

Vorlesung Kommunikationssysteme
Wirtschaftsinformatik
VWA Stuttgart

Armin Herzer¹

11. März 2009

¹armin@herzer-online.de

Inhaltsverzeichnis

1	Internet: Wie das Netz der Netze entstanden ist	7
1.1	Von der Idee zum ersten Netzwerk	9
1.2	Globale Ausdehnung	10
1.3	Gegenwart und Zukunft	11
1.4	Größe von Netzwerken	12
2	Schichtenmodelle: Wie man ein Netzwerk beschreiben kann	13
2.1	Einleitung	15
2.2	Das OSI-Schichtenmodell	16
2.2.1	Das Prinzip der Datenkapselung	17
2.3	Das TCP/IP-Schichtenmodell	18
2.3.1	Aufgaben der einzelnen Schichten:	19
2.4	Anwendungsmöglichkeiten der Schichtenmodelle	20
3	OSI-Schichtenmodell: Eine detaillierte Betrachtung	21
3.1	Layer 1: Bitübertragungsschicht	23
3.1.1	Übertragungsmedien	24
3.1.1.1	Kupfer	24
3.1.1.1.1	Kabel-Klassen	24
3.1.1.1.2	Twisted Pair: 10Base-T und 100Base-TX (Fast-Ethernet)	25
3.1.1.1.3	Dämpfung	26
3.1.1.2	Optische Medien	27
3.1.1.2.1	Lichtwellenleiter	27
3.1.1.2.2	Multi- und Monomode	28
3.1.1.3	Trägerwelle	29
3.1.1.3.1	WLAN	29
3.1.1.4	Weitere Übertragungsmöglichkeiten	30
3.1.2	Netzwerkgeräte	31
3.1.2.1	Repeater und Hubs	31
3.1.2.2	Transceiver	32
3.1.3	Physische NetzwerkTopologie	33
3.1.3.1	Bus-Topologie	33
3.1.3.2	Stern-Topologien	34
3.1.3.3	Ring	35
3.1.3.4	Vermascht	36

3.1.4	Wichtige Begriffe	37
3.1.4.1	Autonegotiation	37
3.1.4.2	Duplex	38
3.1.4.3	Kollisionsdomäne	39
3.2	Layer 2: Sicherungsschicht	40
3.2.1	Wichtige Begriffe	41
3.2.1.1	LAN-Schichtenmodell	41
3.2.1.2	CSMA/CD	42
3.2.1.3	MAC-Adresse	43
3.2.1.4	Broadcastdomäne	44
3.2.2	Logische Netzwerktopologie	45
3.2.2.1	Logische Topologien	46
3.2.3	Netzwerkgeräte	47
3.2.3.1	Netzwerkkarten	47
3.2.3.2	Bridges	48
3.2.3.3	Switches	49
3.2.3.3.1	Geräteklassen: Managed und Unmanaged Switches	50
3.2.3.3.2	Switching Methoden	51
3.2.3.3.3	Spanning Tree	52
3.2.3.3.4	Zusammenfassung	53
3.3	Layer 3: Vermittlungsschicht	54
3.3.1	IP	54
3.3.1.1	Das überholte Klassenschema der IP-Adressen	54
3.3.1.2	CIDR und VLSM	55
3.3.1.3	Probleme des IPv4-Adressraums und Gegenmaßnahmen	56
3.3.1.3.1	Einleitung:	56
3.3.1.4	Private und besondere IP-Adressen	57
3.3.1.5	Multicast-Adressen	58
3.3.1.6	Das IP-Protokoll	59
3.3.2	Netzwerkgeräte	60
3.3.2.1	Router	60
3.3.2.1.1	Mögliche Entscheidungskriterien für die beste Route	61
3.3.2.2	Routingprotokolle	62
3.3.2.2.1	RIP	62
3.3.2.2.2	Vergleich von RIPv1 und RIPv2	63
3.3.2.2.3	OSPF	64
3.3.2.2.4	Vergleich von Distance Vector und Link State	65
3.3.2.3	Default Gateway	66
3.4	Layer 4: Transportschicht	67
3.4.1	TCP	68
3.4.2	UDP	69
3.4.2.0.1	Allgemein:	69
3.4.2.0.2	Funktionsweise:	69
3.4.3	Vergleich TCP/UDP	70
3.4.3.0.3	Vor- und Nachteile:	70
3.5	Layer 5: Kommunikationssteuerungsschicht	71

3.5.1	Ports und Portnummern	72
3.6	Layer 6: Darstellungsschicht	73
3.7	Layer 7: Anwendungsschicht	74
3.7.1	Network File System - NFS	75
3.7.2	Server Message Block - SMB (Samba)	76
4	Netzwerkanwendungen	77
4.1	VLAN	79
4.1.1	Funktionsweise	80
4.1.2	VLAN Tagging	81
4.2	Firewall	82
4.3	DMZ	83
4.4	DHCP	84
4.5	PAT	85
4.6	MD5	86
4.7	SNMP	87
4.8	Kryptographische Verfahren	88
4.8.1	Einleitung	88
4.8.2	Symmetrische Verschlüsselung	89
4.8.3	Asymmetrische Verschlüsselung	90
4.8.4	Hybride Verschlüsselung	91
4.9	Datensicherheit und Verschlüsselung im WLAN	92
A	Zusätzliche Folien	93
B	Weitergehende Informationen	111
B.1	Bolt Beranek and Newman	113
B.2	Robert Metcalfe	113
B.3	Internet	113
B.4	Usenet	113
B.5	MILNET	113
B.6	BITNET	114
B.7	RFC	114
B.8	Wi-Fi	114

Kapitel 1

Internet: Wie das Netz der Netze
entstanden ist

1.1 Von der Idee zum ersten Netzwerk

Im Herbst des Jahres 1957 gelang es der damals noch existierenden UdSSR mit dem Sputnik den ersten künstlichen Satelliten auf eine Umlaufbahn um die Erde zu schicken. Dieser "Sputnik-Schock" wurde in den USA als Niederlage im Wettlauf um die "Beherrschung des Weltalls" empfunden. Die offensichtliche technische Überlegenheit der Sowjetunion konnte deshalb auch auf militärischem Gebiet angenommen werden. Als Konsequenz aus dieser Annahme wurde im darauffolgenden Jahr die ARPA¹, als Unterabteilung des DoD² gegründet. Ziel war, einen Weg zu finden der es ermöglichte, im "Katastrophenfall" ein Steuer- und Kommandonetz in den USA aufrecht zu erhalten.

Es sollten aber noch mehr als 10 Jahre vergehen (in denen viele Ansätze entwickelt wurden) bis ein erstes kleines Testnetzwerk seinen Betrieb aufnehmen konnte.[1, S.68], [2, S.68] sowie [7].

Gleichzeitig mußte aber auch ein Host-zu-Host-Protokoll entwickelt werden, damit die einzelnen Knotenpunkte überhaupt kommunizieren konnten. Diese Aufgabe erwies sich als so schwierig, daß Ende 1969 mit telnet³ nur eine minimale Notlösung präsentiert werden konnte.

1957 Verschiedene Ideen für ein **▶ fehlertolerantes** Design werden diskutiert. AT&T lehnt ein digitales Paketvermittlungssystem als nicht realisierbar ab.

1965 Die ARPA will selbst ein Kommunikationsnetzwerk erstellen. In Großbritannien wurde unabhängig davon ein ähnliches System im Labormaßstab erfolgreich getestet.

1966 **▶ Angebotsausschreibung** zur Erstellung des Netzes: BBN bekommt mit dem **▶ IMP** den Zuschlag und kann sich damit gegenüber Firmen wie z.B. IBM durchsetzen.

1967 In den letzten vier Monaten dieses Jahres waren bereits **▶ 4 Standorte** untereinander verbunden: Über große Entfernungen und mit unterschiedlichen Hosts! Der erste RFC "Host Software" von Steve Crocker wurde am 07.05.69 publiziert.

Zu Beginn der 70er Jahre begann auch die Entwicklung des Betriebssystems **▶ Unix** das sich zum Ende des Jahrzehnts als Standard für vernetzte Rechner etablieren sollte. Parallel dazu wurde die Programmiersprache **▶ C** entwickelt und begann systemspezifische Assemblerprogramme abzulösen.

¹Advanced Research Projects Agency

²Department of Defense

³telnet ermöglicht einen remote console login, d.h. man kann sich auf einer entfernten Maschine einloggen und im Textmodus so darauf arbeiten, als säße man direkt vor ihr.

1.2 Globale Ausdehnung

Zu Beginn der siebziger Jahren wurde an der Universität von Hawaii eine Technik entwickelt, bei der sich mehrere Benutzer ein Datenübertragungsmedium teilen. Jeder Benutzer sollte, wann immer er wollte, Daten senden können. Dabei mußte berücksichtigt werden daß, wenn zwei oder mehr Stationen parallel senden, sich die Datenpakete gegenseitig "zerstören" (sogn. Kollisionen) und deshalb nicht mehr ausgewertet werden können. Die technische Umsetzung dieser Idee wurde in Bodenfunktechnik realisiert und ALOHA genannt. 1972 wurde dieses Netz mit dem ARPANET verbunden. Etwa zur gleichen Zeit führte Ray Tomlinson, der bei BBN ein Jahr zuvor ein E-Mail Programm entwickelt hatte, das @-Zeichen (mit der Bedeutung at) ein.

Nachdem 1979 das USENET⁴, in Betrieb gegangen war, folgte mit dem BITNET⁵ 1981 ein weiteres großes Netzwerk. 1983 wurde schließlich in Europa mit dem EARN⁶ eine dem BITNET ähnlich Struktur geschaffen. Beide Netze konnten so verbunden werden und zu Beginn der 90er Jahre kommunizierten ca 3500 Rechner in über 1400 Organisationen über dieses Netzwerk miteinander.

- | | |
|------|--|
| 1972 | Das @-Symbol taucht zum ersten Mal in Emails auf |
| 1973 | Bob Metcalfe beschreibt in seiner PhD Thesis das Ethernet . Das ARPANET hat sich inzwischen über die Weltmeere ausgedehnt! |
| 1974 | V. Cerf u.a. publizieren als RFC 675 die erste technische Beschreibung von TCP/IP. |
| 1983 | Stuttgart wird vernetzt und das ARPANET wird aus Sicherheitsgründen in die Teile ARPANET und MILNET aufgeteilt. TCP/IP wird zum Standardprotokoll. Desktop-Workstations kommen in Mode und werden in den folgenden Jahren zu einer wahren Host-Explosion führen. |
| 1984 | DNS wird eingeführt, die Zahl vernetzter Hosts liegt bei ca. 1000. Das OSI-Schichtenmodell wird von der ISO veröffentlicht. |
| 1985 | Die IEEE veröffentliche Standards der 802.X Protokolle (Ethernet, Token Ring...) |
| 1988 | Anfang November gräbt sich der Internet Worm durch das Netz und legt $\approx \frac{1}{10}$ der ca. 60000 Hosts lahm |

⁴Schwerpunkt Newsgroups

⁵Schwerpunkt: Großrechnerverbund von Forschung und Wissenschaft in den USA

⁶IBM Europa stellte dazu einen Gateway zur Verfügung

1.3 Gegenwart und Zukunft

Ab 1989 wurden von Tim Berners-Lee am CERN in Genf neben URL und HTML auch das darauf aufsetzende HTTP entwickelt. Damit war der Grundstein zum WWW wie wir es heute kennen gelegt. Rasanten Auftrieb erhielt das Internet ab 1993, als der erste grafikfähige Webbrowser Mosaic⁷ veröffentlicht und zum kostenlosen Download angeboten wurde. Damit wurde das Internet auch für Laien benutzbar.

Die tatsächliche Größe des Internets ist nicht bestimmbar⁸: Jeder Rechner, der über einen Netzwerkanschluß verfügt kann zu jeder Zeit Teil des Internets werden. Die maximale Größe wird jedoch durch den beschränkten Adressraum von 32 Bit limitiert. In Deutschland verfügen 68%⁹ der Erwachsenen über einen Internetanschluss.

Mitte der 90er Jahre begann das Internet immer schneller zu wachsen und war zu diesem Zeitpunkt auch schon größeren Teilen der (nicht-akademischen) Bevölkerung ein Begriff.

Kupfergebundene Internetanbindungen sind heute als Standardzugangsarten für Firmen und Privatpersonen anzusehen. Mehr und mehr setzen sich aber auch Glasfaserkabel durch, um dem ständig wachsenden Anspruch an Geschwindigkeit neuen Techniken (z.B. VoIP) und immer größeren Datenmengen gerecht zu werden.

1990	Das ARPANET wird in das Internet überführt
1991	WWW , ein Browser mit dem Ziel Daten im Netz in einer einheitlichen Sprache austauschen und darstellen zu können wird am CERN der Öffentlichkeit vorgestellt.
1994	Die ersten Internetshops erscheinen
1995	WWW überholt im März ftp als datenintensivsten Dienst auf Paket-Basis und im April auf Byte Basis.
1996	Die Browser-Schlacht ist mittlerweile voll entbrannt
1998	Netscape entscheidet sich den Quellcode des Navigators freizugeben. Mit diesem Schritt wird die Ära der Open Source Software eingeläutet.

Das rasche Wachstum des Internets, hier einmal als [log.](#) Anzahl der Hosts dargestellt, liefert u.a. **eine** Begründung dafür, warum in den nächsten Jahren das Adressierungsschema von IPv4 nach [IPv6](#) umgestellt werden soll.

⁷Marc Andreessen gründete die Firma Mosaic Communication Corporation, später Netscape Communication

⁸Im Jahr 2005 besaßen ca. 2% der Erdbbevölkerung einen Internetanschluss

⁹Stand 2005/2006

1.4 Größe von Netzwerken

PAN: Ein Netz, welches Kleingeräte wie PDAs oder Mobiltelefone ad-hoc, entweder drahtgebunden oder per Funk, auf- und abbauen können.

LAN: Ein Computernetzwerk innerhalb eines räumlich begrenzten Bereiches. LANs sind als feste Installation kabelgebunden und intern bei Firmen zu finden. Zunehmend werden sie auch in privaten Haushalten aufgebaut. Verwendet keine öffentlichen Leitungen.

MAN: Ein breitbandiges Netzwerk das meist in Glasfasertechnologie realisiert und überwiegend in ringförmiger Struktur aufgebaut ist. Die wichtigsten Bürozentren einer Großstadt können so miteinander verbunden werden. Äußerste Ausdehnung max. 50km.

WAN: Ein Netzwerk, in dem Computer und Peripheriegeräte über große Entfernungen verbunden sind. Typischerweise sind dabei verschiedene LAN's, miteinander verbunden.

GAN: Ein weltweites System von miteinander verbundenen Computern. In der Regel ein Netzwerk das über Satellitenverbindungen zustandekommt.

Distance Between CPUs	Location of CPUs	Name
0.1 m	Printed circuit board Personal data asst.	Motherboard Personal Area Network (PAN)
1.0 m	Millimeter Mainframe	Computer Systems Network
10 m	Room	Local Area Network (LAN) Your classroom
100 m	Building	Local Area Network (LAN) Your school
1000 m = 1 km	Campus	Local Area Network (LAN) Stanford University
100,000 m = 100 km	Country	Wide Area Network (WAN) Cisco Systems, Inc.
1,000,000 m = 1,000 km	Continent	Wide Area Network (WAN) Africa
10,000,000 m = 10,000 km	Planet	Wide Area Network (WAN) The Internet
100,000,000 m = 100,000 km	Earth-moon system	Wide Area Network (WAN) Earth and artificial satellites

© Cisco Systems, Inc. 1999

Abbildung 1.1: Verschiedene Größen von Netzwerk und ihre Bezeichnung

Kapitel 2

Schichtenmodelle: Wie man ein Netzwerk beschreiben kann

2.1 Einleitung

Um verstehen zu können wie ein Netzwerk funktioniert und um die Datenübertragung bzw. Netzwerktechnik zu "standardisieren" wurde eine schematische Beschreibung aller dafür notwendigen Bedingungen entworfen. Diese Beschreibung nennt man ganz allgemein **Schichtenmodell**. In der Netzwerktechnik gibt es zwei verschiedene Schichtenmodelle:

- OSI-Schichtenmodell:
- TCP/IP-Schichtenmodell

Das TCP/IP-Schichtenmodell entstand parallel zur Entwicklung der Netzwerktechnik, ist also in der Praxis über viele Jahre hinweg entstanden. Das OSI-Schichtenmodell ist ein eher künstliches Gebilde, mit dem man versucht hat einige Schwächen des TCP/IP-Schichtenmodells zu beseitigen, es ist eher theoretische geprägt. Ganz allgemein kann man sagen, daß das OSI-Schichtenmodell versucht, schärfe Grenzen zwischen den einzelnen Schichten festzulegen und die Funktionen einer Schicht genauer zu beschreiben.

In der Praxis muß man beide Schichtenmodelle kennen. Beide haben Vor- und Nachteile und keines von beiden ist leichter zu verstehen oder einfacher zu handhaben: Vielmehr ist es so, daß man von Fall zu Fall entweder das eine Modell oder das andere Modell heranziehen wird.

Vergleichende Übersicht von TCP/IP- und OSI-Modell

OSI-Schichtenmodell		TCP/IP-Modell				
7	Anwendungsschicht	4	Anwendungsprotokolle		Anwendungen und Dienste, Clients und Server	
6	Darstellungsschicht		HTML, FTP, Telnet, SMTP, SMB, ...	DNS		NFS, NIS, SNMP
5	Sitzungsschicht					
4	Transportschicht	3	Transmission Control Protocol (TCP) User Datagram Protocol (UDP)		Netzwerk-API im Kernel Sockets, Ports	
3	Netzwerkschicht	2	Internet Control Message Protocol (ICMP)		Router	
			Internet Protocol (IP)			
			Address Resolution Protocol (ARP)			
2	Verbindungsschicht	1	Ethernet, WLAN	Logical Link Control (LLC) Media Access Control (MAC)	Hubs, Switches, ...	
1	Physikalische Schicht		Hardware, Übertragungsmedien, Modulationen		Kabel, Funk, Laser, ...	

http://www.linux-magazin.de/heft_abo/ausgaben/2007/09/verkehrsregelung?category=364

Abbildung 2.1: Erstellen und Auspacken eines Ethernet-Frames.

2.2 Das OSI-Schichtenmodell

Das OSI-Schichtenmodell ist im Grunde genommen als eine Vorschrift anzusehen die beschreibt, wie Daten so umgewandelt werden, daß sie als digitale Signale in einem Netzwerk transportiert werden können.

