuten:
der neue CCM 4

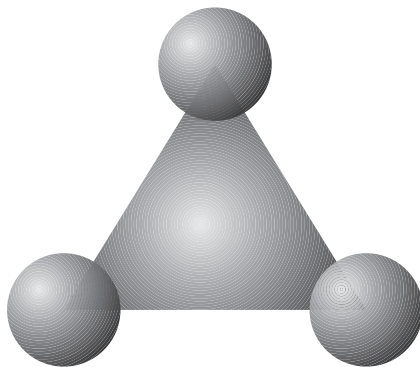
signwerkzeuges, dem Formulardesigner, anpassbar und erweiterbar. Für jedes Benutzerprofil lassen sich damit individuell Informationen und Funktionen bereitstellen.

Die wesentlichen Leistungsmerkmale des neuen CCM 4 sind jedoch die große Flexibilität des Systems und individuelle Anpassbarkeit zur Erhöhung der Gesamtfunktionalität für die Dokumentation und das Bestandsmanagement.

Mit WDS, dem neuen Workflow-Development-System des CCM, wurden typische Beispiele zur Unterstützung des Beschaffungswesens, des Konfigurationsmanagements und des Changemanagements vorgestellt. Das WDS orientiert sich am WfMC-Standard. Die Integration des CCM in die Managementsysteme Tivoli und HP OpenView und in das Helpdesktool AR-System von Remedy wurde intensiviert.

Anzeige

Erleben Sie live die erweiterten Möglichkeiten
WDS (Workflow-Development-System)
und technisches Data-Warehousing
für den Netz- und Systembetrieb
vom 21. bis zum 25. September
auf der ORBIT 99 in Basel
mit dem neuen CCM 4



Sie finden uns am Stand
der Firma NetDesign
Halle 1.1, Stand C 40
Wenden Sie sich bitte
an Silke Karius
Telefon 02408/149-272
und wir nehmen uns
extra Zeit für Sie.

ComConsult
Kommunikationstechnik GmbH

Sie sind gerade nicht in Basel,
oder der Termin ist ungünstig?
Kontaktieren Sie uns:
heiko.soldan@comconsult.de
oder Sie besuchen uns im Web:
www.comconsult.de
Gerne senden wir Ihnen
ausführliche Informationen.

Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

Fortsetzung von Seite 1

Die Technologie-Entwicklung der letzten Jahre im Bereich Switching hat nun in den Switch-Systemen Zusatz-Technologien geschaffen, die den Basis-Technologien Ethernet und Token Ring gefehlt haben. Im Endeffekt werden alle Steuerungs- und Gestaltungsmöglichkeiten im Netzwerk-Design in die Switch-Systeme verlagert. Dies geht zum Teil auch erheblich über Reservierungs- und Filterbedingungen einzelner Geräte hinaus und reicht bis zu komplexen Verbundschaltungen aller eingesetzten Switch-Systeme.

Für den Anwender hat diese Entwicklung nicht nur Vorteile. Eine Reihe der neuen Möglichkeiten werden in der Software der Switch-Systeme geschaffen. Zum einen kann dies unvorhersehbar die Performance der Systeme beeinflussen, zum anderen ermöglicht der Wechsel von der Basis-Technologie Ethernet / Token Ring in die Software der Switches extrem schnelle Technologie-

Wechsel im Wochen und Monatsbereich. Nur so ist zu erklären, dass Technologien wie RSVP im Bereich der Reservierung von Kapazität noch bevor sie überhaupt ernsthaft eingesetzt werden bereits wieder in einer Nachfolgediskussion sind.

Es stellt sich nun die Frage, welche Konsequenzen sich für das Design zukunftsorientierter Netzwerke ergeben.

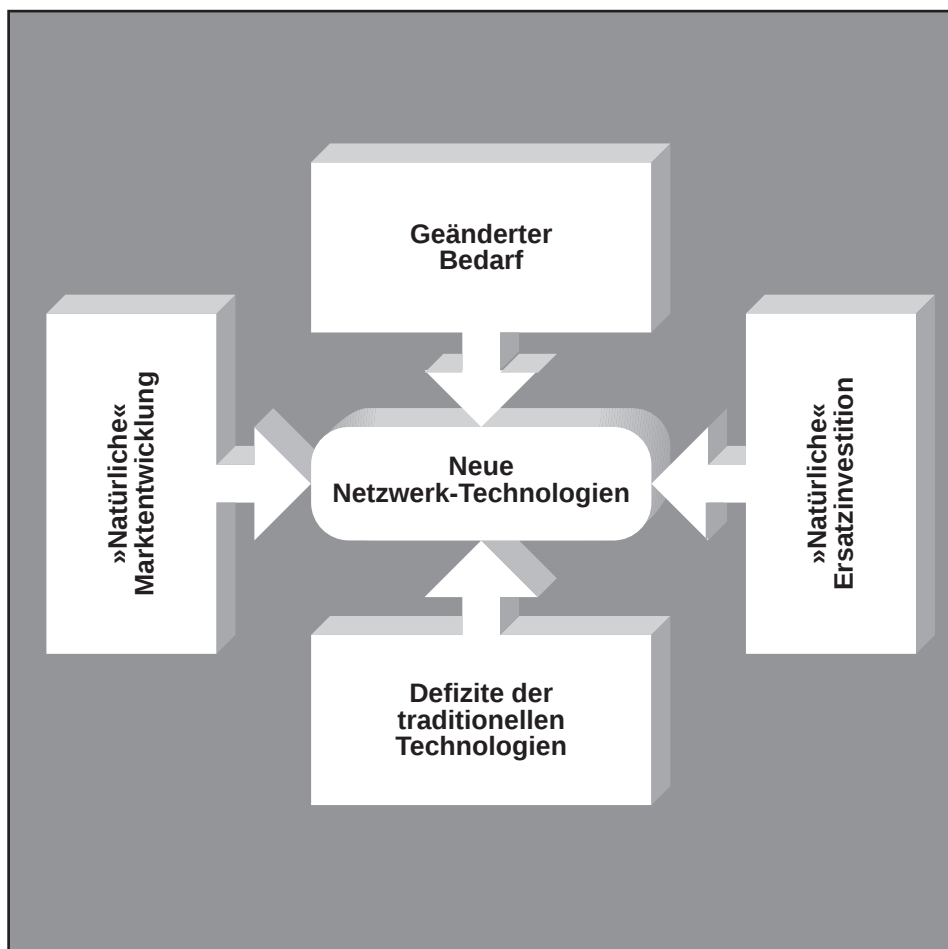
Auf den ersten Blick lassen sich mindestens zwei Design-Ansätze unterscheiden, die sich allerdings partiell überlappen:

- Service-Level-Netzwerk-Design basiert auf Gestaltungsbausteinen, die je nach gewünschten Leistungs-Niveau räumlich und technisch kombiniert werden. Insbesondere bei geringen und normalen Ansprüchen kann daraus ein auch einfaches Design abgeleitet werden. In den komplexeren Anwendungsbereichen würden auch Policy-basierte Konzepte zum Einsatz kommen. Somit kann Service-Level-

Design auch als Obermenge von Policy-based-Design angesehen werden.

- Policy-based-Netzwerk-Design basiert auf der Idee, daß ein Designer eine physikalische Switch-Struktur (Verbindung von Switch-Systemen) schafft, deren konkretes Verhalten auf der Basis von Regeln durch eine Policy-Console festgelegt wird. Der Anwender definiert dabei tendenziell abstrakte Strategien und der Policy-Server leitet daraus die Detail-Konfigurationen für die Switch-Systeme ab. Dies könnte beispielsweise bedeuten, daß der Designer ein Netz als hochverfügbar kennzeichnet und der Policy-Server daraus ein OSPF-Design automatisch ableitet, ohne dass dies der Anwender explizit gefordert hätte. Mit diesem Design-Ansatz will man insbesondere der Tatsache Rechnung tragen, dass viele der neuen Gestaltungs- und Konfigurations-Möglichkeiten für eine breite Masse von Anwendern eventuell nicht beherrschbar sind. Allerdings hat die Erfahrung mit ähnlichen Abstraktions-Ansätzen im Umfeld von ATM gezeigt, dass Trouble Shooting und Protokoll-Analyse durch diesen Ansatz nicht gerade erleichtert werden. Zudem ist zu beachten, dass der Policy-based Ansatz nicht genormt ist. Zur Zeit existieren mindestens 9 Ansätze, die sich im Inhalt und in der Skalierbarkeit in großen Netzen erheblich unterscheiden (siehe unten). Es besteht die ernstzunehmende Gefahr, dass Policy-based Ansätze zu einer sehr starken Hersteller-Abhängigkeit führen. Dies macht auch den großen Marketing-Aufwand einzelner Hersteller zu diesem Thema verständlich.

Bild 1 Warum neue Netzwerk-Technologien?



Wer nun meint, beide Ansätze würden auf seine heutige Problematik nicht zutreffen und er könnte ruhig noch ein wenig warten, der irrt natürlich gewaltig. Schließlich kauft er regelmäßig neue Switch-Systeme ein, deren Eignung für Service-Level-Gestaltung oder Policy-Based-Networking bei einem späteren Ausbau der Netze in diese Richtung elementar sein kann. Da die meisten der heute verkauften Midsize- und High-End-Switches sowohl Service-Level als auch Policy-Funktionselemente haben, werden durch falsche Kaufentscheidungen heute bereits implizit Entscheidungen für das zukünftige Policy-Konzept getroffen, ohne dass sich der Käufer ggf. darüber im klaren ist.

Um diese neuen Design-Ansätze zu bewerten, ist es notwendig, noch einmal auf die Grundfrage "wie definiert sich Netzwerk-Leistung" einzugehen. Dies soll auch mit der Frage verbunden werden, für welchen Anwendungsbereich welche Leistung benötigt wird.

Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

Bei der Bewertung des Bedarfs für neue Technologien und vor allem für mehr Kontrolle und Steuerung sind zwei Faktoren zu unterscheiden:

- die normative Kraft des faktischen. Wesentliche Elemente der Technologie-Entwicklung der letzten Jahre wurde nicht vom Bedarf sondern von den Strategien der Hersteller bestimmt. Fast Ethernet hat sich auch durchgesetzt, weil es durch massiv niedrig gehaltene Preise in den Markt gedrückt wurde. Bei Einstiegspreisen für Gigabit-Ethernet von 600 DM in Verbindung mit den typischen Preisentwicklungen ist somit klar, dass Gigabit-Technik in spätestens 12 bis 24 Monaten im Server-Bereich normal sein wird. Im Klartext: bei Neuinvestitionen muss bereits heute von der Gigabit-Fähigkeit mindestens der Backbone-Switch-Systeme ausgegangen werden.

- Der wirklich bestehende Neubedarf durch die Integration neuer Anwendungen in unsere Netze. Dazu zählen

- Mission Critical Server und Cluster mit angestrebten Verfügbarkeiten von deutlich über 99%
- Jede Form von Datenbank-Anwendung mit SQL übers Netzwerk
- Sprache und Video

Analysiert man, warum dieser neue Bedarf mit traditionellen Netzen nicht immer erfüllt werden kann, so stößt man mindestens auf zwei Kernthemen

- Verzögerung
- Quality of Service

Dies macht deutlich, dass die früher übliche Optimierung von Netzwerken nach reinen Kapazitätsgesichtspunkten für ein zukunftsorientiertes Netzwerkdesign unzureichend sind. Aus diesem Grund setzt sich die technische Leistung eines modernen Netzwerks heute in der Planung zusammen aus

- Kapazität
- Verzögerungszeit
- Verzögerungs-Varianz
- Verfügbarkeit und Stabilität