Dazu definiert es in der untersten Schicht 1 reine Hardwareaspekte und endet in der obersten Schicht 7 bei rein softwaretechnischen Aspekten. Die dazwischen liegenden Schichten übernehmen ganz exakt umrissenen Aufgaben, die nötig sind, um die oberste Schicht (also z.B. das Arbeiten mit einem Web-Browser, d.h. darstellen von Text und Grafik auch mit länderspezifischen Sonderzeichen) mit der untersten Schicht zu verbinden. Dazu kommunizieren immer zwei benachbarte Schichten miteinander bis das Modell von oben nach unten oder umgekehrt durchlaufen wurde.

Das OSI-Schichtenmodell

Schicht	Bezeichnung	Layer Name	Implementierung
7	Anwendungsschicht	Application	Dateisystem-
6	Darstellungsschicht	Presentation	treiber, Anwen-
5	Kommunikationssteuerschicht	Session	dungsprogramme
4	Transportschicht	Transport	Transport-
3	Vermittlungsschicht	Network	protokolle
2	Sicherungsschicht	Data Link	Treiber für
1	Bitübertragungsschicht	Physikal	Netzwerkkarten

Beispiel

Alle durstigen Kölner trinken viel schäumendes Bier

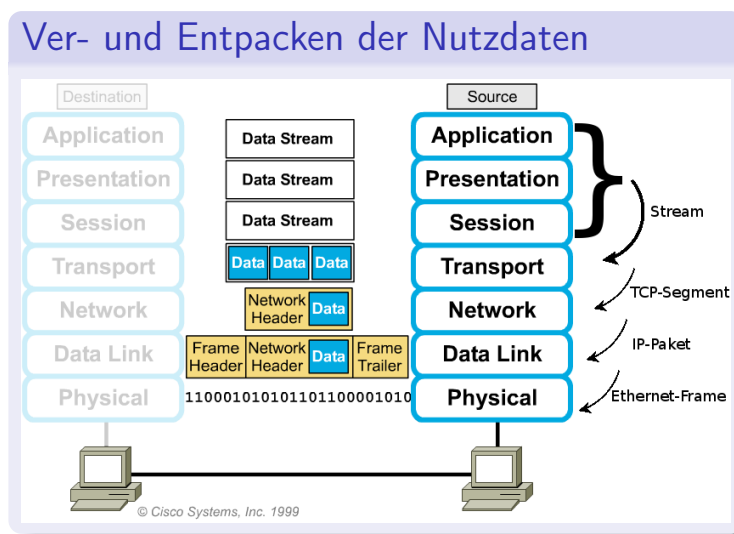
All People Seem To Need Data Processing

Abbildung 2.2: Mit dem OSI-Schichtenmodell sind Netzwerke leichter zu verstehen.

2.2.1 Das Prinzip der Datenkapselung

Beim Durchlaufen des Schichtenmodell werden die Daten immer zwischen zwei benachbarten Schichten ausgetauscht und modifiziert. Eine tieferliegende Schicht (z.B. IP auf Layer 3) interpretiert dabei alles, was von der darüberliegenden Schicht (TCP auf Layer 4) kommt, als Daten. Den TCP-Header von Schicht 4 erkennt Schicht 3 nicht! Jede Schicht erkennt nur ihre eigenen Header – beim Empfang der Daten (und vor dem weiterreichen an Schicht 4) muß er ja wieder entfernt werden! Es haben sich folgende Begriffe eingebürgert:

- Alle Daten, die von Netzwerkanwendungen kommen heißen **Stream** (bei TCP-Anwendungen wie Telnet oder FTP) bzw. **Message** (bei UDP- Anwendungen, wie NFS oder DNS)
- Wenn der TCP- bzw. UDP-Header hinzugefügt wurde, spricht man von einem **TCP-Segment** bzw. **UDP-Datagram**
- Ist auch der IP-Header hinzugefügt, so nennt man diese Struktur **IP-Packet**
- Mit dem Ethernet-Header entsteht schließlich ein **Ethernet-Frame**



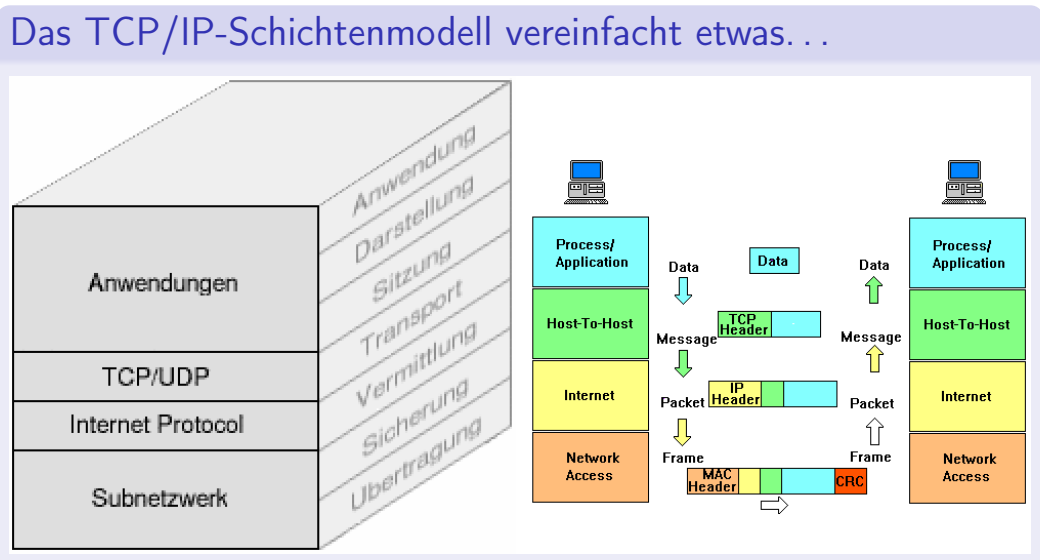
Schematisch Darstellung des Hinzufügens von Headern, zB. beim Sender. Der Empfänger der Daten muß die Header in der jeweiligen Schicht wieder entfernen, damit die Netzwerk-anwendung (z.B. Telnet oder FTP) mit den Daten arbeiten kann.

Abbildung 2.3: Erstellen und Auspacken eines Ethernet-Frames.

2.3 Das TCP/IP-Schichtenmodell

Die OSI-Normierung ist ein theoretisches Gebilde das in dieser Form in der Praxis nicht anzutreffen ist (und praktisch auch nie erprobt wurde). Der Grund dafür ist, daß TCP/IP längst etabliert war (und in der Praxis eingesetzt wurde) als die ISO mit OSI einen neuen Standard zu setzen versuchte. Bei der Entwicklung des OSI-Modells gingen die praktischen Erfahrungen des ARPANET's zwar mit in die Normierung ein, TCP/IP wurde von der ISO aber nie als ISO-Standard akzeptiert.

Es gibt keine generelle Regel dafür, wie die Protokollfamilie TCP/IP mit dem OSI-Referenz-Modell zu beschreiben ist. Eine sinnvolle Beschreibung ist erst mit dem Übergang zu einem 4-Schichten-Modell möglich. Hierbei werden die oberen drei Schichten des OSI-Modells, Kommunikations-, Darstellungs- und Verarbeitungsschicht, zu einer Schicht zusammengefaßt und als Anwendungsschicht (application layer) bezeichnet. Ebenso werden die unteren zwei Schichten, Bitübertragungs- und Sicherungsschicht, zu einer sogenannten Netzwerkschicht (network layer) zusammengefaßt.



1983 wurde TCP/IP das offizielle ARPANET-Protokoll. 1988 stellte OSI sein Schichtenmodell vor. Gegenüber dem offen entwickelten und jahrelang in der Praxis bewährten TCP/IP konnte es sich jedoch nie richtig durchsetzen.

Abbildung 2.4: Aufbau des TCP/IP-Schichtenmodell.

2.3.1 Aufgaben der einzelnen Schichten:

Anwendungsschicht: Sie umfaßt alle höherschichtigen Protokolle des TCP/IP-Modells. Hierzu zählen zunächst alle ursprünglichen ARPANET-Anwendungen wie z.B. telnet, ftp und SMTP. Im Laufe der Zeit kamen mit z.B. DNS oder HTTP auch neuere Protokolle dazu.

Transportschicht: Wie im OSI-Modell ermöglicht die Transportschicht im TCP/IP-Modell die Kommunikation zwischen den Quell- und Zielrechnern. Auf dieser Schicht werden zwei grundlegende Ende-zu-Ende-Protokolle definiert: TCP und UDP.

Vermittlungsschicht: Hier sind z.B. die Protokolle IP und ICMP definiert. ICMP dient der Übertragung von Diagnose- und Fehlerinformationen für IP.

Netzwerkschicht: Unterhalb der Vermittlungsschicht beschreibt das TCP/IP-Schichtenmodell nicht mehr exakt, was genau geschehen soll. Festgelegt ist lediglich, daß zur Übermittlung von IP-Paketen ein Rechner an ein Netz angeschlossen werden und ein Protokoll beherrschen muß. Das Protokoll ist nicht weiter spezifiziert. Das TCP/IP-Schichtenmodell macht hier vielmehr Gebrauch von bereits vorhandenen Protokollen, wie z.B. Ethernet oder PPP.

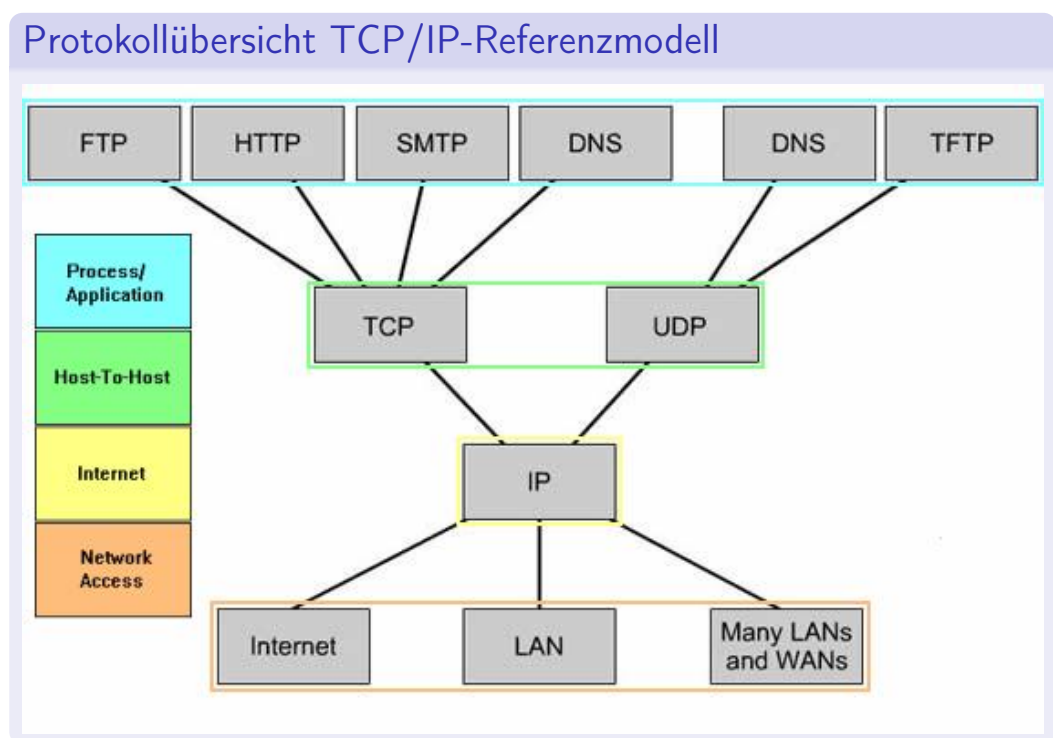


Abbildung 2.5: Das TCP/IP-Schichtenmodell beschränkt sich auf vier Ebenen.

2.4 Anwendungsmöglichkeiten der Schichtenmodelle

Mit der Kenntnis von OSI- und TCP/IP-Schichtenmodells und deren Funktionsweise ist z.B. eine systematische Fehlersuche bei Netzwerkproblemen möglich:

- Bitübertragungsschicht: Typische Probleme sind falsche oder defekte Kabel, Kontaktprobleme an Steckverbindungen oder Geschwindigkeitsproblem (10/100) im Ethernet.
- Sicherungsschicht: Doppelte MAC-Adressen (passiert eigentlich nie unabsichtlich) und inkompatible Frametypen bilden den Großteil der Probleme. Die Verbindung zum Network Layer stellt ARP her.
- Vermittlungsschicht: Doppelte oder inkorrekte IP-Adressen sowie falsche Netzwerkmasken sind hier sehr häufige Fehlerquellen. Auch Routing findet in diesem Layer statt und stellt eine Fehlerquelle dar.
- Transportschicht: Im Transport Layer sind UDP und TCP angesiedelt.
- Schichten 5-7: Die Anwendungsorientierten Schichten sind sehr komplex aufgebaut. Eine Fehlersuche muß auch ev. nicht richtig installierte oder konfigurierte Software mit einschließen.

Fehlerhäufigkeiten und -suche

Schicht	Fehlerhäufigkeit	Fehleranalyse
Anwendungsorientierte Schichten	ca. 15%	telnet
Transportschicht	ca. 10%	ethereal
Netzwerkschicht	ca. 15%	arp, ifconfig, ping, traceroute
Sicherungsschicht	ca. 20%	arp, ifconfig
Bitübertragungsschicht	ca. 40%	z.B. Kabeltester

Die Zahlenwerte sind als ungefähre Richtwerte aufzufassen die, je nach Netzwerk, mehr oder weniger stark variieren können.

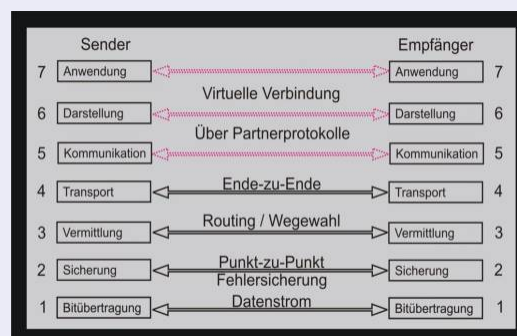


Abbildung 2.6: Strukturierte Fehleranalyse mit Hilfe des OSI-Schichtenmodells.

Kapitel 3

OSI-Schichtenmodell: Eine detaillierte Betrachtung

3.1 Layer 1: Bitübertragungsschicht

Die unterste Schicht im OSI-Schichtenmodell, auch *Bitübertragungsschicht* oder *physikalische Schicht* (*physical layer*) genannt, ist für alles zuständig, was mit Signalen und Übertragungsmedien zu tun hat. Sie spezifiziert z.B. elektrische und mechanische Eigenschaften der Übertragungsmedien. Unter elektrischen Eigenschaften versteht man z.B. Spannungspegel und deren definierte Änderung in definierten Zeitabständen. Analoges gilt natürlich, wenn die Datenübertragung mit Lichtimpulsen, z.B. in Lichtwellenleitern stattfinden soll. Nur werden hier eben Lichtintensitäten anstelle von Spannungspegeln beschrieben. Unter mechanischen Eigenschaften versteht man z.B. die Art und Weise wie die Steckverbindung zwischen Netzwerkkarte und Datenkabel beschaffen sein muß. Auch die Datenrate des Datenübertragungsmediums wird in dieser Schicht festgelegt. Weiterhin wird die maximale Übertragungslänge, bei der diese maximale Datenübertragungsrate funktionieren muß, festgelegt.

Diese Tabelle gibt einen Überblick über einige Mitglieder der Standardfamilie 802.3. Man sieht, daß sich zwar die Geschwindigkeit und ggf. die Topologie ändert, das Zugriffsverfahren aber immer das gleiche ist.

Beispiel für einige Layer 1 Eigenschaften von Ethernet

Ethernet Type	Max. Länge	Kabeltyp	Interface Typ	Standard	Phys. Topologie, Zugriff
10Base5	500 m	Dickes Coax	AUI	802.3	Bus CSMA/CD
10Base2	185 m	Dünnes Coax	BNC	802.3a	Bus CSMA/CD
10BaseT	100 m	UTP, CAT-3, Cat-4, CAT-5	RJ-45	802.3i	Stern CSMA/CD
100BaseTX	100 m	UTP CAT-5	RJ-45	802.3u	Stern CSMA/CD
1000BaseT	100 m	UTP CAT-5e und CAT-6	RJ-45	802.3ab	Stern CSMA/CD

Abbildung 3.1: Einige Eigenschaften, die auf Layer 1 definiert sind.

3.1.1 Übertragungsmedien

3.1.1.1 Kupfer

3.1.1.1.1 Kabel-Klassen

Datenkabel werden anhand ihrer Übertragungseigenschaften (z.B. Dämpfung, Übersprechen) und ihrer elektromagnetischen Eigenschaften in verschiedene Klassen eingeteilt. Wichtige Klassen für moderne Ethernet-Technologie sind:

Cat 5e: Zur Zeit Standard bei der horizontaler 100Mbps-Vernetzung (100Base-TX, 1995, IEEE 802.3u). Im Prinzip gigabitfähig (1000Base-T, 1999, IEEE 802.3ab)!

Cat 6: Ist für Ethernet 1000Base-T (Gigabit) bis zu einer Länge von 100 m verwendbar. Cat 6 Kabel sind für Betriebsfrequenzen von 250 MHz pro Adernpaar ausgelegt. Im wesentlichen handelt es sich um eine verschärfte Cat 5e Normierung.

Cat 7: Vier einzeln abgeschirmte Adernpaare innerhalb eines gesamten Schirms. Sie können mit Frequenzen bis 600 MHz betrieben werden und sind 10GB-fähig. 10GBASE-T (IEEE P802.3an) wurde Mitte 2006 von der IEEE 2006 verabschiedet.

Aufbau verschiedener TP-Kabel

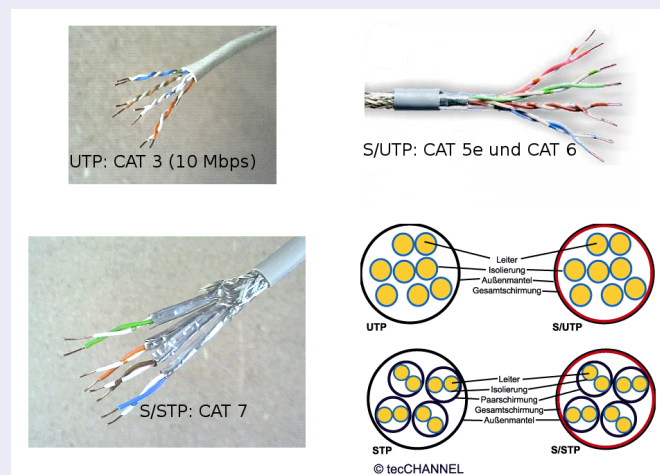


Abbildung 3.2: Verschiedene Kabelklassen.