Kapazitäts-Planung

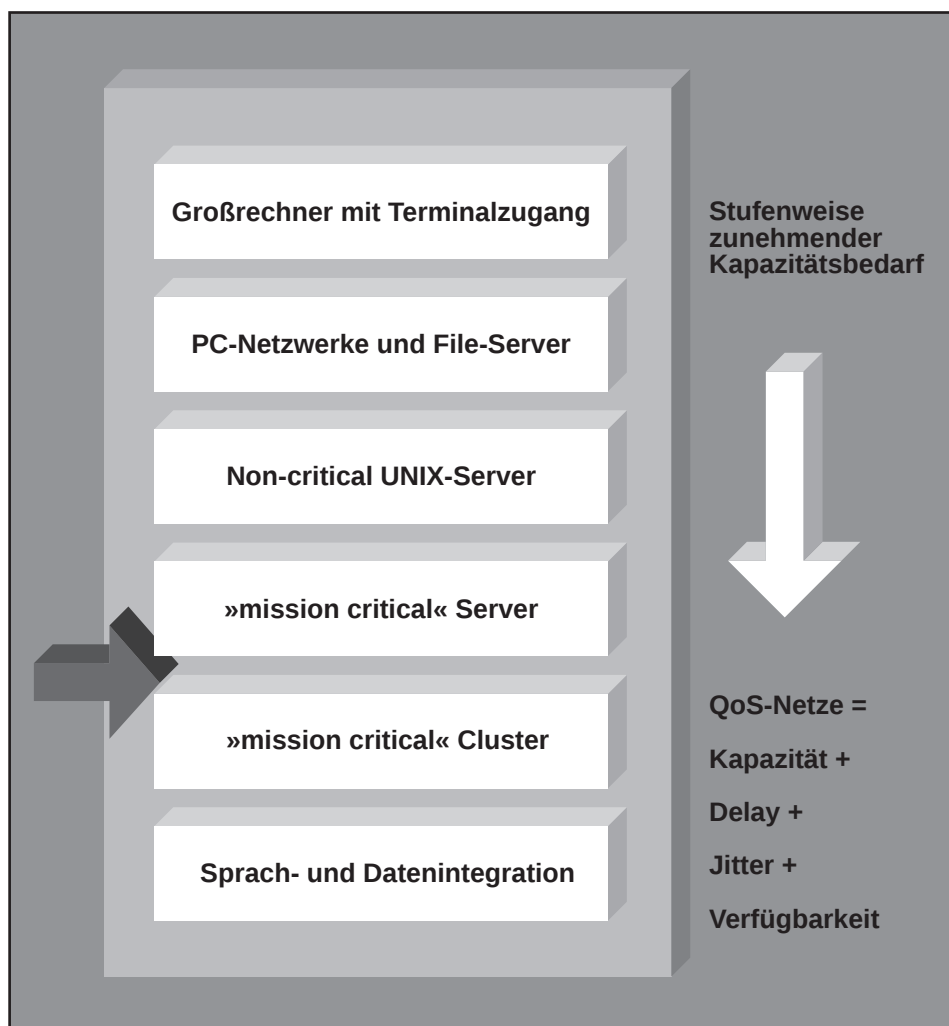
Bei der Kalkulation der notwendigen Netzwerk-Kapazität geht man immer noch von Server-zentrierten Datenströmen aus (Endgeräte untereinander kommunizieren nicht, Kommunikation erfolgt vom Endgerät zum Server oder zwischen Servern). Dies wird für einen Zeitraum von mehr als 3 Jahren ggf. nicht gültig sein, da mit Sprache und Video wahrscheinlich nur der Verbindungsaufbau über einen Server erfolgen wird und danach eine direkte Kommunikation zwischen den Endsystemen erfolgt. Die Basis dafür werden voraussichtlich Multicast-basierte Protokolle sein (Unterstützung von Gruppen-Kommunikation).

Die Frage, welche Anwendungen in Zukunft sehr hohe Datenströme erfordern werden, lässt sich abschließend nicht beantworten. Zur Zeit werden folgende Anwendungen als besonders intensiv angesehen:

- Server-Cluster (bis in den höheren Gigabit-Bereich, je nach Art des Clusters)
- Internet/Intranet-Anwendungen in Verbindung mit E-Commerce
- Betriebssysteme
- Ausdehnung des Server-Betriebssystems bis in die Endsysteme zur Erreichung einer verbesserten Total Cost of Ownership TCO

Grundsätzlich wird der Kapazitätsbedarf von Endsystemen und Servern natürlich von der Plattenleistung bestimmt. Während hier auf der Serverseite kein Ende der Entwicklung absehbar ist und Übertragungsraten von 1 Gigabit/s und mehr möglich erscheinen (abhängig von der I/O-Leistung des Betriebssystems, siehe auch I2O-Konzept für PC-Server u.a.), ist das Endsystem sicher zur Zeit eher leistungsschwach. Normalerweise sind von einem Endgeräte-PC unter Einbeziehung der bremsenden Wirkung des

Bild 2 Entwicklungsstufen der Netzwerk-Technologie



Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

Betriebssysteme nicht dauerhaft mehr als 30 bis 40 Mbit/s zu erwarten. Messungen in Laborumgebungen haben gezeigt, dass sich zur Zeit bis zu 8 PCs ohne Probleme 100 Mbit/s teilen könnten, ohne dass dies die Plattenleistung beeinträchtigt wird.

Von besonderer Bedeutung für die Planung der Netzwerk-Kapazität ist deshalb die Kommunikation zwischen Servern, im Spezialfall innerhalb eines Clusters. Hier muss einerseits ein schnell wachsender Bedarf und andererseits die schnelle Normalität von Gigabit-Adaptoren eingeplant werden. Dies spricht grundsätzlich dafür, zwischen den Servern eigene Netzwerke zu konzeptionieren. Das ist unabhängig von der Standortfrage zu realisieren.