3.1.1.1.2 Twisted Pair: 10Base-T und 100Base-TX (Fast-Ethernet)

Ein Nachteil von 10Base5 und 10Base2 (und auch den frühen Mischnetzwerken aus 10/100 Mbps zu Zeiten von Hubs bzw. Dualspeed-Hubs) war jedoch, daß sich alle Benutzer, die an einem Kabel angeschlossen waren, die zur Verfügung stehende Bandbreite teilen mußten. Mit einem 10Base-T Kabel konnte dieses Problem gelöst werden. Mit seinen beiden Adernpaaren war es in der Lage, Workstations, die bis zu 100 m von einem Anschlußpunkt entfernt waren, mit einer reservierten 10 Mbit-Leitung zu bedienen. Später wurde bei gleichbleibender Entfernung die Geschwindigkeit auf 100 Mbit erhöht (durch eine Änderung des Kodierungsverfahrens der Daten: Manchester → NRZ (non return zero)). Diese Technik (100Base-TX) erhielt den Namen Fast Ethernet und ist immer noch Standard in vielen Netzwerken. Kupfergebundene Datenübertragung hat heute schon die 1 Gbps-Marke übersprungen. Zur Zeit wird an einem Standard für 10 Gbps gearbeitet. Diese hohen Geschwindigkeiten erzielt man durch die Benutzung der beiden noch freien Adernpaare eines TP-Kabels. Auf jedem Adernpaar werden die Daten mit 250 Mbps übertragen und ergeben so in der Summe 1 Gbps. Auch hier war es nötig das Kodierungsverfahren erneut anzupassen.

Schematisierte Signalverläufe

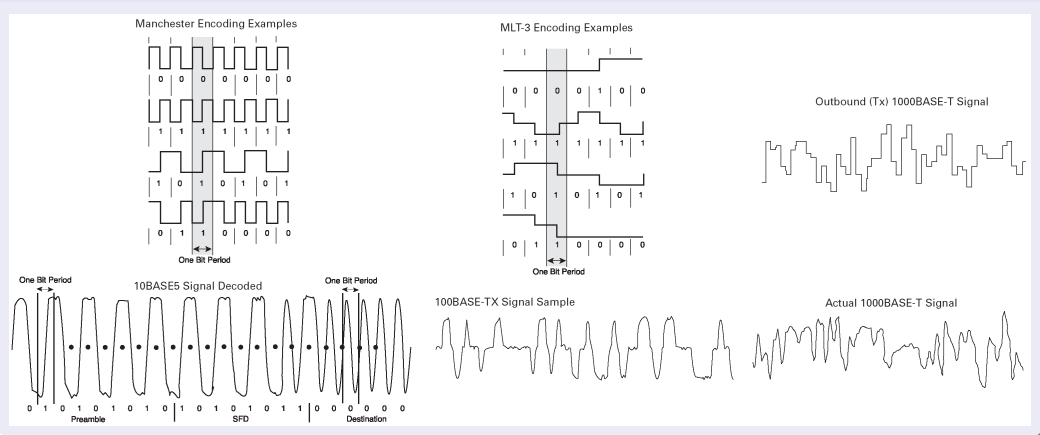


Abbildung 3.3: Signaldarstellungen von Twisted-Pair Netzwerken.

3.1.1.1.3 Dämpfung

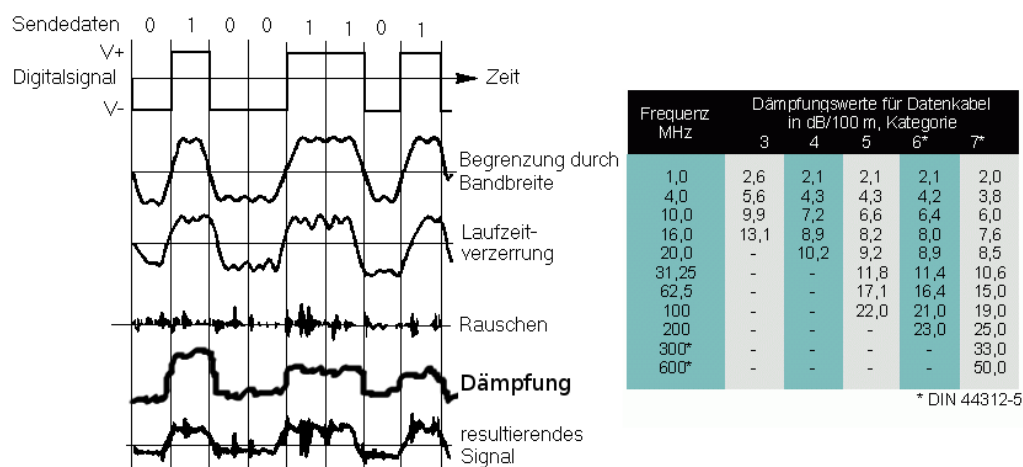
Als Dämpfung bezeichnet man die Abnahme der Signalstärke also eine Verringerung der Amplitude. Gedämpft werden sowohl optische als auch elektrische Signale. Durch Dämpfung wird die Reichweite begrenzt. Verstärker können die Reichweite erhöhen. Die Dämpfung $\mathcal{D}$ wird in Dezibel (dB) angegeben. Sie berechnet sich aus der Eingangsspannung U_1 und der Ausgangsspannung U_2 :

$$\mathcal{D} = 20 \cdot \lg \frac{U_1}{U_2} \text{ dB} \quad (3.1)$$

Man kann die Dämpfung auch mit Hilfe der Leistung berechnen, da $P \sim U^2$. Sie berechnet sich aus der Eingangsleistung P_1 und der Ausgangsleistung P_2 :

$$\mathcal{D} = 10 \cdot \lg \frac{P_1}{P_2} \text{ dB} \quad (3.2)$$

Beispiel: Wenn die Spannung auf den Wert von 50% absinkt, also die halbe Spannung verloren geht, dann hat das Kabel eine Dämpfung von 6,02 dB. Wenn die Spannung auf den Wert 70,7%, also $\frac{100}{\sqrt{2}}$ absinkt beträgt die Dämpfung nur 3,01 dB (was natürlich gegenüber dem ersten Beispiel ein besserer Wert ist).



Die Dämpfung ist nur eine von vielen Möglichkeiten, die ein Signal längs des Weges störend beeinflusst. Alle Störmöglichkeiten in Summe sind im letzten Beispiel zu sehen.

Abbildung 3.4: Signalbeeinflussung durch verschiedene Faktoren.

3.1.1.2 Optische Medien

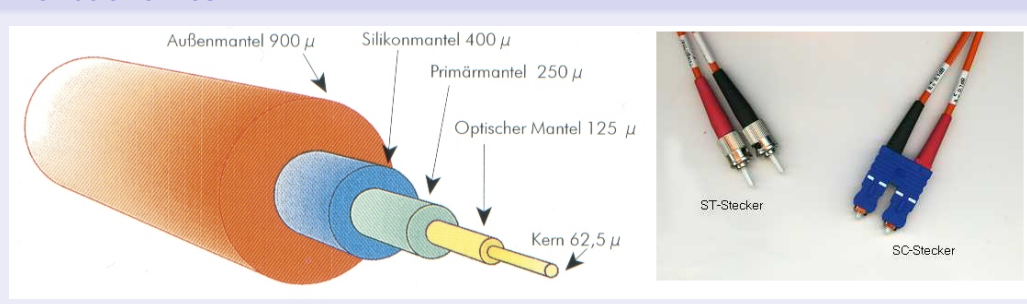
3.1.1.2.1 Lichtwellenleiter

Glasfaserkabel, oft auch Lichtwellenleiter (LWL) genannt, sind flexible, hochtransparente optische Medien aus Mineralglas (meist Kieselglas, SiO_2) oder organischem Glas (Kunststoff), in denen Licht kontrolliert geleitet werden kann.

Die Faser besteht aus einem Kern (Core), einem Mantel (Cladding) und einer Beschichtung (Coating oder Buffer). Der lichtführende Kern dient zum Übertragen des Signals. Der Mantel ist auch lichtführend, hat jedoch eine niedrigere optische Brechzahl als der Core. Durch diese Konstruktion erreicht man eine Totalreflexion an der Grenzschicht und somit eine Führung der Strahlung im LWL-Kern. Die äußere Beschichtung ist ein Schutz vor mechanischen Beschädigungen (z.B. Bruchsicherheit) bzw. Belastungen (z.B. Zugfestigkeit) und Feuchtigkeit.

Die Signalübertragung erfolgt unidirektional über Lichtimpulse, d.h. nur in eine Richtung, weshalb mindestens zwei Fasern (Stränge) pro Kabel notwendig sind. Die Lichtimpulse werden mittels einer Laser-Diode (1300 nm) oder einer Lumineszenz-Diode (LED = light emitting diode bei einer Wellenlänge von 850/860 nm oder 1300 nm) in den Kern übertragen.

Aufbau eines LWL



Zusammenfassung

Vorteile	Nachteile
- hohe Übertragungsraten (GB-, TB-Bereich) - sehr große Reichweiten - keine Beeinflussung durch äußere Störfelder - keine Erdung nötig - Verlegung in explosionsgefährdetem Umfeld möglich (keine Funkenbildung) - kein Übersprechen	- hoher Konfektionierungsaufwand - teure Gerätetechnik - Schwachstelle Stecker-technologie - aufwendige Messtechnik - empfindlich gegenüber mechanischer Belastung

Abbildung 3.5: Glasfaserkabel.

3.1.1.2.2 Multi- und Monomode

Multimode: Hier sind viele einzelne Lichtwellen (Moden) an der Signalübertragung beteiligt. Aufgrund mehrerer möglicher Lichtwege kommt es zu Laufzeitunterschieden, daher sind Multimode-Fasern für große Distanzen und bei hoher Bandbreite ungeeignet.

Monomode oder Singlemode: Die Signalübertragung erfolgt hier über eine einzige Lichtwelle (Mode). Bei großen Distanzen und/oder Bandbreiten werden nur Singlemode-Fasern eingesetzt, da hierbei Signalverzerrungen nur aufgrund der Dispersion (Streuung an Verunreinigungen) auftreten.

Merkmal	Multimode	Mono-, Singlemode
innerer Core-Durchmesser	62.5 (US) bzw. 50 μm (EU)	5 - 9 μm
äußerer Faser-Durchmesser	ca. 125 μm	ca. 125 μm
Kosten	günstig	teuer
Ummantelung	meist orange	meist grau oder gelb
Lichtquelle	LED günstig, wenig Leistung	Laserdiode teuer, hohe Leistung
Wellenlänge	ca. 850 nm bzw. ca. 1300 nm	ca. 1310 bzw. 1510 nm

Hochgeschwindigkeits-Glasfaserverbindungen

Singlemode-Fasern für Weitverkehrsnetze werden im DWDM-Verfahren betrieben. Dabei wird über mehrere Laser Licht verschiedener Wellenlängen in eine Faser eingekoppelt: Bandbreite-Längen-Produkte bis $> 10 \text{ (Tbit/s)*km!}$

Bandbreiten-Längen-Produkt

Bei Glasfaserkabeln ist die maximal nutzbare Bandbreite umgekehrt proportional zur Entfernung, die ohne Signalregeneration überbrückt werden kann. Das Produkt aus Bandbreite und Länge bildet daher eine Maßzahl, die hauptsächlich von der Art des verwendeten LWL abhängig ist.

- bei 500 m Länge mit einer Bandbreite von 2 GHz gearbeitet werden
- bei 1 km Länge mit einer Bandbreite von 1 GHz gearbeitet werden usw. . .

Typische Werte sind:

- 100 MHz*km für Stufenindex-Multimode-Fasern,
- 1 GHz*km für Gradientenindex-Multimode-Fasern und
- 10 GHz*km für Stufenindex-Monomode-Fasern

Abbildung 3.6: Signalverläufe in verschiedenen Typen von Glasfaserkabeln.

3.1.1.3 Trägerwelle

3.1.1.3.1 WLAN

Beschreibung: Ein drahtloses lokales Funknetzwerk, mit dem meist ein Standard der IEEE 802.11-Familie gemeint ist. Üblicherweise koordiniert eine spezielle Basisstation (Wireless Access Point), die einzelnen Wireless Clients. Häufig ist diese Basisstation dann auch Router in ein weiteres Netz (z.B. ein weiteres Funknetz oder ein klassisches Kabelnetz).

Datenraten: Hierbei ist zu berücksichtigen, daß sich alle Geräte im Netz die Bandbreite teilen. Geräte werden oft nur mit der unter optimalen Bedingungen erzielten Bruttorate (alle Bytes/s) beworben. Die erreichbare Netto-Datenrate (Nutzbytes/s) ist aber nur meist halb so hoch.

Diskussion um gesundheitliche Risiken: Die von WLAN-Geräten benutzten Funkfrequenzen liegen im Mikrowellenbereich (2.4 bzw. 5 GHz). Nach einer Studie des Bundesamts für Strahlungsschutz gibt es (innerhalb der gesetzlichen Grenzwerte) nach dem aktuellen Stand der Wissenschaft keine Nachweise, daß hochfrequente elektromagnetische Felder gesundheitliche Risiken verursachen.

Frequenzen, Datenraten und Normen

Norm	Frequenz	Max. Mbps
802.11a	5 GHz	54
802.11b	2.4 GHz	11
802.11g	2.4 GHz	54

802.11a hat sich wg. höherer Hardwarekosten und stärkerer Störung durch Satelliten bzw. Radargeräten noch nicht etabliert.

Leichtbauwände mindern die Reichweite, sind aber einzeln kein Hindernis; dagegen werden Stahl und Beton nicht durchdrungen. Bäume, insbesondere dicht belaubte, sind ebenfalls Hindernisse für WLAN-Verbindungen.

Reichweite

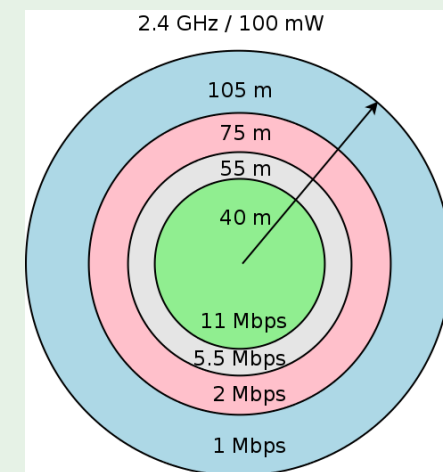


Abbildung 3.7: Reichweiten und Frequenzen moderner WLAN-Techniken.

3.1.1.4 Weitere Übertragungsmöglichkeiten

Optischer Richtfunk Bei Optischem Richtfunk (engl.: FSO), handelt es sich um eine Technologie zur Übertragung von digitalen Daten jeglicher Art (also z.B. auch Video oder Sprache)mit Lichtimpulsen. Es können Datenraten bis zu 2.5 Gbps realisiert werden.

Homeplug Diese Technik wird manchmal auch *Inhouse-Powerline* genannt und wird im Haus über 230V Stromleitungen betrieben. mit dem Computer verbunden. Der Einsatzbereich der Homeplug-Technologie liegt bei kleinen LANs v.a. im privaten Bereich oder wenn keine Netzkabel verlegt werden können bzw. dürfen.

10Base5 (Thicknet) bzw. 10Base2 (Thinnet) Hierbei handelt es sich um die älteren Ethernet-Implementierungen. 10Base5 wurde ab 1980 eingesetzt. Ab 1985 wurde es allmählich durch 10Base2 ersetzt. Als Übertragungsmedium wurden Kupfer-Koaxialkabel eingesetzt. Heute nur noch in sehr alten Installationen anzutreffen!

Kabelmodem Hierbei handelt es sich um ein Gerät, mit dem man über das TV/Radionetz einen Breitband-Internetzugang realisieren kann.

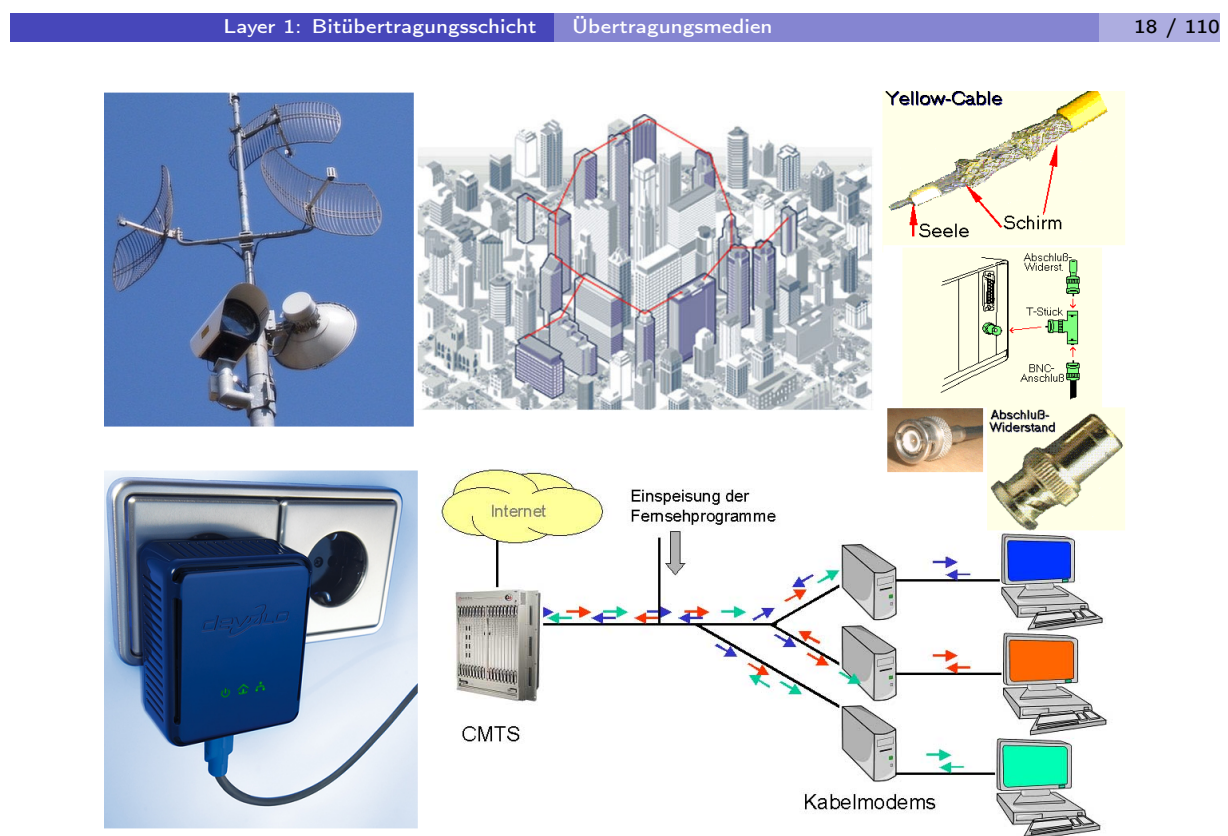


Abbildung 3.8: Technische Grundlagen von 10Base2/5 Netzwerken.

3.1.2 Netzwerkgeräte

3.1.2.1 Repeater und Hubs

Ein Repeater ist eine Art Verstärker: Er wandelt Binärdaten, in z.B. Spannungs- bzw. Lichtimpulse um. Erst jetzt sind Bits "transportfähig". Außerdem kann er, in gewissen Grenzen) schwache Signale wieder auffrischen (z.B. das Timing korrigieren). Repeater können auch zur Verbindung zweier Netzwerksegmente (heute v.a. in Form von Medienkonvertern) verwendet werden.

Repeater werden deshalb dem Layer 1 des OSI-Schichtenmodells zugeordnet, weil sie nur auf Bit-Ebene arbeiten, die Bedeutung der Signale aber nicht interpretieren können.

Hubs regenerieren Netzwerksignale wie Repeater, besitzen aber im Unterschied zum klassischen Repeater, mehr als nur zwei Ports. Hubs senden Pakete, die sie auf einem Port empfangen, an alle anderen Ports weiter. Da die Eigenschaften von Hubs und Repeatern so ähnlich sind, werden Hubs oft auch als **Multiport-Repeater** bezeichnet. In modernen Netzwerken sind Hubs heute nur noch selten zu finden.

Mit einem Hub ist es sehr leicht möglich den Netzwerkverkehr in einem geschalteten Netzwerk mitzuschneiden (was aber u.U. von den Betroffenen bemerkt werden kann - deshalb besser einen Monitorport verwenden, falls vorhanden).

Hubs und ihre Möglichkeiten



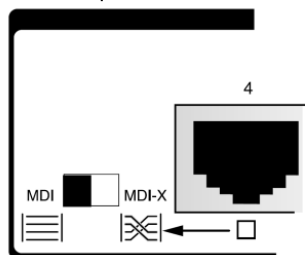
Managebarer HUB AT-3612TR
bzw. AT-3624TR



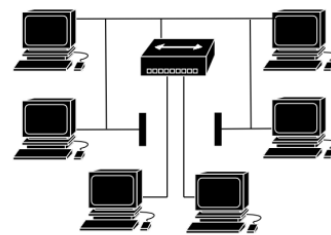
Management Interface,
hier 9 pol. RS-232



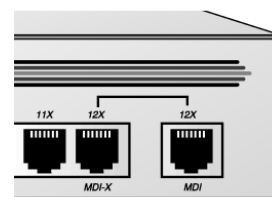
Nicht managebarer Hub AT-420TR
bzw. AT-820TR



Wahlschalter MDI (Für den Anschluß
von z.B. PC's und Routern) und
MDI-X (für den Anschluß eines we-
teren Hubs)



Typisches Netzwerk
mit einem Hub



Stromlos: Verbindet zwei PC's mit
zwei Straight Through-Kabeln
Mit Strom: Uplink und PC-Anschluß
mit zwei Straight Through-Kabeln
bzw. Anschluß von zwei PC's wenn
der rechte mit einem Crossover-
Kabel angeschlossen wird!

Abbildung 3.9: Managebare und nicht-managbare Hubs.

3.1.2.2 Transceiver

Transceiver ist ein Kunstwort, gebildet aus den Begriffen Transmitter (Sender) und Receiver (Empfänger). Er stellt die Schnittstelle zwischen einem Datenendgerät (Computer) und dem Übertragungsmedium (Koaxkabel, Twisted Pair, Lichtwellenleiter usw.) zur Verfügung. Der Transceiver hat mehrere Aufgaben. Er wandelt digitale in analoge Signale um, übernimmt Steuerungsaufgaben beim Senden bzw. Empfangen von Daten und sorgt für die Übergänge zwischen unterschiedlichen Übertragungsmedien, wie z.B. von Lichtwellenleiter auf Kupferkabel. In vielen Fällen ist der Transceiver bereits in das Endgerät (z.B. eine Netzwerkkarte) integriert.

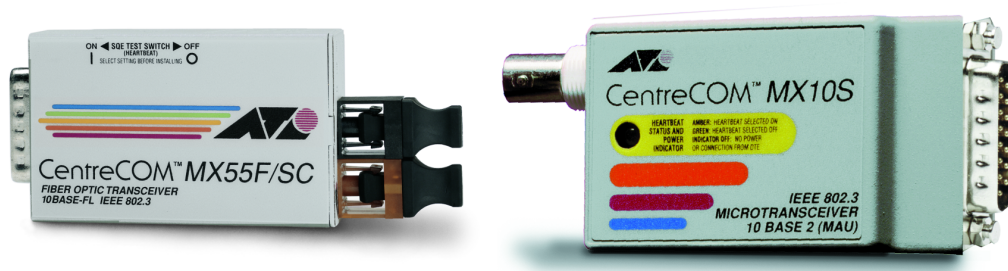


Figure 18-28: Cisco 2503 Rear View

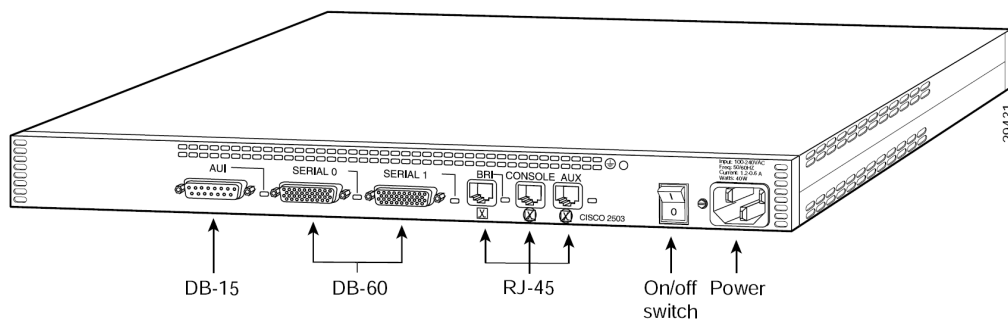


Abbildung 3.10: Transceiver werden mit den 15-poligen Sub-D-Steckern von Netzwerkgeräten verbunden.

3.1.3 Physische NetzwerkTopologie

3.1.3.1 Bus-Topologie

Die einfachste aller Topologien. Es handelt sich dabei um eine passive Technologie, weil für die Datenübertragung zwischen mehr als zwei Endgeräten kein Verstärker oder sonstiges Verbindungselement benötigt wird. Das vom Endgerät auf das Kabel gegebene Signal breitet sich in beiden Richtungen aus und wird längs des Wegs gedämpft. Dies ist u.a. ein Grund für die Längenbegrenzung der Ethernet-Technologie. Es existiert ein Kabel (der Bus), an das alle Netzwerkgeräte angeschlossen sind und das an beiden Enden terminiert werden muß. Sonst würde das elektrische Signal in die Leitung zurückreflektieren und dadurch die Datenübertragung verhindern. Diese Topologie liegt allen Ethernet-Netzwerk zugrunde und stammt aus den Gründungsjahren der LAN-Netzwerke. 10Base-2 ist z.B. eine Technologie, die auf genau dieser Topologie beruht. Heute wird sie in dieser Form (Koax-Kabel) eigentlich nicht mehr verwendet.

Die Bustopologie ist, wegen der geringen Kabelmengen, preislich relativ günstig. Fällt eine Station aus, wird das übrige Netz nicht gestört.

Nachteile

- Alle Daten werden in nur einem Kabel übertragen
- Ist das Übertragungsmedium gestört, geht nichts mehr
- Fehlersuche aufwendig
- Jede Station sieht alle Pakete

Bustopologie

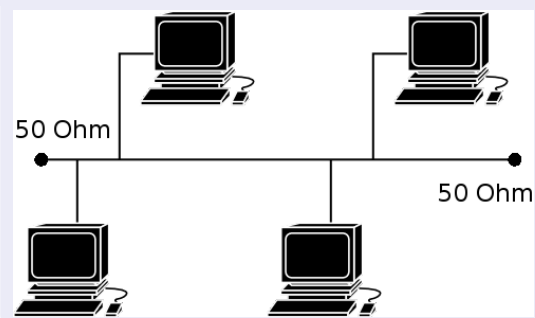


Abbildung 3.11: Das klassische Ethernet: Ein Koax-Bus an dem alle Hosts hängen!.

3.1.3.2 Stern-Topologien

Stern: Die star- oder Stern-Topologie ist die heutzutage am meisten verwendete Form um Ethernet-Lan's physisch zu realisieren. Hier wird jeder Host mit einem eigenen Kabel an einen zentralen Punkt (z.B. Hub oder Switch) angeschlossen. Es existiert somit eine Punkt-zu-Punkt-Verbindung zwischen Endgerät und Datenverteiler. Funktioniert also ein Kabel (oder ein Netzwerkgerät) nicht, ist nur dieser eine Kabelstrang nicht funktionsfähig. Der (managbare) Switch bzw. Hub hat den Vorteil, daß er zentrale Sicherheitsaspekte zur Verfügung stellen kann und den Nachteil eines "central point of failure".

Erweiterter Stern: Diese Form entsteht immer dann, wenn eine Sterntopologie erweitert wird (z.B. um Switches oder Hubs). Die neu entstandene Topologie bezeichnet man dann als "erweiterter Stern".

Die Sterntopologie verschlingt große Kabelmengen. Fällt der Verteiler aus, steht das Netz. Die Vorteile wiegen die Nachteile im praktischen Einsatz aber auf!

Vorteile

- Der Ausfall einer Station oder der Defekt eines Kabels hat keine Auswirkung auf das restliche Netz
- Volle Bandbreite für alle
- Weitere Stationen und/oder Verteiler können problemlos hinzugefügt werden (extended star)

Sterntopologie

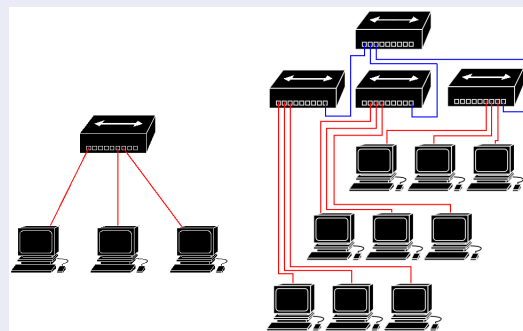


Abbildung 3.12: Modernes Ethernet: Trotz TP-Kabeln immer noch ein Bus!

3.1.3.3 Ring

Die logische Ring-Topologie ist (bzw. war) eine weitere wichtige LAN-Technologie. Alle Hosts sind hier, wie der Name andeutet, durch eine kreisförmige Leitung miteinander verbunden. Der Mechanismus der Datenübertragung ist von der Bustopologie völlig verschieden und funktioniert vereinfacht beschrieben wie folgt: Ein Token kreist im Ring und hält an jeder Station an. Will eine Station senden, dann wird es mit Nutzlast und der Zieladresse beladen. Danach wird es weitergereicht bis es die Zielstation erreicht hat, die es dann von der Nutzlast befreit. Der Vorteil dieser Technologie ist, daß es zu keinen Kollisionen kommen kann, weil ja nur ein Frame unterwegs ist. Bei der Single Ring Technologie teilen sich alle Netzwerkgeräte dasselbe Kabel, der Datentransport findet immer in die gleiche Richtung statt. Jedes Gerät muß eine definierte Zeitspanne warten, bis es an der Reihe ist. Die meisten Single Ring Topologien wurden als physische Sterntopologie realisiert. Token Ring ist ein Beispiel dafür.

Die Dual Ring Topologie kann Daten in beiden Richtungen senden. Die Aussfallsicherheit wird dadurch wesentlich erhöht. Wenn beide Ringe ausfallen entstehen durch eine Art Selbstheilungsprozess zwei unabhängige Ringe.

Die Ring-Topologie ist der Bus-Topologie, vom technischen Ansatz her, überlegen. Vor allem die hohen Kosten im Vergleich zu Ethernetprodukten ließen sie aber scheitern.

Vorteil

- deterministisch

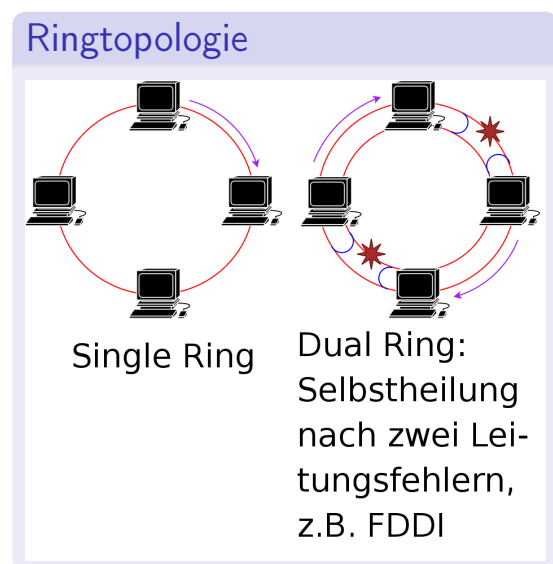


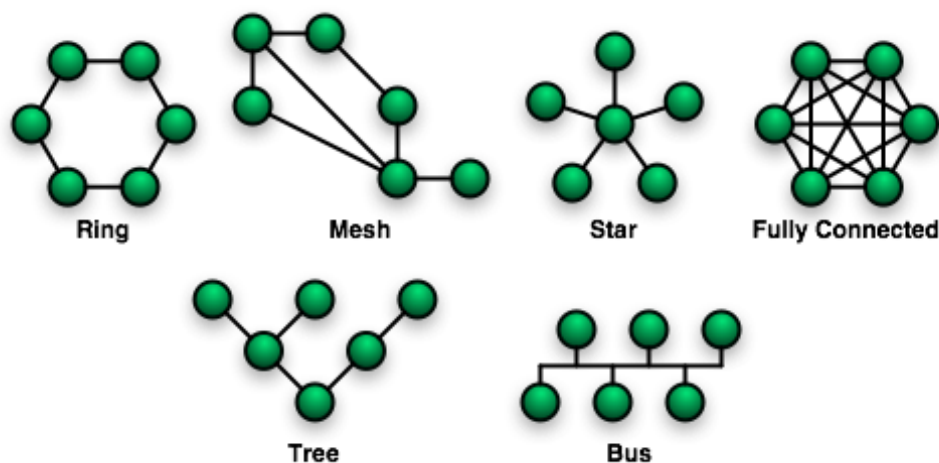
Abbildung 3.13: Ring: Hat sich trotz überlegener Technologie nie richtig etabliert.

3.1.3.4 Vermascht

Vermaschte Netzwerke sind beim Ausfall eines Endgerätes oder einer Leitung im Regelfall in der Lage, durch Umleiten (Routing, Spanning Tree) der Daten weiter zu funktionieren.

In großen Netzen findet man oftmals Strukturen, die sich aus mehreren verschiedenen Topologien zusammensetzen. So ist das Internet in weiten Teilen ein vermaschtes Netz, trotzdem gibt es "Hauptverkehrsadern" (sogn. Backbone-Leitungen), die bevorzugt benutzt werden.

- Vorteile
 - Sicherste Variante eines Netzwerkes
 - Bei Ausfall eines Endgerätes ist durch Umleitung die Datenkommunikation weiterhin möglich (Redundanz)
- Nachteile
 - Hoher Kabelverbrauch. Installation auch bei nicht vollständig vermaschten Netzwerken sehr aufwendig.
 - Sehr hoher Energieverbrauch durch die Switches und Router



In einem vermaschten Netz ist jedes Endgerät mit einem oder mehreren anderen Endgeräten verbunden. Wenn jeder Teilnehmer mit jedem anderen Teilnehmer verbunden ist, spricht man von einem vollständig vermaschten Netz.

Abbildung 3.14: Es sind viele physischen Topologien vorstellbar und realisiert.

3.1.4 Wichtige Begriffe

3.1.4.1 Autonegotiation

Die auf Layer 1 im OSI-Schichtenmodell angesiedelte Autonegotiation (auch als NWAY bezeichnet) soll es zwei Linkpartner im Ethernet ermöglichen, sich optimal zu konfigurieren. Die Stationen tauschen dazu über den sog. Fast Link Pulse (FLP) Informationen aus, mit denen sie dem jeweils anderen Linkpartner ihre unterstützten Modi mitteilen. Auto-Negotiation handelt die folgenden Parameter aus:

- Übertragungsgeschwindigkeit (10/100)
- Duplexeinstellungen (Voll-, Halbduplex)

Auto-Negotiation führt keine Überprüfung des Links durch. Unterstützen beide Karten Fast Ethernet, werden auch über ein Cat3-Kabel 100 MBit/s gefahren. In der Praxis gibt es mitunter Probleme, Komponenten unterschiedlicher Hersteller, mit Auto-Negotiation zueinander zu bringen.

Sollte eine Seite keine Auto-Negotiation unterstützen, dann erkennt die Auto-Negotiation-fähige Station am Link Pulse des Gegenübers die Geschwindigkeit und stellt sich korrekt darauf ein. Der Übertragungsmodus wird dann in so einem Fall **immer auf Halbduplex** gestellt.

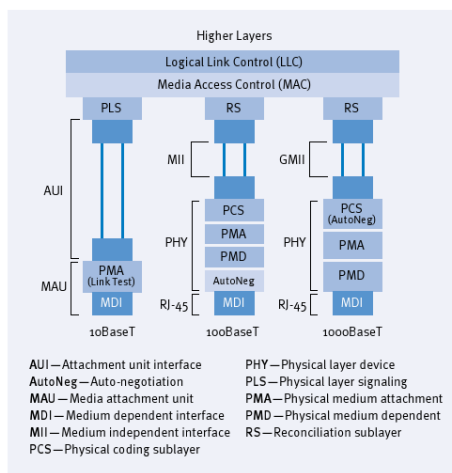


Figure 1. Data terminal equipment layer model (Redrawn from the IEEE Std 802.3, 1998 Edition)

Priority	Technology
1 (highest)	1000BaseT — Full duplex
2	1000BaseT — Half duplex
4	100BaseTX — Full duplex
7	100BaseTX — Half duplex
8	10BaseT — Full duplex
9 (lowest)	10BaseT — Half duplex

Figure 6. Priority resolution table

10/100/1000 BaseT switch configuration		Auto	10	10	100	100	1000	1000
Speed	Duplex		Half	Full	Half	Full	Half	Full
NIC configuration								
Speed	Duplex							
Auto		1000 FDX	10 HDX	Duplex conflict	100 HDX	Duplex conflict	1000 HDX	Duplex conflict
10	Half	10 HDX	10 HDX					
10	Full	Duplex conflict		10 FDX				
100	Half	100 HDX			100 HDX			
100	Full	Duplex conflict				100 FDX		
1000	Half	1000 HDX					1000 HDX	
1000	Full	Duplex conflict						1000 FDX

Figure 12. Configuration table for 10/100/1000 BaseT devices

» Applications will timeout and retransmit continuously, causing a very slow connection.