Diese Überlegung hat eine sofortige Konsequenz für die Auswahl eines geeigneten Switch-Systems. Für den Backbone und den Verbindungs-Bereich der Server sollte die Schalteistung eines Switches nach der einfachen Formel

$$\frac{\text{Schalteistung Switch}}{\text{Anzahl Server} \times 1 \text{ Gbit/s}}$$

geplant werden. Neueste Untersuchungen und mathematische Berechnungen der ComConsult Technologie Information in Zusammenarbeit mit Dr. F.J. Kauffels bestätigen diese Formel*.

* Veröffentlichung in einem der nächsten Insider.

Bei der Auswahl eines Switch-Systems ist dabei zu beachten, dass die Hersteller in der Regel nur die Summenleistung eines Switches angeben. Dies ist in der Regel nicht mit der Transferleistung einer Einschubkarte identisch. Häufig können die Einschubkarten nur 2 Gbit/s zur zentralen Schaltmatrix übertragen. Wenn dann die Einschubkarte 8 Ports oder mehr in Gigabit-Technik hat, kann es theoretisch zu Stauwirkungen kommen. Natürlich ist dabei zu beachten, dass ein Server zur Zeit kaum in der Lage ist, dauerhaft Gigabit zu übertragen. Mindestens werden mehrere parallele Bus-Systeme erforderlich sein.

Diese Entwicklung ist noch nicht zu Ende. Die ersten Vorzeichen der Verfügbarkeit von 10 Gigabit-Systemen nehmen an Konkretheit zu. Deren Anwendung dürfte aber auf den Wide Area Bereich und Power-Cluster reduziert sein. Immerhin liefert 10 Gigabit/s Transferdaten und Zugriffszeiten im Bereich schneller RAM-Bausteine.

Verzögerungs-Planung

Viele Betreiber von 10 Mbit/s Ethernet oder 16 Mbit/s Token Ring werden die Überlegungen zur Kapazität mit einer Mischung aus Schrecken und der Erwartung, dass eine derartige Bedarfsentwicklung sie nicht betrifft, gelesen haben. Mit der Einbeziehung einer bewussten Planung der Verzögerung von Netzwerken zeigt sich, dass es weitere starke Argumente für eine Erhöhung

der Datenraten von Netzwerken gibt.

Die in einem Netz gegebene Verzögerung wird gemäß der Formel in Bild 3 bestimmt. Es zeigt sich, dass die Speicher- und Sendezeit einen elementaren Einfluss auf die erreichbare Verzögerung hat. Für die meisten Netzwerke kann mit der Faustformel

$$\text{Faustformel Verzögerung} = \text{Speicherzeit} \times 10 \text{ bis } 50$$

die gegebene Verzögerung abgeschätzt werden. Die Speicherzeit ist die Zeit, die benötigt wird, um ein Paket zu versenden oder im Empfang bzw. der Verarbeitung zu speichern. Sie beträgt in Abhängigkeit von der Datenrate für ein maximal langes Ethernet-Paket von 1500 Byte (12000 Bit):

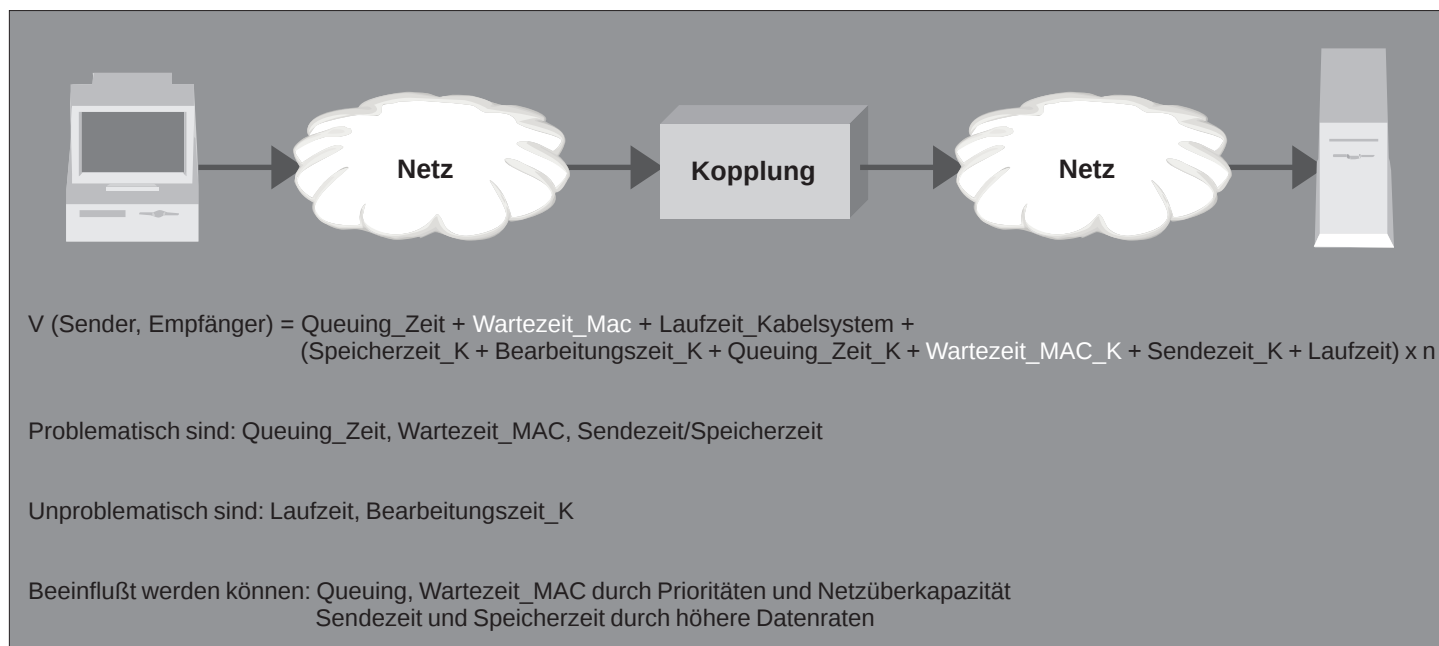
10 Mbit/s	=	1,2 ms
100 Mbit/s	=	0,12 ms
1000 Mbit/s	=	0,012 ms