Figure 12 summarizes all possible combinations of speed and duplex settings, both on 10/100/1000-capable switch ports and on NICs. For example, connecting a 10/100/1000-capable NIC configured for auto-negotiation with a 10/100/1000 switch also configured for auto-negotiation results in the ports for both the NIC and the switch being configured at 1000 Mbps full duplex. Figure 12 also shows combinations that would yield no link or link fail conditions, as well as combinations that would yield a duplex mismatch.

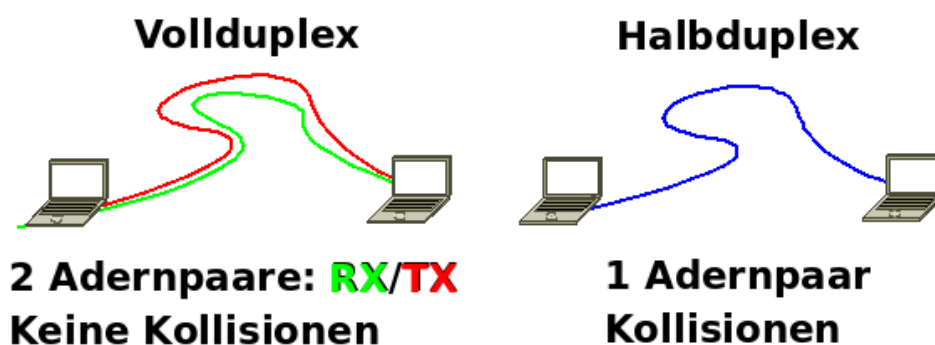
Abbildung 3.15: Technische Grundlagen der Autonegotiation

3.1.4.2 Duplex

Mit Voll- bzw. Halbduplex bezeichnet man in der Kommunikationstechnologie die Richtungsabhängigkeit von Kommunikationskanälen.

Halbduplex bedeutet, daß Daten nicht gleichzeitig gesendet und empfangen werden können. Wenn nun zwei Stationen gleichzeitig senden treffen die beiden Datenpakete irgendwann aufeinander. In diesem Fall "vernichten" sich die Datenpakete gegenseitig. Man spricht dann von einer "Kollision". Erkennt eine Station im Netzwerk eine solche Kollision dann erzeugt sie ein Jam-Signal. Durch das Jam-Signal werden die jeweiligen Sendestationen veranlaßt, ihre Datenpakete nochmals zu einem zufälligen Zeitpunkt zu senden. Hubs arbeiten immer mit Halbduplex.

Vollduplex läßt den Datenaustausch in beide Richtungen gleichzeitig zu. Dadurch werden Kollisionen verhindert und die Netzwerkleistung "verdoppelt". In einem 10 Mbps-Netz wird somit eine Steigerung auf 20 Mbps erreicht, in einem 100 Mbps-Netz auf 200 Mbps. Gigabit-Ethernet (1000Base-T) ist z.B. eine reine Vollduplex-Technik. Fast Ethernet kann sowohl in Voll-, als auch in Halbduplex realisiert sein.



Gerät A	Gerät B	Resultat
100 half	100 half	Halfduplex, bei Last Kollisionen
100 half	100 full	Duplex-Mismatch , viele Fehler, keine Performance
100 half	auto	Glück gehabt, B geht auf 100 half
100 full	100 full	immer noch die sicherste Methode
100 full	auto	Duplex-Mismatch, B geht auf 100 half, viele Fehler, keine Performance
auto	auto	mit etwas Glück gehen A und B auf 100 full
100	10	Speed-Mismatch, kein Link

Kollisionsfragmente oder runt frames sind kürzer als 64 Byte. Solche Frames werden von den Switches automatisch aussortiert!

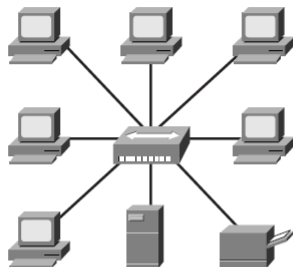
Abbildung 3.16: Der Unterschied zwischen Voll- und Halbduplex

3.1.4.3 Kollisionsdomäne

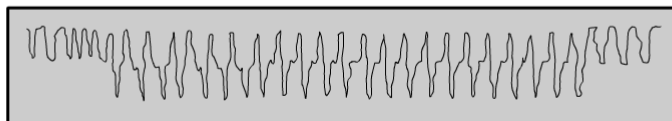
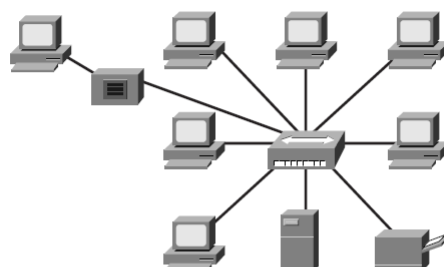
Ethernet ist eine Übertragungsmethode, bei der sich alle Teilnehmer ein Übertragungsmedium teilen. Je mehr Teilnehmer vorhanden sind, um so wahrscheinlicher kommt es zu Kollisionen auf dem Übertragungsmedium. Kollisionen treten immer dann auf, wenn zwei Teilnehmer, praktisch gleichzeitig, ihre Daten auf das Medium legen. Kollisionen zerstören die übertragenen Daten und bewirken (im besten Fall), daß sie erneut gesendet werden. Dazu warten die beteiligten Stationen eine zufällig gewählte Zeitspanne ab und senden ihre Daten dann erneut. Besonders anfällig für Kollisionen sind nicht geschaltete Netzwerke. Sind mehr als ca. 30 Stationen an einem Ethernet-Bus angeschlossen, ist wegen der andauernden hohen Anzahl an Kollisionen praktisch kein Netzwerkverkehr mehr möglich.

Kollisionsdomänen

Collision Domain Created by Hub



Collision Domain Extended by Repeater



Kollision im 10Base2 Netzwerk: Links sieht man das normale Signal (Manchester kodiert). Im Mittelteil die Kollision (erkennbar am negativen Ausschlag), rechts wieder das normale Signal. Die Kollision wird von einem Schaltkreis auf der Netzwerkkarte erkannt und sofort allen Stationen mitgeteilt: Alle hören auf zu senden!

Abbildung 3.17: Kollisionen lassen sich in 10Base2/5 Netzwerken nicht vermeiden.

3.2 Layer 2: Sicherungsschicht

Was passiert nun mit dem Strom von Einsen und Nullen aus der Bitübertragungsschicht? Der Data Link Layer

- arbeitet mit **Frames**, d.h. er faßt anhand bestimmter Kriterien (Preamble und Frame Checksum, FCS) eine Serie von Bits als zusammenhängenden Datensatz auf
- koordiniert mit Hilfe von **CSMA/CD** die Übertragung der Daten längs des physikalischen Verbindungsweges

Zur Adressierung der Frames verwendet er die sogn. MAC-Adresse, eine 48 Bit breite Kennzeichnung, die innerhalb eines Netzwerks eindeutig sein muß! DIX (Ethernet II, 1982) und IEEE (Ethernet 802.2/3, 1983) unterscheiden sich wie folgt:

DIX: Ethernet II	IEEE: Ethernet 802.2/3
keine Sublayer	LLC und MAC
type field	length field

Layer 2 Protokolle

Heutzutage ist **Ethernet II** mit dem IEEE Standard **Ethernet 802.3** verschmolzen. Die Bytes 13 und 14 haben (aus Kompatibilitätsgründen zu **Ethernet II**) bei **Ethernet 802.3** eine Doppelfunktion die vom Host, der den Frame empfängt, interpretiert werden muß: Ist ihr Zahlenwert ≤ 1536 bzw. $0x0600$ dann wird der gesamte Frame als **Ethernet 802.3** interpretiert (mit dem Zahlenwert als Framelänge). Zahlenwerte > 1536 bedeuten, daß es sich um einen **Ethernet II**-Frame handelt (mit dem Zahlenwert als Ethertype, also z.B. $0x0800$ für das IP-Protokoll IPv4 oder $0x0806$ für ARP). **Ethernet II** kennt kein Längenfeld!

Preamble	Dest.	Src.	Type	Data/Pad	FCS
8	6	6	2	46 to 1500	4

Im Vergleich dazu der Aufbau des IEEE 802.3 Ethernet Frame:

Preamble	SFD	Dest.	Src.	Type/Length	Data/Pad	FCS
7	1	6	6	2	46 to 1500	4

Ethernet II wurde ursprünglich von DIX (Digital, Intel und Xerox) entwickelt und ist das am häufigsten anzutreffende Frame Format in TCP/IP basierten Ethernets. Weitere Layer 2 Protokolle sind Token Ring, ISDN und PPP.

Abbildung 3.18: Frames vom Typ Ethernet II sind am häufigsten anzutreffen.

3.2.1 Wichtige Begriffe

3.2.1.1 LAN-Schichtenmodell

Der Layer 2 wird vom OSI-Schichtenmodells in zwei Hälften aufgeteilt:

- Teilschicht für Mediumzugangskontroll: Media Access Control (MAC)
- Teilschicht für die logische Verbindungskontrolle: Logical Link Control (LLC)

MAC: MAC ist die untere Teilschicht der Sicherungsschicht im LAN-Schichtenmodell. Das Mediumzugangsverfahren umfasst z.B. die Frame-Aufbereitung beim Senden (Encapsulation) und Empfangen (Decapsulation). Außerdem kann auf der MAC-Schicht u.a. eine Trägererkennung (Carrier Sense/No Carrier Sense, s.a. CSMA/CD, IEEE 802.3) erfolgen.

LLC: LLC ist ein OSI-Protokoll, das von der IEEE-Arbeitsgruppe 802¹ entwickelt wurde und für alle LAN-Subsysteme im Rahmen des Standards IEEE 802 gleich ist. Es regelt die Datenübergabe zwischen den Schichten 2 und 3 und kann so z.B. die Grundlage für den Aufbau logischer Verbindungen vorbereiten.

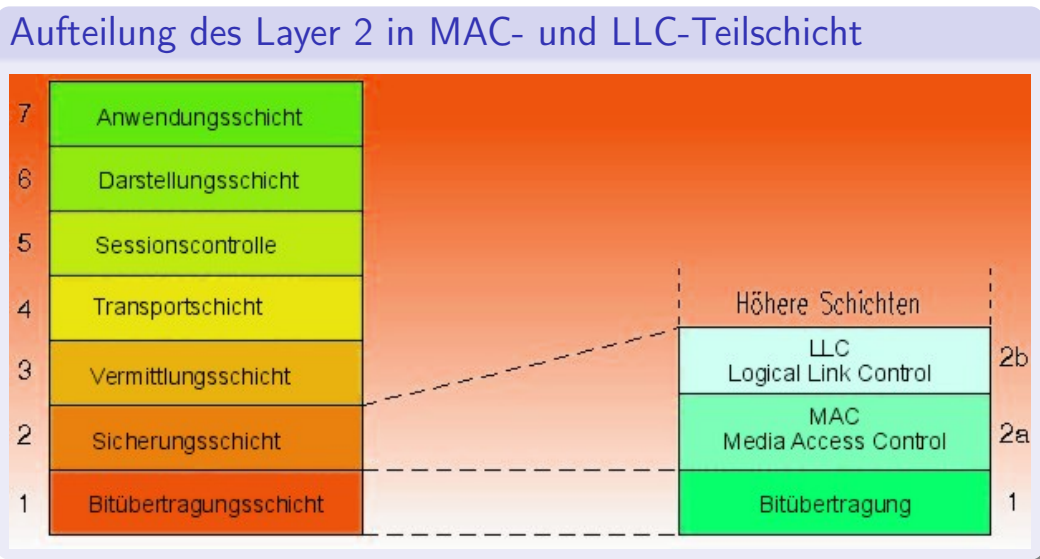


Abbildung 3.19: Auf Layer 2 sind die MAC-Adressen zu Hause.

¹Die erste Beschreibung dieses Protokolls wurde im Februar 1980 festgelegt, deshalb IEEE 802!

3.2.1.2 CSMA/CD

CSMA/CD ist **die** Zugangstechnik bei Ethernet und funktioniert wie folgt:

An das Kabel angeschlossene Endgeräte horchen (**carrier sense**) zunächst das Medium ab. Wenn es frei ist, beginnen sie mit der Übertragung ihrer Daten. Sollten nun zwei oder mehr Endgeräte gleichzeitig (**multiple access**) mit der Datenübertragung beginnen, führt dies zu einer Überlagerung der Signale. In der Netzwerktechnik spricht man dann von Kollisionen. Jedes Gerät im Netzwerk kann erkennen ob Kollisionen aufgetreten sind (**collision detection**). In einem solchen Fall wird das Senden der Daten sofort beendet. Damit jedes Gerät im Netzwerk-Segment erkennen kann daß es zu einer Kollision kam, wird für einen gewissen Zeitraum ein Störsignal (Jam) auf die Leitung gegeben. Nach Aussenden des Jam-Signals wird von jedem Endgerät ein zufällig gewählter kleiner Zeitraum gewartet, bevor die Übertragung wieder aufgenommen wird. Auf diese Weise kommt es normalerweise zu einer "Entzerrung" der Datenübertragung.

Der Zeitpunkt (und die Dauer) der Datenübertragung bei CSMA/CD ist also nicht vorhersagbar (nicht deterministisch). Man könnte auch sagen: Wer zuerst kommt malt zuerst. Alle Endgeräte, die untereinander Kollisionen erzeugen können befinden sich in einer Kollisionsdomäne (s.a. S. 39).

Bedeutung

CSMA/CD: **C**arrier **S**ense **M**ultiple **A**ccess with **C**ollision **D**etection

Ablauf von CSMA/CD

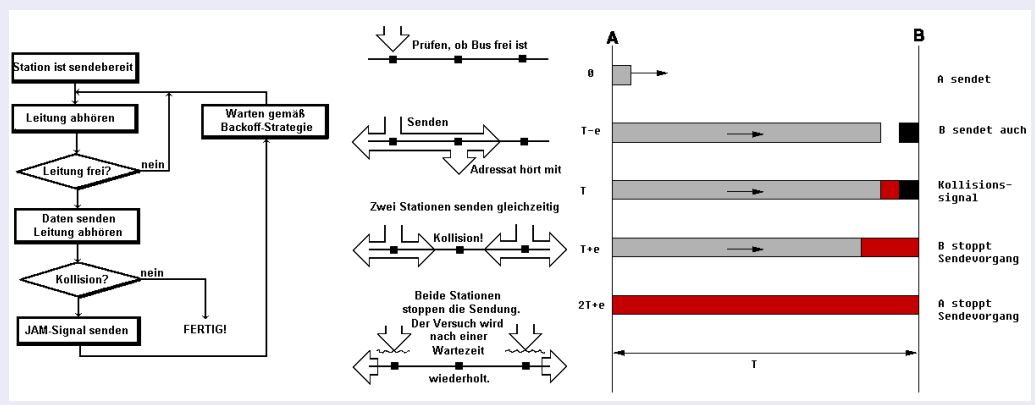


Abbildung 3.20: CSMA/CD ist **die** Zugriffstechnik bei Ethernet.

3.2.1.3 MAC-Adresse

Die MAC-Adresse ist eine weltweit eindeutige Kennzeichnung von Netzwerkschnittstellen und wird bei deren Produktion fest in die Hardware integriert. Man nennt sie deshalb auch *Burned In Addresses (BIA)*. Damit soll verdeutlicht werden, daß diese Adresse nachträglich nicht mehr zu ändern ist. Manche Netzwerkgeräte erlauben zwar die Eingabe einer anderen MAC-Adresse, die festeingetragene wird dadurch aber nicht überschrieben!

MAC-Adressen werden von vielen Layer 2 Technologien, z.B: Ethernet², Token Ring, Bluetooth um nur einige zu nennen, zur eindeutigen Identifizierung von Netzwerkgeräten auf Layer 2 Ebene benutzt. MAC-Adressen sind 48 Bit breit, was theoretisch 2^{48} (mehr als 281 Billionen) verschiedenen MAC-Adressen erlaubt. Besondere MAC-Adressen sind:

- 00:00:00:00:00:00 wird nicht verwendet
- FF:FF:FF:FF:FF:FF ist der ARP-Broadcast

Die ersten drei Bytes der MAC-Adresse werden den Netzerkäufern von der IEEE zugewiesen. Die letzten drei Bytes vergibt der Hersteller dann selbständig.

MAC-Adressen

Using this page you can search for an OID if you know the manufacturer, or for a manufacturer if you know its OID. This way you can find who made your network adapter, given the mac address and vice-versa.

Type the OID (first six hex digits of your mac address) you're looking for, using no space or other symbols, i.e. **000001**

Type the manufacturer you're looking for, or just a part of its name, i.e. **xerox corporation** or **intel**

[Back to goriani.com public portal home](#)

[Mac Makeup, the free mac address spoofing tool for windows](#)

Also jump to the [portal](#) if you think this database is incomplete (I know it is) and want to send me more info

There are **2** manufacturers matching **Zyxel**

zyxel communication	00a0c5
zyxel communications, inc.	004001



Weitere Informationen z.B. unter
<http://standards.ieee.org/regauth/oui/index.shtml>

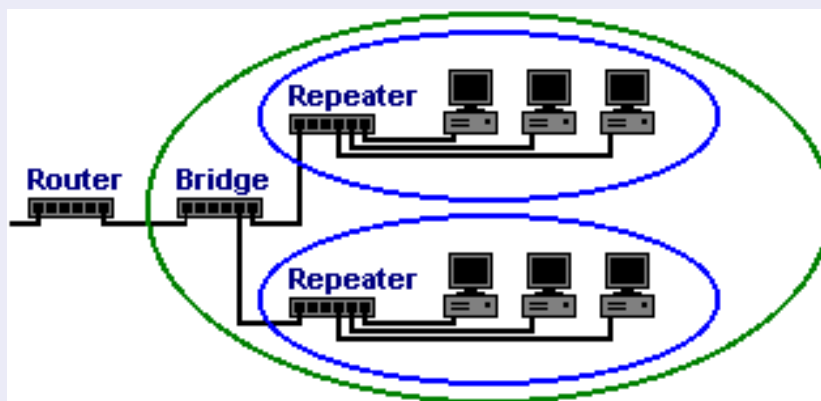
Abbildung 3.21: Zuordnung Hersteller/MAC-Adresse

²Tatsächlich wird die MAC-Adresse manchmal auch *Ethernet-Adresse* genannt, weil sie in der Ethernet-Spezifikation zum ersten Mal auftauchte.