Diese Zeiten nach der Faustformel multipliziert mit 5 bis 10 ergeben die in der Praxis typischen Verzögerungswerte für die Netzwerke mit den jeweiligen Übertragungsraten:

10 Mbit/s	=	10 bis 50 ms
100 Mbit/s	=	1 bis 10 ms
1000 Mbit/s	=	0,1 bis 1 ms

Bild 3

Verzögerung/delay



Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

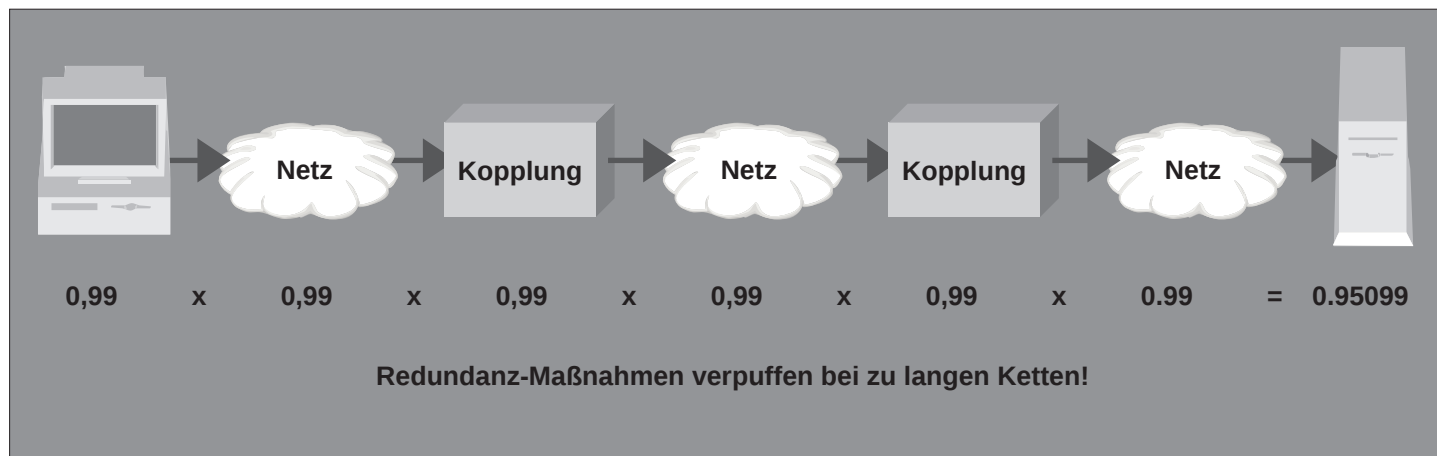


Bild 4 Ketten-Problematik

Vergleicht man dies mit dem Bedarf einzelner Anwendungsbereiche, so ergibt sich eine klare Planungssituation

Application Sharing Cluster = 0,1 bis 1 ms

Verteilte DB-Applikationen = 1 bis 10 ms

High End Video = unter 10 ms

Low End Video = unter 40 ms

Sprache = unter 100 ms
(PBX hat zur Zeit ca. 40 ms)

Damit ist klar: die notwendigen Verzögerungszeiten für eine Reihe von Anwendungsgebieten lassen sich nur durch Kapazitätsausbau auf 100 und 1000 Mbit/s erreichen. Wer also argumentiert "wir brauchen kein Gigabit, so viele Daten haben wir gar nicht", der liegt ggf. völlig daneben, weil er den

Verzögerungs-Bedarf nicht kalkuliert hat. Das Ergebnis wäre eine stark schwankende und zum Teil auch untragbar schlechte Performance im Cluster-Einsatz oder bei Datenbank-Anwendungen mit SQL übers Netzwerk.

Planung der Verzögerungs-Varianz

Die Übertragung von Sprache und Video gestattet keine schwankenden Verzögerungswerte. Um einen Verlust von Information zu vermeiden, müssen die Puffergrößen nach der maximal zu erwartenden Verzögerung (oder 99-Prozent-Wert) ausgelegt werden. Größere Puffer erhöhen aber die Laufzeiten. Höhere Laufzeiten werden bei bidirektionaler Kommunikation für den Endbenutzer ab einer Größenordnung von über 300 ms wie ein Satelliteneffekt wirken.

Ansonsten gelten die gleichen Überlegun-

gen wie bei der Planung der mittleren Verzögerung.

Planung der Verfügbarkeiten

Spätestens seit der Einbeziehung von "mission critical" Servern in unsere Netze in Kombination mit der Zentralisierung von Serverleistung (Stichwort Re-Hosting) ist die bewusste Planung der Verfügbarkeit eines Netzwerks ein zwingendes Muss. Im Datenbank-Umfeld gelten die in Tabelle 1 dargestellten Verfügbarkeits-Klassen.

Berücksichtigt man den in Bild 4 dargestellten Multiplikator-Effekt in Netzwerken, der sich aus der unvermeidbaren Kettenbildung im Netzwerk-Design ergibt, so wird deutlich, dass die Erwartungen der Betreiber von Datenbanken und verteilten Applikationen im Netz ohne besondere Zusatzmaßnahmen nicht erfüllt werden können.