3.2.1.4 Broadcastdomäne

Ein Broadcast ist ein Datenpaket, daß an alle Teilnehmer in einem Netzwerk gesendet wird. Eine Broadcastdomäne besteht aus allen Netzwerkgeräten, die über einen Broadcast zu erreichen sind. Weil Layer 2-Netzwerkgeräte die MAC-Adresse FF:FF:FF:FF:FF:FF nicht lernen können sondern an alle anderen Geräte weiterleiten bildet die zweiten Schicht des OSI-Schichtenmodells mit ihren Bridges und Switches und Netzwerkkarten eine große Broadcast-Domäne.

Broadcast-Domain



Abgrenzung der Domänen

- Layer 1 → Kollisionsdomäne → begrenzt durch L2-Geräte
- Layer 2 → Broadcastdomäne → begrenzt durch L3-Geräte

Eine Broadcast-Domäne (grün) kann aus mehreren Kollisionsdomänen (blau) bestehen.

Abbildung 3.22: Router erzeugen neue Broadcast-Domänen. Vergleiche dazu auch die Abbildung auf S. 39!

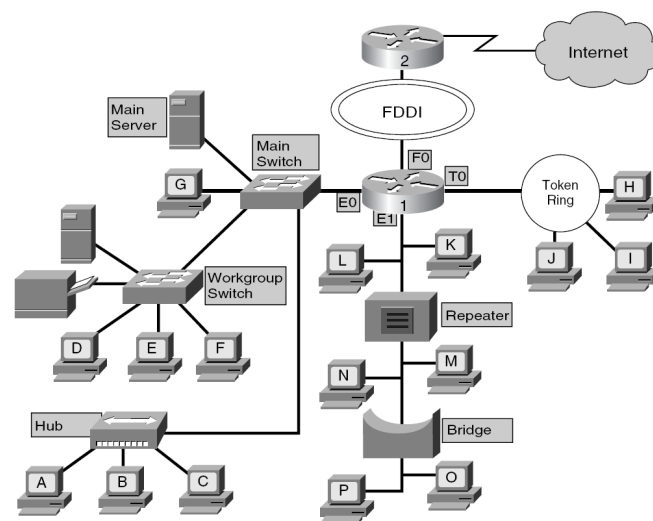
3.2.2 Logische Netzwerktopologie

Die Netzwerktopologie soll grafisch darzustellen, wie die Geräte eines Netzwerks³ logisch zusammenhängen. Die Kenntnis der Netzwerktopologie ist sehr wichtig um die Funktionsweise des Netzwerks zu verstehen.

Physische Topologie: Unter der physischen (manchmal liest man auch physikalischen) Topologie eines Netzwerks wird der Lauf der Kabel, also die Verkehrswege der Daten, sowie die Anordnung der einzelnen Netzwerkgeräte verstanden. Es handelt sich dabei um eine reine Layer-1-Angelegenheit.

Logische Topologie: Unter der logischen Topologie eines Netzwerkes versteht man, wie auf das Übertragungsmedium zugegriffen wird. Diese Zugriffsverfahren sind über die damit verbundenen Zugriffsprotokolle (CSMA/CD bei Ethernet und Token Ring bei Ring-Netzwerken) definiert.

Figure 2-40 Networking Topologies



In dieser Abbildung sind viele verschiedene physische und logische Topologien dargestellt. Es handelt sich um ein mittelmäßig komplexes Netzwerk, wie man es etwa in einer Schule oder einem kleinen Unternehmen finden könnte.

Abbildung 3.23: Verschiedene Topologien in einem Netzwerk

³z.B. Drucker, Arbeitsplatzrechner, Server. usw.

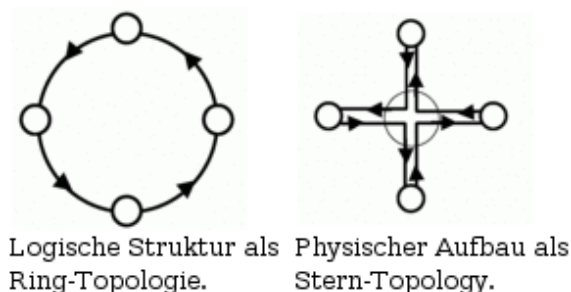
3.2.2.1 Logische Topologien

Die logische Topologie beschreibt das Zugriffsverfahren auf das Medium. Die beiden wichtigsten Zugriffsverfahren sind *broadcast* und *token passing*.

broadcast: Bei der Broadcast-Topologie schickt ein Host seine Daten mit Hilfe einer zuvor per ARP-Request ermittelten MAC-Adresse an den Ziel-Host. Sollen die Daten für alle Hosts bestimmt sein, dann werden sie an die Broadcast-Adresse gesendet. Für das Versenden (also z.B. wer wann wieviel und wie lange darf) der Daten gibt es keine Vorschriften, d.h wer zuerst kommt, mahlt zuerst. Ethernet ist ein typischer Vertreter dieser Topologie.

token passing: Hier wird die Zugriffsabfolge auf das Netzwerkmedium durch ein Token geregelt, das nacheinander alle Stationen passiert. Wenn ein Host das Token erhält, kann er Daten senden, ansonsten gibt er das Token an die nächste Station weiter. Token Ring und FDDI sind zwei Vertreter der token passing-Topologie.

Logische und physische Topologie können voneinander abweichen. So kann Ethernet physisch als Stern oder als Bus aufgebaut sein - logisch gesehen ist es immer eine Bus-Topologie, da der Datenfluss von einem Endgerät gleichzeitig zu allen anderen Endgeräten erfolgt.



Token Ring kann physisch auch als Stern über einen Ringleitungsverteiler realisiert werden, ist logisch jedoch immer eine Ring-Topologie, da der Datenfluss nacheinander von einem Endgerät zum Endgerät läuft. Bei token passing wurde dies in der Praxis auch fast immer so realisiert, weil es aus praktischen Gründen meist sehr schwierig war, die Kabel wirklich als Ring zu verlegen.

Abbildung 3.24: Logische und physische Topologien müssen nicht identisch sein.

3.2.3 Netzwerkgeräte

3.2.3.1 Netzwerkkarten

Auf Layer 2 findet man Netzwerkgeräte, die über "Intelligenz" verfügen, d.h. den Datenstrom schon auf Inhalte untersuchen können. Im Gegensatz dazu liest man oft, daß Layer 1 Netzwerkgeräte "dumm" sind, und zwar deshalb, weil sie Signale lediglich regenerieren und weiterleiten, den Datenstrom aber **nicht** weiter analysieren.

Klassische Netzwerkkarten in Form von Einsteckplatinen für den PCI- oder ISA-Bus verlieren immer mehr an Bedeutung. Jedes moderne Motherboard verfügt im Regelfall über eine fest eingebaute Netzwerkkarte. Bei Laptops kommt zur klassischen Ethernet-Buchse dann meist noch ein ebenfalls fest integrierter WLAN-Adapter hinzu, so daß diese Geräte dann sogar über zwei Netzwerkkarten verfügen!

Vorsicht ist geboten bei ganz neuen Motherboards und/oder neuen Netzwerk-Chipsätzen: Oftmals müssen die Treiber unter Windows dann von Hand mit der beigelegten CD installiert werden, während für Linux ev. noch gar keine Treiber vorhanden sind!

Einige Netzwerkkarten...



Abbildung 3.25: Eine Auswahl verschiedener Netzwerkkarten und Zubehör.

3.2.3.2 Bridges

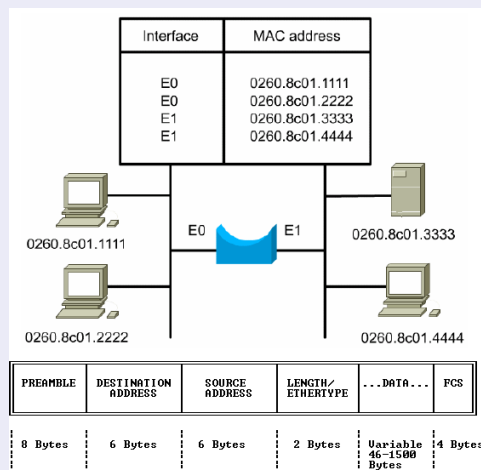
Bridges waren die ersten Geräte mit denen die Anzahl der Kollisionsdomänen in einem Netzwerk erhöht werden konnte. Bridges laden einen Ethernet-Frame vollständig in den Speicher, ermitteln die Ziel-MAC-Adresse und senden den Frame nur auf dem Port raus, hinter dem der Zielrechner ist. Da Bridges lernfähig sind, speichern sie diese Information in einer CAM-Tabelle⁴. Diese Vorgänge waren bei Bridges komplett in Software realisiert und wurden deshalb entsprechend langsam abgearbeitet. Für 10Base2-Netzwerke reichte diese Geschwindigkeit jedoch völlig aus.

Heutzutage sind Bridges in den meisten Fällen durch Switches ersetzt worden.

Vorteile

- durch Filterfunktion wird die Anzahl der Kollisionsdomänen erhöht
- Broadcast- und Multicast-Frames werden an alle Ports weitergegeben

Funktionsweise



Beim Bridging laufen folgende Vorgänge ab:

- zu Beginn ist die Tabelle leer
- der erste Frame jeder Station wird an alle Ports weitergeleitet (flooding, die Quelladresse des Frames wird mit der Portnummer in die Tabelle geschrieben)

Im laufenden Betrieb werden die Frames nur noch weitergeleitet, wenn sich das Ziel in einem anderen Segment befindet. Auch unbekannte Ziel-MAC-Adressen werden weitergeleitet, und zwar an alle Ports

Animation der Funktionsweise einer CAM-Tabelle

Abbildung 3.26: Bridges analysieren MAC-Adressen und erzeugen Kollisionsdomänen.

⁴Content Addressable Memory: Zuordnung MAC-Adresse/Port

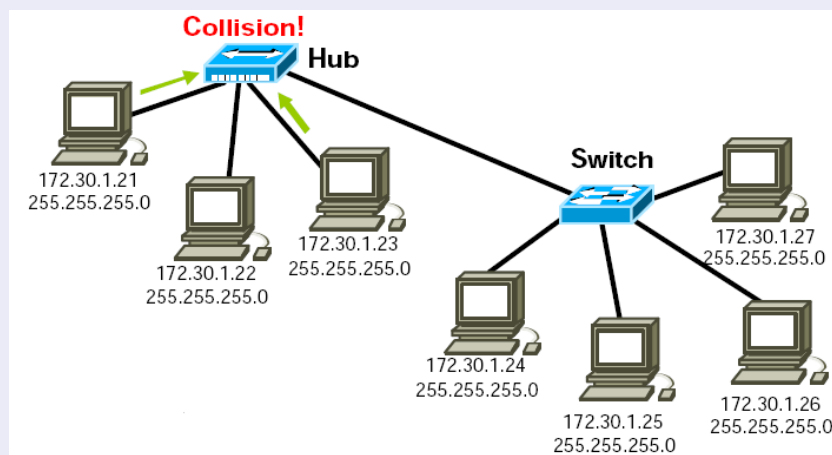
3.2.3.3 Switches

Gestiegene Anforderungen an die Geschwindigkeit (wegen zunehmender Datenmengen), das Aufkommen des Fast-Ethernet-Standards und höhere Portdichten am Gerät (immer mehr Arbeitsstationen) führten zu einem neuen Gerätetyp, dem Switch. Switches arbeiten die Frame-Zustellung in Hardware ab und erreichen deshalb extrem hohe Geschwindigkeiten. Durch den Einsatz spezialisierter Netzwerkprozessoren können Switches auch Zusatzfunktionen wahrnehmen, wie das Überwachen eines Ports auf geänderte MAC-Adressen (Port-Security), VLAN-Filterung usw.. Der Speicherbereich eines Switches beträgt mehrere tausend MAC-Adressen (von 1024 über 2048 bis zu 4096 oder noch mehr).

Eine heutzutage eher selten anzutreffende Geräteklasse ist der **Switching Hub**. Man verstand darunter ein Gerät, an das gleichzeitig 10 und 100 Mbit Netzwerkgeräte angeschlossen werden konnten. Zwischen 10 und 100 Mbit-Backplane war ein Switch eingebaut. Die 10 bzw. 100 Mbit Backplane arbeitete intern aber als Hub!

Moderne Switches handeln ihre Übertragungsparameter (Duplex, Geschwindigkeit) an den einzelnen Ports selbstständig aus. Selbst Crossover-/Straight Through-Mischsysteme sind kein Problem mehr (AUTO-MDI). Geschwitchte Netzwerke arbeiten kollisionsfrei. Jedes Gerät ist in einer eigenen Kollisionsdomäne.

Switches und ihre prinzipielle Funktion



Animation der Funktionsweise eines Hubs

Animation der Funktionsweise eines Switches

Abbildung 3.27: Switches sind hochperformante Bridges.

3.2.3.3.1 Geräteklassen: Managed und Unmanaged Switches

Diese Gerätegruppe ist aus modernen Ethernet-Netzwerken nicht mehr wegzudenken. LAN-Switch, Workgroup-Switch, oder einfach nur Switch sind dabei nur verschiedene Ausdrücke für ein und denselben Gerätetyp.

Switches funktionieren wie Bridges, erledigen ihre Arbeit aber wesentlich schneller⁵ und können deshalb auch für Fast- oder Gigabit-Ethernet eingesetzt werden. Aufgrund der Ähnlichkeit in der Arbeitsweise werden Switches oft auch als Multiport-Bridges bezeichnet. Damit wird zum Ausdruck gebracht, daß sich jeder Port eines Switches wie eine Bridge verhält und damit den angeschlossenen Host in eine eigene Kollisionsdomäne stellt. Das wird dann oft auch als Microsegmentation⁶ bezeichnet.



Unmanaged Switches:

- Plug and Play
- Meist zwischen 4 bis 24 Ports
- Desktop tauglich da oft schicke Gehäuse und lüfterlos
- int./ext. Stromversorgung
- Einsatzbereich von schnell mal ein paar Rechner verbinden bis Workgroup-Switch

Managed (oder intelligente) Switches:

- Vorwiegend im professionellen Umfeld eingesetzt
- meist VLAN-fähig
- meist nur ab höheren Portzahlen (12 bzw. 24 aufwärts) verfügbar
- meist modular erweiterbar (Module für mehr Ports, Lichtwellenleiter, Gigabit usw...)

Abbildung 3.28: Einige Switches für unterschiedliche Ansprüche.

⁵u.a. auch deshalb, weil der Switching-Algorithmus (das Auswerten der MAC-Tabelle) mit ASIC's, also speziell programmierten Hardwarebausteinen, erledigt wird. Bridges mußten diese Arbeit in Software erledigen.

⁶Microsegmentation bedeutet dann natürlich auch, daß dem Host die volle Bandbreite zur Verfügung steht. In diesem Segment gibt es dann auch keine Kollisionen weil konkurrierende Hosts fehlen.

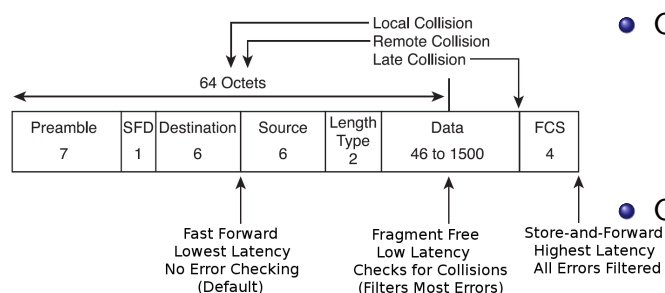
3.2.3.3.2 Switching Methoden

Switches kennen zwei prinzipiell unterschiedliche Zustellmöglichkeiten für Frames:

- Store-and-forward: Der gesamte Frame wird in den Puffer des Switches eingelesen. Dann wird die Prüfsumme (CRC) berechnet. Entspricht sie der mitgelieferten, wird der Frame weitergeleitet, sonst verworfen.
- Cut-Through: Der Frame wird bereits weitergeleitet bevor er vollständig empfangen wurde. Man unterscheidet zwei Technologien:
 - Fast-Forward switching: Es werden nur die ersten 14 Bytes (8 Byte Preamble und 6 Byte Destination MAC-Address) eingelesen, dann wird der Frame weitergeleitet.
 - Fragment-free switching: Es werden nur die ersten 64 Bytes eingelesen (um Kollisionsfragmente herauszufiltern) und der Frame dann weitergeleitet.

Viele Hersteller bieten zusätzlich noch einen dynamischen portbasierten Switching-Mode an: Hier werden alle Pakete zunächst per Fast-Forward zugestellt. Übersteigt die Fehlerrate einen einstellbaren Grenzwert, wird in den Store-and-Foreward Mode geschaltet, sinkt die Fehlerrate, wird wieder im Fast-Forward-Mode weitergemacht.

Collision Types: Local, Remote, and Late



- Store-and-Foreward
 - langsam
 - maximale Sicherheit beim Erkennen defekter Frames
- Cut-Through (Fast-Forward)
 - sehr schnell
 - fehlerhafte Pakete werden nicht aussortiert
- Cut-Through (Fragment-Free)
 - schnell
 - Pakete kleiner 64 Bytes (also Kollisionsfragmente) werden aussortiert

Abbildung 3.29: Die 3 wichtigen Betriebsarten von Switches.

3.2.3.3.3 Spanning Tree

Das STP (Spanning Tree Protocol) ist ein nach IEEE 802.1d standardisiertes Protokoll mit dem Switches verhindern, daß es in Netzwerken zu Schleifen kommt. Schleifen in Netzwerken bedeuten immer, daß das Netzwerk nicht mit optimaler Leistung läuft (Performance-, Durchsatzprobleme usw. . .). Mit STP können Switches (durch versenden kleiner Datenpakete, sogn. BPDU) Netzwerk erkennen und deaktivieren. Damit kann man jetzt sogar absichtlich Schleifen erzeugen (Ersatzleitungen, sogn. redundante Datenleitungen). Fällt die Hauptverbindungsleitung aus, dann wird vom Switch die Ersatzleitung aktiviert. Der optimale Weg im geschwitchten Netzwerk wird anhand des *Cost*-Parameters ermittelt:

In modernen Geräten findet man häufig auch RSTP implementiert. Der Vorteil von RSTP ist, daß neue Wege innerhalb des Netzwerks im Bedarfsfall innerhalb weniger Sekunden berechnet werden können. Bei STP kann dieser Vorgang bis zu 30 Sekunden dauern!