Tabelle 1 Verfügbarkeitsklassen von Server-Systemen

	Verfügbarkeit	Typische Ausfalldauer	Ausfallzeit pro Jahr
Unterbrechungsfrei	100 %	Keine	Keine
Fehlertolerant	99,9999 %	Taktzyklen	0,5 Minuten
Failover (Cluster)	99,999 %	Sekunden	5 Minuten
Fehlerfest	99,99 %	Sekunden/Minuten	max 53 Minuten
Hochverfügbar	99,90 %	Minuten	max 8,7 Stunden
Standard1	99,50 %	Minuten bis Stunden	2 Tage
Standard 2	99 %	Stunden	3,5 Tage

Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

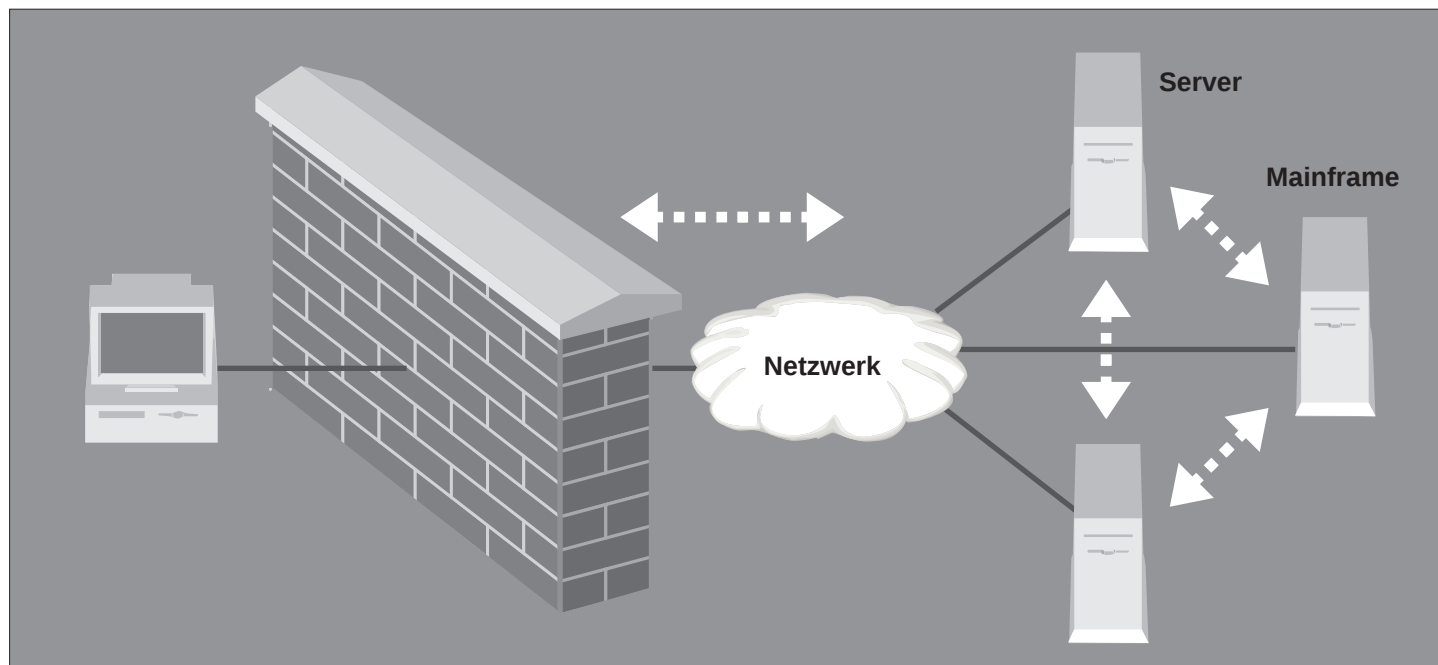


Bild 5 Client-Server-Technologien schaffen Technologie-Verbundsituationen mit erheblichen Abhängigkeiten

Unter Berücksichtigung der typischen Verbundsituation einer modernen DV-Architektur, in der mehrere Server und Großrechner zusammen eine Applikation umsetzen (inklusive Anmeldung, Namensauflösung usw.) wie in Bild 5 dargestellt, wird klar, wie hoch die bestehenden Risiken für instabile Applikationen sind.

Die Erfahrungen des ComConsult Troubleshooting-Teams belegen, dass viele Störungen handwerkliche Ursachen haben:

- Kabel / Stecker / Rangiersysteme

- Fehlkonfigurationen von Switch-Parametern
- Fehlkonfigurationen von Endsystemen mit Rückwirkung auf den Server
- Unsachgemäß geplante Redundanz (kein Test, Prinzip Hoffnung)

Um dem Verfügbarkeitsbedarf von "mission critical" Anwendungen zu entsprechen, werden mindestens folgende Maßnahmen ins Auge gefasst werden müssen:

- Redundanz durch parallele Wege, nicht durch hot-stand-by

- Hochwertige Kabelsysteme (Risikofaktor RJ 45, neue Anforderungen von Gigabit an Twisted Pair beachten, am besten Gigabit nur auf LWL-Basis realisieren)

- Redundante Switch-Architekturen

- Netzwerk-Management mit leistungsstarker Überwachung und Steuerung

Fasst man die diskutierten Planungselemente für eine moderne Netzwerk-Auslegung zusammen, so ergibt sich Tabelle 2.

Tabelle 2

Leistungsparameter und Anwendungsbereiche

	Kapazität	Delay	Delay-Varianz	Verfügbarkeit
File-Server	X			(X)
Terminal-Emulation				(X)
Applikations-Server		X		X
Applikations-Cluster	X	X		X
Web-Server	(X)			(X)
Sprachübertragung		X	X	X
Videoübertragung	X	X	X	X

Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

Service-Level-Design

Die Planung von Netzwerken im Service-Level-Design trägt der Tatsache Rechnung, dass die an der kritischsten Stelle gewünschten Leistungswerte nicht auf das gesamte Netzwerk übertragen werden können. Das Endgerät hat in der Regel wesentlich niedrigere Leistungs-Erwartungen als ein Server-Cluster.

Aus diesem Grund geht man im Service-Level-Design davon aus, dass ein Netz aus Bausteinen unterschiedlicher Leistung zusammen gesetzt wird. Die Übergänge zwischen den Bausteinen werden durch Switch-Systeme einer ausreichenden Leistungsklasse realisiert.

Die in Tabelle 3 dargestellten Bausteine werden im Service-Level-Design benötigt und müssen im Rahmen der Planung erarbeitet werden.

Im Rahmen der Nutzung dieser Bausteine ergibt sich ggf. der zusätzliche Bedarf nach Feinparameterisierung. Dieser entsteht dann, wenn eine gegebene Netzwerk-Infrastruktur kurzzeitig überlastet werden kann, aber Applikationen existieren, die auch in dieser Überlast-Situation eine definierte Mindest-Leistung erbringen müssen.

Service-Level-Design erlaubt folgende Strategien zur Vermeidung bzw. bewussten Steuerung derartiger Überlasten:

- Bewusste Realisierung von Überkapazität
 - Mindestens 100 Mbit/s, im Backbone und an Servern mehr
- Lastverteilung / Loadsharing
 - Link Aggregation
 - OSPF
- Bandbreiten-Vermeidung
 - Typischerweise Multicast-Steuerung
 - Auch Namens-Auflösung
- Priorisierung
 - IEEE 802.1P
 - TOS
- Dynamische Reservierung
 - RSVP
 - ???

Tabelle 3

Bausteine im Service-Level-Design

Endgeräte-Versorgung	Dedicated	
	Shared	
Steig-/Distributionszone	Pro Workgroup eine Anbindung an HVT	
	Pro Workgroup redundante Anbindung an HVT	
	Pro Bereichsverteiler konzentrierte Versorgung durch High-Speed	
	Pro Bereichsverteiler redundante High-Speed-Versorgung	
Hauptverteiler-Gebäude-Distributionszone	Immer modular (Klimatisierung unbedingt beachten)	
	Nicht-redundant	
	Redundant	
	Multi-media-fähig	
Gebäude-lokaler Server	Nicht redundant	
	Redundant	Hotstand by
		Parallel / Load Sharing
		Skalierbar
Backbone	Redundant	Hotstand by
		Load Sharing
		Skalierbar
	Entfernungsbausteine	Kleiner 200 m
		Kleiner 500 m
		Kleiner 2 km
		Kleiner 5 km
	Multi-media-fähig	
Zentraler Server	Nicht redundant	
	Redundant	Hotstand by
		Parallel / Load Sharing
		Skalierbar

Änderungen im Netzwerk-Design: Service-Level-Design contra Policy-based-Networking

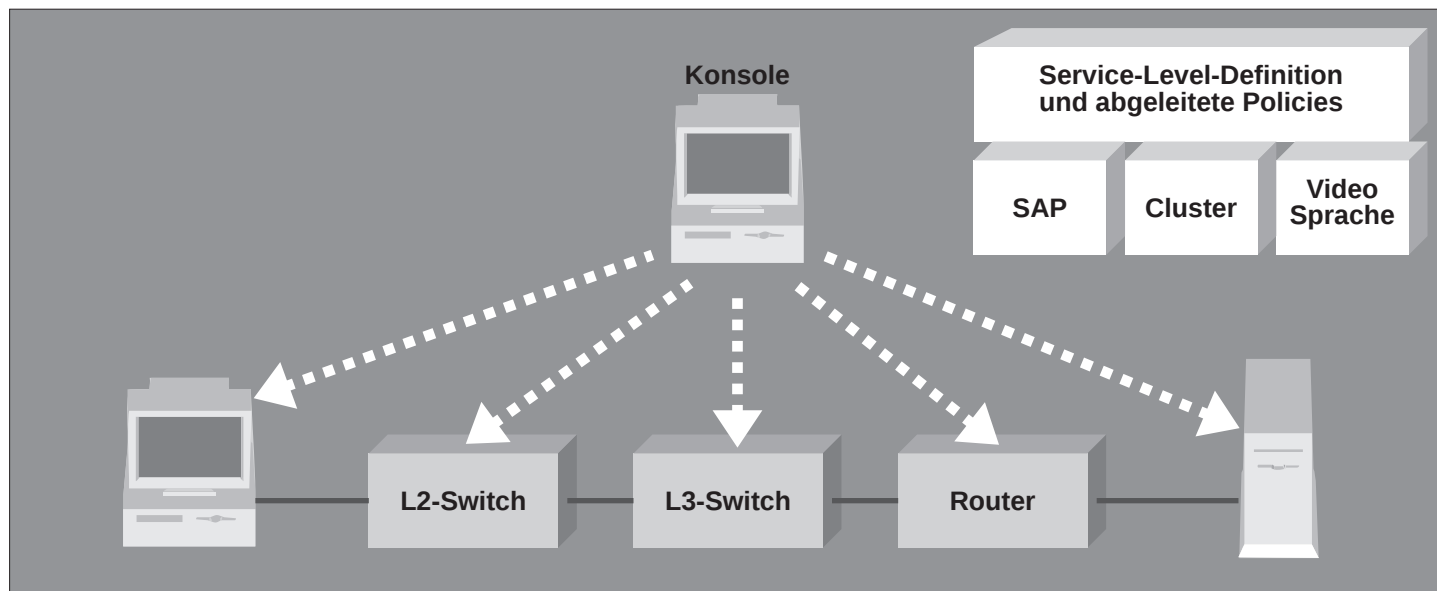


Bild 6 Zentraler Steuerungs-Server

Sinnvollerweise wird man auch genau in der Reihenfolge der Aufzählung vorgehen. Dies bedeutet, dass es wenig Sinn macht, mit Prioritäten zu arbeiten, wenn nicht zuvor durch hohe Bandbreiten ausreichend Reserven angelegt worden sind. Für viele Anwender wird auch die Strategie der Überkapazitäten bedeuten, dass für mehrere Jahre keine der anderen Strategien umgesetzt werden muss.

Jede dieser Strategien wird durch eine Redundanzplanung begleitet. Dabei wird grundsätzlich unterschieden zwischen

- Layer 2 Auslegungen
 - Spanning Tree nach IEEE 802.1D letzte Fassung aus 1998
 - Link Aggregation
 - Multicast-Steuerungen
 - Layer-2-Priorisierung
- Layer 3 Auslegungen
 - Basis in der Regel IP-Switching
 - HSRP / VRRP
 - OSPF
 - Layer-3/4-Priorisierung
 - RSVP

Auch hier gilt, dass die Netzwerk-Auslegung in Bausteinen erfolgt. So wird sich die Layer-3-Planung in der Regel nicht auf die Etage beziehen.