Link Speed	Cost
10 Gbps	2
1 Gbps	4
100 Mbps	19
10 Mbps	100

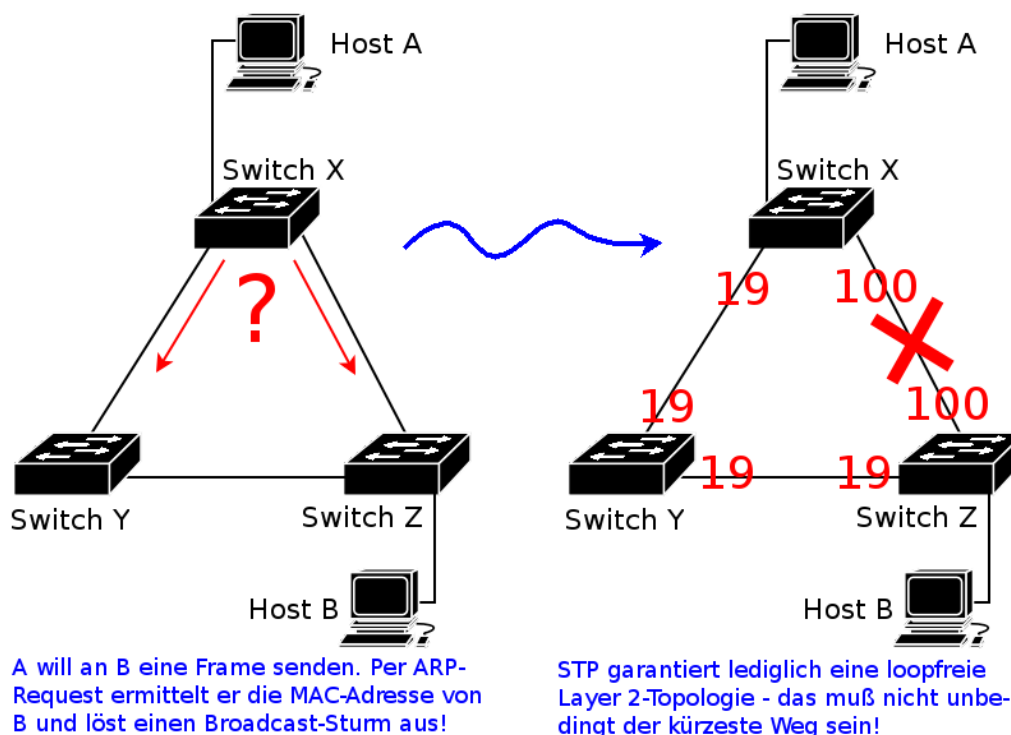


Abbildung 3.30: Prinzipielle Funktionsweise von STP

3.2.3.3.4 Zusammenfassung

Moderne Switches beherrschen asymmetrisches Switching. Sie sind in der Lage, aufgrund eines Speichers, Pakete zu puffern. Solche Geräte können in Mischnetzwerken (10/100 Mbit/s) als Bindeglieder eingesetzt werden. Symmetrische Switches können nur in Netzwerken eingesetzt werden, die eine Geschwindigkeit (10 oder 100 Mbps) haben.

Manche Switches (z.B. Cisco Catalyst 2950) kennen nur den Store- and Forward-Mode. Häufig können solche Switches trotzdem (wegen hochleistungsfähigen ASICs)⁷ performanter sein also solche, die zusätzlich noch dynamisches Switching bzw. Cut-Through-Mechanismen beherrschen. Beim Kauf eines Switches sollte also zwischen Kosten, Performance und Features des Gerätes eine Abwägung stattfinden.

VLAN-fähige Switches heben die Sicherheit eines Netzwerks auf ein höheres Niveau und sind heute praktisch als Standard bei Neuvernetzungen anzusehen.

Vergleich von Bridges und Switches

Funktion	Bridges	Switches
Art des Switching	Software	Hardware (in ASIC's)
Switching Methode(n)	Store-and-forward	Store-and-forward Cut-Through fragment-free
Ports	2 – 16	bis zu mehreren 100
Duplex	Halb-Duplex	Voll-Duplex
Collision Domains	1 per Port	1 per Port
Broadcast Domains	1	1 (bzw. 1 per VLAN)

Abbildung 3.31: Unterschiede und Gemeinsamkeiten von Bridges und Switches.

⁷Application Specific Integrated Circuit

3.3 Layer 3: Vermittlungsschicht

3.3.1 IP

3.3.1.1 Das überholte Klassenschema der IP-Adressen

Im Dezember 1974 wurde von V. Cerf und zwei weiteren Mitarbeitern der Stanford University als RFC 675 die erste technische Beschreibung von TCP/IP veröffentlicht. Sie sah eine 32 Bit breite "IP-Adresse" vor. Die ersten 8 Bit sollten dabei das Netzwerk adressieren, die restlichen 24 Bit waren für Host-Adressen vorgesehen. So konnte man 256 Netze mit 16.777.216 eindeutigen Host-Adressen betreiben.

Gegen 1981 kamen erste Zweifel auf, ob ein Netzwerk mit maximal 256 Netzen zukunftsfähig ist. RFC 791 schlug daraufhin die Aufteilung in drei Netzwerkklassen A, B und C, mit Netzwerkmasken von jeweils 8, 16 bzw. 24 Bits vor. Ende der 1980er Jahre erschöpfte sich der Vorrat an Klasse B-Netzen allmählich. Aus diesem Grund wurde 1992 RFC 1338 und ein Jahr später RFC 1519 auf den Weg gebracht, um das Classless Inter-Domain Routing (CIDR) einzuführen. Heute ist RFC 4632 der dafür gültige Standard.

Auch wenn RFC 791 formal nie abgeschafft wurde, spielt die Drei-Klassen-Gesellschaft in der Internet-Realität keine Rolle mehr. [15][16]

Klasseneinteilung der IP-Adressen

Klasse	Addressbereich	Anzahl der Netzwerke	Hosts pro Netzwerk
A	1.0.0.0 ... 126.255.255.255	126	16.777.214
B	128.0.0.0 ... 191.255.255.255	16.384	65.534
C	192.0.0.0 ... 223.255.255.255	2.097.152	254
D	224.0.0.0 ... 239.255.255.255	-	-
E	240.0.0.0 ... 254.255.255.255	-	-

Abbildung 3.32: Die wichtigsten IP-Klassen A, B und C im Überblick.

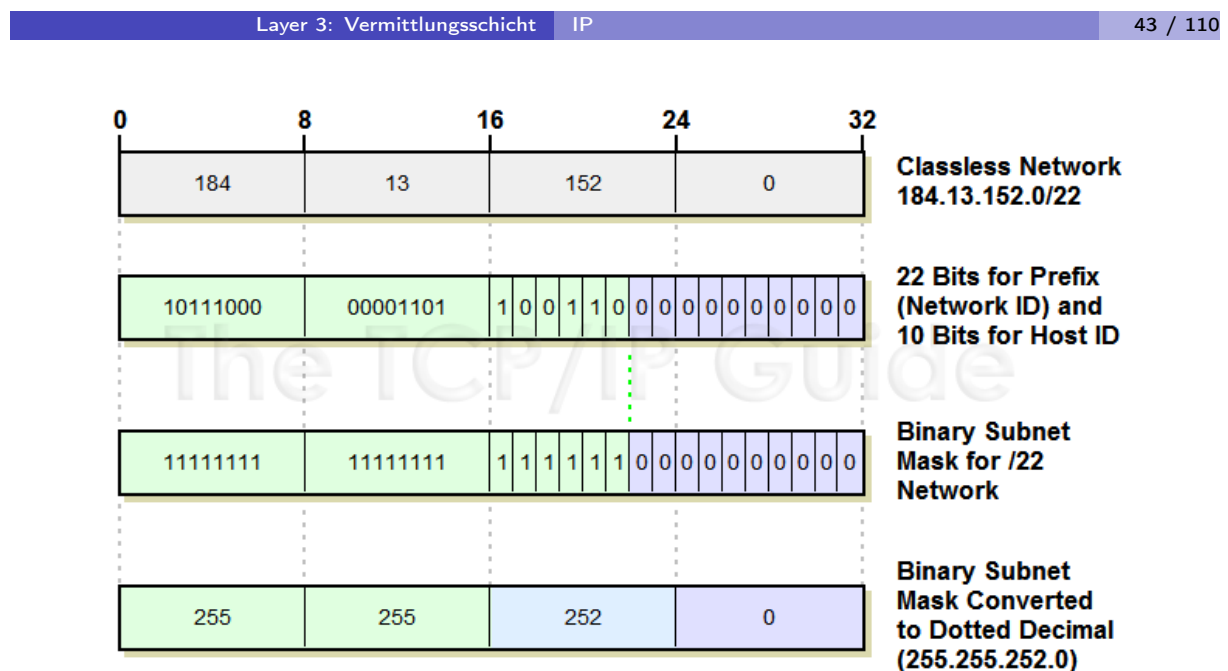
3.3.1.2 CIDR und VLSM

Ursprünglich wurde bei IP-Netzwerken die Trennlinie zwischen Netzwerk- und Hostanteil immer nach 8 Bits gezogen. Jedes Netzwerk bestand damit aus 2^{24} Hosts. Durch die Einführung von klassenbasierten Netzwerken wurde etwas mehr Flexibilität erreicht. Allerdings galt die Subnetzmaske immer für die gesamte Netzwerkadresse!

Klassenlose Netzwerke können nun innerhalb **einer** Netzwerkadresse verschiedene Subnetzmasken verwenden. Damit ist es möglich, die verfügbaren IP-Adressen des Netzwerks sehr effizient zuzuteilen.

Um zu verdeutlichen, daß es sich bei einer IP-Adressen-/Netzwerkmasken-Kombination um eine klassenlose Angabe handelt, wird die Netzwerkmaske nicht mehr in der Punktform angegeben. Vielmehr wird der IP-Adresse nach einem Slash-Symbol die Anzahl der Subnetzmask-Einsen nachgestellt. Diese Schreibweise wird Slash- oder CIDR-Notation (RFC 1817) genannt.

CIDR und VLSM beschreiben im Prinzip den gleichen Vorgang: Ein großes Netzwerk soll in kleinere Netzwerke, ggf. auch unter Verwendung unterschiedlicher Subnetzmasken, zerlegt werden. VLSM wird vom Eigentümer eines IP-Adressblocks durchgeführt, CIDR macht der Internet Service Provider!



Ein klassenloses Netzwerk wird üblicherweise in CIDR- oder Slash-Notation dargestellt. Damit kann man Routing-Tabellen drastisch verkleinern.

[► Details...](#)

Abbildung 3.33: Ohne CIDR wäre das Internet in der heutigen Form nicht denkbar.

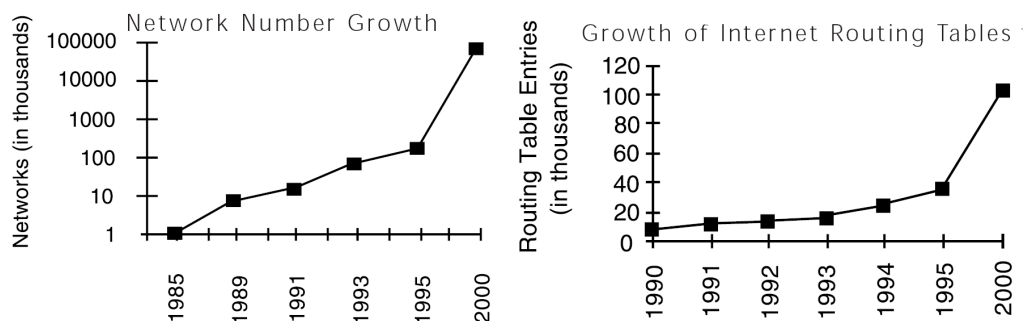
3.3.1.3 Probleme des IPv4-Adreßraums und Gegenmaßnahmen

3.3.1.3.1 Einleitung: Durch das rasante Wachstum des Internets in den letzten 10-15 Jahre traten zwei Probleme immer deutlicher in den Vordergrund:

- Der, gemessen an heutigen Ansprüchen, zu kleine Adreßraum von IPv4. Eine IP-Adresse wird hier durch eine 32-Bit-Zahl dargestellt. Die damit möglichen $2^{32} = 4294967296$ Adressen reichen aber nicht mehr aus um allen Geräten eine routbare IP-Adresse zuzuweisen.
- Der immer größere Aufwand, die permanent steigenden Anzahl von Netzen zu routen.

Zur Zeit behilft man sich deshalb u.a. mit Subnetting und NAT. Langfristig wird IPv4 durch die 128 Bit breite IPv6 IP-Adresse ersetzt werden.

Mehr als die Hälfte aller [IP-Adressen](#) sind bereits vergeben. Aktuell genutzt werden davon aber nur ca. 69 Millionen.



Ende September 2000 gab es ca. 88000 Routen: Das Routing-Problem kann nicht nur dadurch gelöst werden, daß immer mehr Speicher und schnellere Prozessoren in den Routern installiert werden um die ständig wachsenden Routingtabellen zu verarbeiten.

Abbildung 3.34: Das rasante Wachstum des Internets stellt eine große Herausforderung dar.

3.3.1.4 Private und besondere IP-Adressen

März 1994: Die Network Working Group veröffentlicht unter dem RFC 1597 ein Dokument, das die Adresszuweisung für private Netzwerke regelt. Es war bereits abzusehen, daß die zur Verfügung stehenden IP-Adressen nicht ausreichen werden, um alle netzwerkfähigen Geräte damit auszustatten. Aus diesem Grund entschloß man sich, aus jeder IP-Addressklasse einen Bereich herauszunehmen und zur privaten Nutzung (also ohne Absprache mit der IANA, welche die IP-Adressvergabe regelt) zur Verfügung zu stellen. Einzige Bedingung: Diese Adressen dürfen nicht ins Internet geroutet werden (weil, wenn sie jeder für sich verwenden kann wie er will, keine Eindeutigkeit mehr vorhanden ist). Dafür tragen heute Provider wie die Telekom, Versatel o.ä., die Verantwortung. Der Vorteil ist, daß man nun auch Netzwerke unter **einer öffentlichen IP-Adresse** ansprechen kann, obwohl sich dahinter vielleicht ein privates Class-C-Netzwerk mit 100 Benutzer befindet. Wer zu Hause DSL nutzt und mit zwei (oder mehr PC's) gleichzeitig im Internet ist, wendet genau diese Technik an.

Noch ein Netzwerk hat besondere Bedeutung: Es handelt sich um 127.0.0.0, ein komplettes Class-A-Netz, aus dem sich jeder TCPI/IP-fähige Rechner beim Start seine localhost-Adresse (i.d.R. 127.0.0.1) holt. Diese IP-Adresse funktioniert übrigens auch in Rechnern **ohne** Netzwerkkarte. Es muß lediglich ein TCP/IP-Protokollstack installiert sein.

3. Private Address Space

The Internet Assigned Numbers Authority (IANA) has reserved the following three blocks of the IP address space for private networks:

10.0.0.0	-	10.255.255.255
172.16.0.0	-	172.31.255.255
192.168.0.0	-	192.168.255.255

We will refer to the first block as "24-bit block", the second as "20-bit block", and to the third as "16-bit" block. Note that the first block is nothing but a single class A network number, while the second block is a set of 16 contiguous class B network numbers, and third block is a set of 255 contiguous class C network numbers.

127.0.0.0 ist für lokale Zwecke reserviert

localhost 127.0.0.1

Abbildung 3.35: Funktionsweise von Multicast-Adressen

3.3.1.5 Multicast-Adressen

Multicast kann immer dann verwendet werden, wenn ein Sender an mehrere (aber nicht alle) Empfänger in einem Netzwerk übertragen werden soll. Zum Einsatz könnte dies z.B. bei der Übertragung von Video- bzw. Audiodaten einer Firmenkonzferenz sein. Damit ähnelt Multicast sehr einer Radio oder TV-Übertragung, weil tatsächlich nur die Empfänger erreicht werden, die in der entsprechenden Multicast-Gruppe sind. Mit anderen Worten: Man hört oder sieht genau das, was einen interessiert. Jedes Mitglied im Netzwerk kann sich seinen "Kanal" selbst auswählen! Damit sind Multicast-Anwendungen typische Vertreter von Punkt-zu-Mehrpunktverbindung.

Standardmäßig arbeitet Multicast mit UDP-Datagrammen, ist also "nicht zuverlässig" wenn die Zustellung der Daten garantiert sein muß. Aus diesem Grund wurde das MTP-Protokoll entwickelt, das sicher stellt, daß auch wirklich alle Mitglieder einer Multicast-Gruppe die Daten erhalten haben!

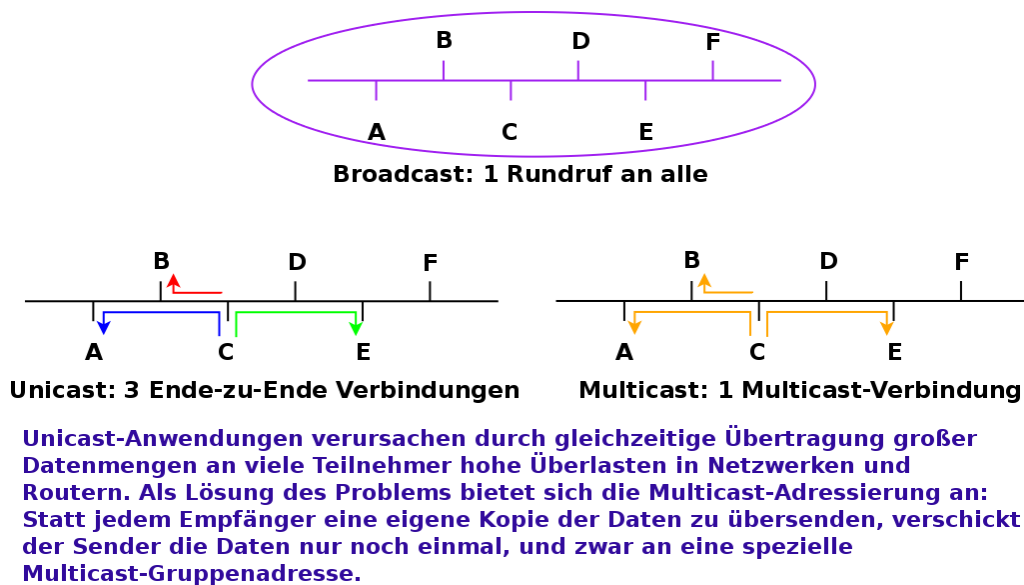


Abbildung 3.36: Private IP-Adressen werden nicht ins Internet geroutet!