Policy-based Networking

Dieser Design-Ansatz geht davon aus, dass an einer zentralen Steuerungs-Konsole Stra-

tegien für den Netzwerk-Betrieb definiert werden (zum Beispiel Antwortzeiten, Kapazitäten für bestimmte Applikationen). Von einer Konsol-Applikation werden diese Strategien in konkrete Konfigurations-Parameter der Switches umgesetzt und in die Switches geladen.

Also sind erforderlich:

- Mindestens eine Konsole, diese sollte zur Erreichung von Skalierbarkeit nicht in einem Switch integriert sein
- Ein Kommunikations-Protokoll zur Übertragung der Konfigurations-Parameter
- Ein Agent in den Switch-Systemen zur Umsetzung der Einstellungen
- Ein Berichtswesen zur Verfolgung der Einhaltung der strategischen Ziele, da deren Abbildung indirekt sein kann (Ziel Antwortzeit wird ggf. durch Priorisierung oder Reservierung umgesetzt, damit ist aber keine Antwortzeit-Garantie verbunden)

Policy-based-Networking ist somit kein direkter Widerspruch zum Service-Level-Design. Allerdings verbirgt der Policy-Ansatz die wirklich durchgeführte Parametrierung. Dies ist nicht in jedem Fall sinnvoll und machbar. OSPF beispielsweise wird immer individuell konfiguriert werden müssen. Im Endeffekt suggeriert das Policy-Prinzip dem Betreiber eine Vereinfachung seines Lebens.

Zur Zeit haben die Policy-Ansätze der verschiedenen Hersteller mehr Nachteile als

Vorteile und müssen als Marketing-Gag angesehen werden:

- es gibt keinen einheitlichen Ansatz, zur Zeit existieren mindestens 9 unterschiedliche Ansätze zur Umsetzung. Dies reicht von Trivialitäten bis zu komplexen Ansätzen
- die Hersteller treiben ihre Kunden in eine erhöhte Hersteller- und Produkt-Abhängigkeit
- die Botschaft ist eine Illusion, wer heute ernsthaft ein Netz betreiben und analysieren will, benötigt das entsprechende Detailwissen. Dies kann nicht durch einen Policy-Server ersetzt werden. Maximal ist dieser als Kopierhilfe sinnvoll.

Fazit

Design, Betrieb und Analyse von Netzwerken erfordern das entsprechende Technologie-Wissen. Dies kann nicht durch eine schöne Oberfläche ersetzt werden. Gerade die Erkenntnis, dass ein Netz aus Bausteinen unterschiedlicher Leistung optimiert werden muss, führt zu sehr individuellen Einstellungen der einzelnen Switch-Systeme. Zur Zeit gibt es deshalb keine Alternative zu einem bewussten Service-Level-Design. Policy-Server sollten nur auf der Basis eines Hersteller-übergreifenden Standards eingesetzt werden und sollten vorher einen erheblichen Reifegrad erreicht haben. Dies wird in den nächsten 2 Jahren nicht der Fall sein. Bis dahin ist Policy-Networking viel heiße Luft um Nichts.

Anzeige

Ausbildung zum ComConsult Certified Network Engineer



**Mit
! Palm IIIx !**

Lokale Netze für Einsteiger

Einzelpreis: DM 3.500,-- (EUR 1.789,52) zzgl. MwSt.

- ▶ 25.10. - 29.10.99 Aachen
- ▶ 22.11. - 26.11.99 Köln
- ▶ 06.12. - 10.12.99 Aachen
- ▶ 24.01. - 28.01.00 Aachen
- ▶ 14.02. - 18.02.00 Aachen
- ▶ 13.03. - 17.03.00 Aachen
- ▶ 10.04. - 14.04.00 Aachen
- ▶ 15.05. - 19.05.00 Aachen
- ▶ 05.06. - 09.06.00 Aachen
- ▶ 26.06. - 30.06.00 Aachen

Internetworking

Einzelpreis: DM 3.790,-- (EUR 1.937,80) zzgl. MwSt.

- ▶ 14.02. - 18.02.00 Aachen
- ▶ 27.03. - 31.03.00 Aachen
- ▶ 05.06. - 09.06.00 Aachen

Neue Ethernet Technologien

Einzelpreis: DM 3.690,-- (EUR 1.886,67) zzgl. MwSt.

- ▶ 27.09. - 01.10.99 Aachen
- ▶ 29.11. - 03.12.99 Aachen
- ▶ 27.03. - 31.03.00 Aachen
- ▶ 22.05. - 26.05.00 Aachen

TCP/IP und SNMP

Einzelpreis: DM 2.990,-- (EUR 1528,76) zzgl. MwSt.

- ▶ 15.09. - 17.09.99 Hamburg
- ▶ 25.10. - 27.10.99 Königswinter
- ▶ 13.12. - 15.12.99 Hamburg

Nutzen Sie unsere Paketpreise!

Zwei 5-tägige Seminare und das 3-tägige Seminar zum Paketpreis von DM 9.150,-- oder EUR 4.678,32 (statt mind. DM 9.990,--) zzgl. MwSt.

Drei 5-tägige Seminare zum Paketpreis von DM 9.600,-- oder EUR 4.908,40 (statt mind. DM 10.500,--) zzgl. MwSt.

Sonderaktion 1999: Palm IIIx

Drei 5-tägige Seminare und das 3-tägige Seminar zum Paketpreis von DM 11.900,-- oder EUR 6.084,37 (statt mind. DM 13.490,--) zzgl. MwSt.

Bei Buchung des Paketes erhalten Sie als Bestandteil des Seminarpaketes zur Planung Ihrer Akademie-Termine den 3Com »Palm IIIx«

ComConsult
Akademie