3.3.1.6 Das IP-Protokoll

IP stellt die Basisdienste für die Übermittlung von Daten in TCP/IP-Netzen bereit und ist im RFC 791 spezifiziert. Hauptaufgaben des Internet Protokolls sind die Adressierung von Hosts und ggf. das Fragmentieren von TCP-Paketen. Die so erzeugten Pakete werden von IP nach bestem Bemühen ("best effort") von der Quelle zum Ziel befördert, unabhängig davon, ob sich die Hosts im gleichen Netz befinden oder andere Netze dazwischen liegen. Garantiert ist die Zustellung allerdings nicht. Das Internet Protokoll enthält keine Funktionen für die Ende-zu-Ende-Sicherung oder für die Flußkontrolle.

IP ist ein verbindungsloses Protokoll, d.h. zur Datenübertragung wird keine Ende-zu-Ende-Verbindung der Kommunikationspartner aufgebaut. Ferner ist IP ein unzuverlässiges Protokoll, da es über keine Mechanismen zur Fehlererkennung und -behebung verfügt. Unzuverlässig bedeutet aber keinesfalls, daß man sich auf das IP Protokoll nicht verlassen kann. Unzuverlässig bedeutet in diesem Zusammenhang lediglich, daß IP die Zustellung der Daten nicht garantieren kann (also daß alle Pakete eines Datenstroms überhaupt bzw. in der richtigen Reihenfolge ankommen). Sind die Daten aber beim Zielhost angekommen, dann sind sie im Regelfall inhaltlich auch korrekt (wg. der FCS die immer überprüft wird).

Aufbau des IP-Headers

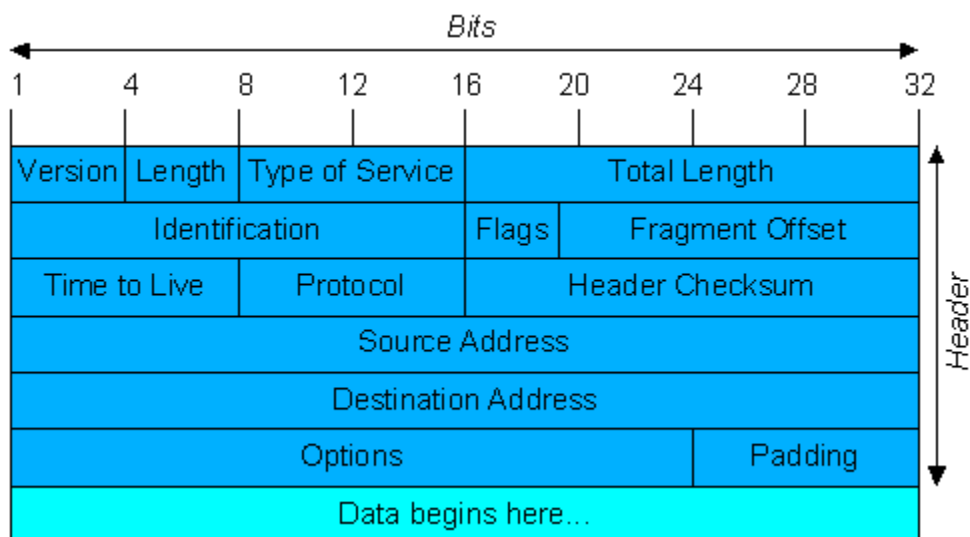


Abbildung 3.37: Aufbau des IP-Headers.

3.3.2 Netzwerkgeräte

3.3.2.1 Router

Router verbinden Netzwerke. Sie erreichen dies mit Hilfe von Netzwerkadressen (einer besonderen IP-Adresse, bei der alle Host-Bits auf Null gesetzt sind) und speziellen Routing-Protokollen. So sind Router jederzeit über den Status des Netzwerks informiert und können, mit Hilfe von Routingtabellen, einen Weg vom Quell- zum Zielnetz ermitteln.

Es gibt verschiedene Ansätze um Routing Protokolle zu implementieren. Zwei der wichtigsten Vertreter dieser Ansätze sind RIP und OSPF. Benachbarte Router müssen das gleiche Protokoll beherrschen um Informationen über ihre angeschlossenen Netze austauschen zu können. Oft laufen auch verschiedene Routing-Protokolle nebeneinander auf einem Gerät.

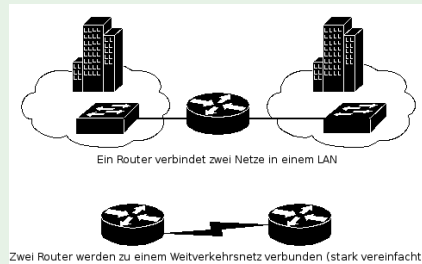
Daneben gibt es noch den Begriff des Routed Protocols. Darunter versteht man ein Protokoll, das sich dazu eignet, ein Weitverkehrsnetz zu durchqueren. Das bekannteste Routed Protocol ist IP.

So erhalten sie Informationen über entfernte Netzwerke und ein weltweiter Datenaustausch wird möglich.

Beschreibung

- Verbindet Netze
 - innerhalb eines LAN's
 - Höhere Sicherheit (z.B. VLANs)
 - unterschiedlicher Technik
 - Ethernet und Token Ring
 - LAN und WAN
 - z.B. PPPoE bei DSL
 - WAN und WAN
 - Telko
- Arbeitet auf Layer 3 des OSI-Schichtenmodells
- Erzeugt neue Broadcast- und Kollisionsdomänen

Skizze



► Router-Modelle

Abbildung 3.38: Einsatzgebiete für Router.

3.3.2.1.1 Mögliche Entscheidungskriterien für die beste Route

Die Ermittlung der besten Route wird durch folgende Parameter beeinflusst:

- Bandbreite: Wie hoch ist die Kapazität einer Leitung (z.B. ISDN oder Gigabit)
- Verzögerung: Wie lange dauert die Datenübertragung (ein Wert, der z.B. von der physikalische Länge der Leitung, der Netzauslastung und einigen weiteren Parametern abhängt)
- Pfadkosten: Wird vom Administrator vergeben und kann sich z.B. an den tatsächlichen Nutzungskosten einer Leitung orientieren
- Hop-Count: Ein Wert der beschreibt, durch wieviele Router ein Datenpaket reisen muß um am Ziel anzukommen
- Zuverlässigkeit: Beschreibt die Fehlerrate einer Netzwerkverbindung
- Netzauslastung: Beschreibt die Aktivität einer Netzwerkverbindung

Routing Protocol Metrics

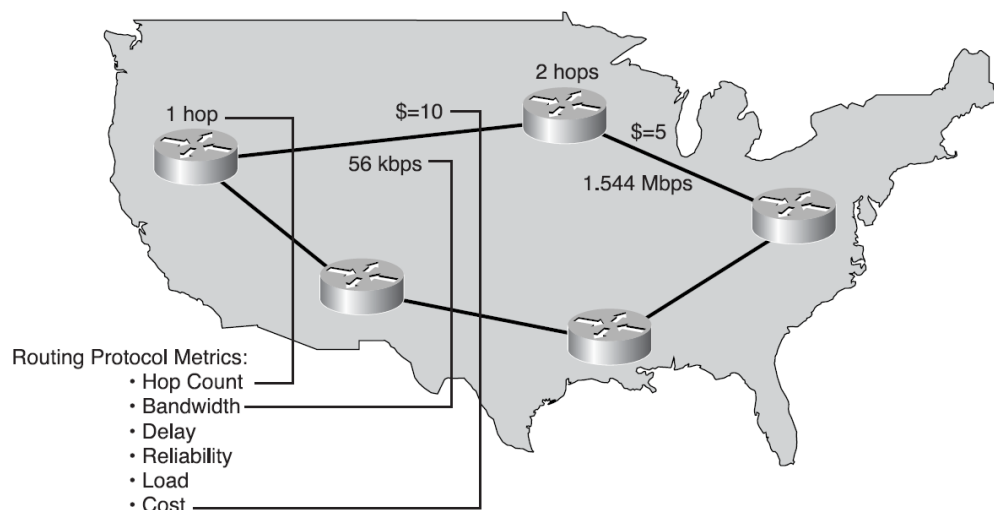


Abbildung 3.39: Prinzipielle Vorgehensweise bei der Wahl einer Route.

3.3.2.2 Routingprotokolle

3.3.2.2.1 RIP

RIP ist ein Routing-Protokoll auf Basis des Distance-Vector-Algorithmus. Es dient zur dynamischen Erstellung einer Routingtabelle. RIP wird vor allem in Verbindung mit dem IP-Protokoll benutzt.

Beim Starten eines Routers kennt dieser nur seine direkt angeschlossenen Netzwerke und erzeugt damit eine interne Routingtabelle. Diese Tabelle sendet er im Abstand von 30 Sekunden an benachbarten Router und erhält von diesen Geräten deren Routingtabelle. So kann die eigene Routingtabelle ergänzt werden und der Router "lernt", welche Netzwerke jeweils über welchen Port erreicht werden können. Dieser Vorgang läuft ohne Eingriff von außen, also dynamisch ab.

Bei Änderungen im Netzwerk (z.B. Ausfall eines Routers) dauert es eine gewisse Zeit, bis alle Router wieder eine einheitliche Sicht auf das Netzwerk haben. Ist dieser Zustand erreicht, spricht man von **Konvergenz** (oder einem konvergierten Netzwerk).

Die Kosten (auch Metrik) bezeichnen den Aufwand, um ein bestimmtes Netz zu erreichen. RIPv1 verwendet dazu den Hop Count. Er bezeichnet die Anzahl der Router, die entlang eines Pfades bis zum Zielnetz durchlaufen werden müssen.

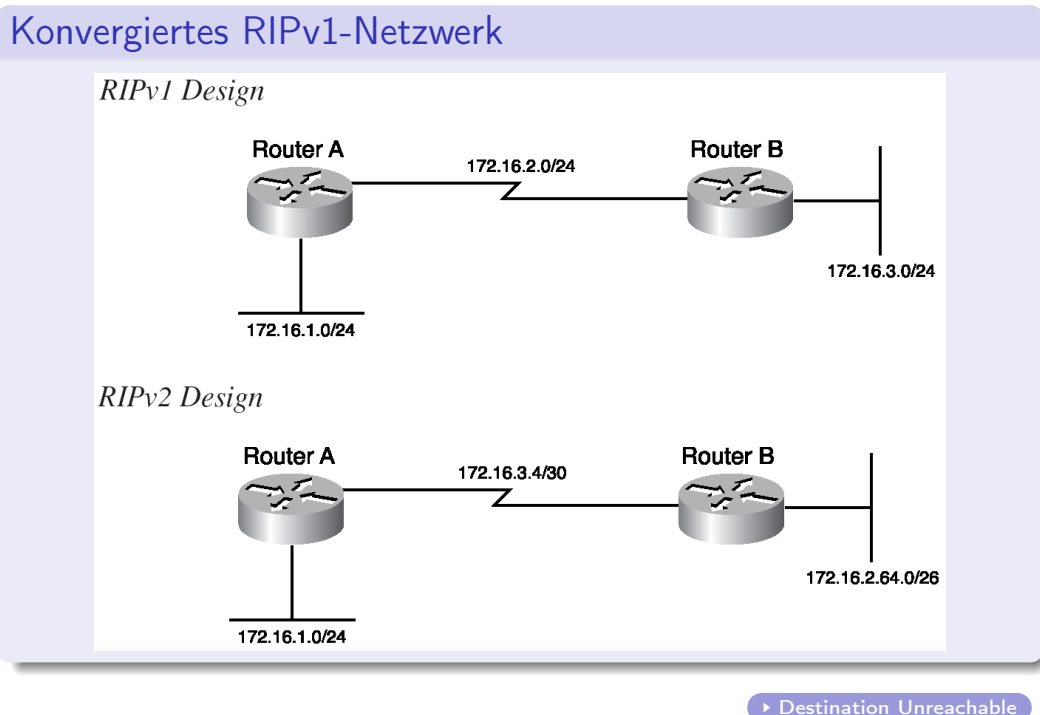


Abbildung 3.40: RIPv2 kann auch mit unterschiedlich langen Subnetzmasken umgehen.

3.3.2.2 Vergleich von RIPv1 und RIPv2

Neben der einfachen Bedienung und dem Quasi-Standard von RIP, der eigentlich (zumindest in der Version 1) von jedem Router beherrscht wird, gibt es folgende Nachteile zu bedenken:

- jeder Router kennt zu Beginn des Routingprozesses nur seine direkte Nachbarn. Wegen dieser lokalen Sicht müssen die Routen nicht immer die besten (z.B. schnellste Übertragung) sein
- Relativ lange Konvergenzzeit bei Änderungen im Netz (Updates alle 30 Sekunden)
- Nur eine Metrik (Hop-Count)
- Nur für kleine Netze geeignet, da max. 15 Hops
- Wegen Routing-Updates alle 30 Sekunden hohe Netzlast
- Nicht skalierbar

Unterschiede zwischen RIPv1 und v2:

Merkmal	RIPv1	RIPv2
RIP-Information	nur Broadcast	Broadcast oder Multicast
IP-Adressen	classfull	classless (weil Subnetzmaske Bestandteil der Routinginformation ist)
Authentifizierung	nein	ja
Updates	alle 30s	alle 30s
Metric	max. 15 Hops	max. 15 Hops

Authentifizierung: Der kann überprüfen, ob die Information aus einer vertrauenswürdigen (also ihm bekannten Quelle) stammt!

Abbildung 3.41: RIPv1 kann ein Ziel über maximal 15 Hops erreichen.

3.3.2.2.3 OSPF

OSPF (RFC 1583) ist ein sehr anspruchsvolles dynamisches Routing-Protokoll. Im Gegensatz zu RIP ist OSPF ein Link-State-Protokoll. Link-State bedeutet, etwas vereinfacht ausgedrückt, daß jeder Router die exakte Topologie des Netzwerks kennt. OSPF ist ein quelloffenes Routing Protokoll, das von der IETF standardisiert ist. OSPF ist skalierbar, d.h. es kann an verschiedene Netzwerkgrößen angepaßt werden:

- Kleine Netzwerke: Single Area OSPF
- Große Netzwerke: Multi Area OSPF

OSPF verwendet *cost* als Routing Metrik und unterstützt natürlich VLSM. Mit den OSPF-Nachbarn werden Beziehungen (adjacencies) aufgebaut und alle 10 Sekunden hello-packets (via multicast auf 224.0.0.5) ausgetauscht. Um das Update-Management zu optimieren, wird in jedem Netzwerk ein Designated Router (DR) und ein Backup Designated Router (BDR) gewählt. Nur der DR (und falls der DR ausfällt der BDR) ist für die Verbreitung von Routing Updates zuständig. OSPF verwendet zur Berechnung der besten Routen den shortest path - Algorithmus, des Informatiker Edsger Dijkstra (deshalb oft auch Dijkstra-Algorithmus).

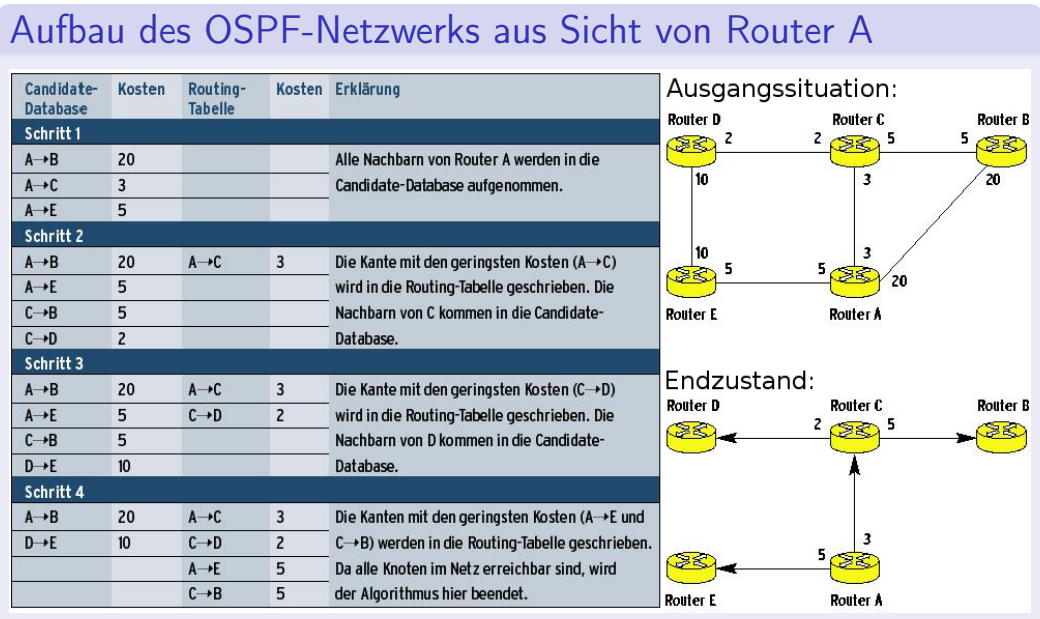


Abbildung 3.42: Erstellen einer OSPF-Routingtabelle für einen Router.

3.3.2.2.4 Vergleich von Distance Vector und Link State

	Distance Vector Routing Protocol	Link State Routing Protocol
Lernweise	Lernt aufgrund der Sichtweise der Nachbarn	Erhält alle Informationen von der Quelle
Updates	periodisch (alle 30 Sekunden bei IP)	Bei Neustart bzw. Änderungen der Netzwerktopologie
Routing-entscheidungen	Übernimmt neue oder bessere Informationen der Nachbarn durch Addition der Vektoren	Erstellt mittels Dijkstra-Algorithmus aus der Datenbank eine Baumstruktur des Netzwerks
Vorteile	Geringe Hardwareanforderungen an den Router und somit günstig. Einfach zu konfigurieren	schnell, schleifenfreie Topologie
Nachteile	langsam und nur für "kleine" Netzwerke geeignet	aufwendige Bedienung, hohe Ansprüche an die Hardware, skalierbar
Wegewahl	kürzester Weg	schnellster Weg

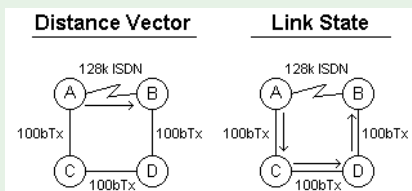
Distance Vector: RIPv1 und RIPv2

Bestimmen den Weg zum Ziel über die Entfernung: **Routing by rumor** (Hop, das einzige Entscheidungskriterium!!!)

Link State: OSPF

Bildet zu Beginn des Routingprozesses ein Abbild der Netzwerktopologie im Speicher ab.

Wirkungsweise

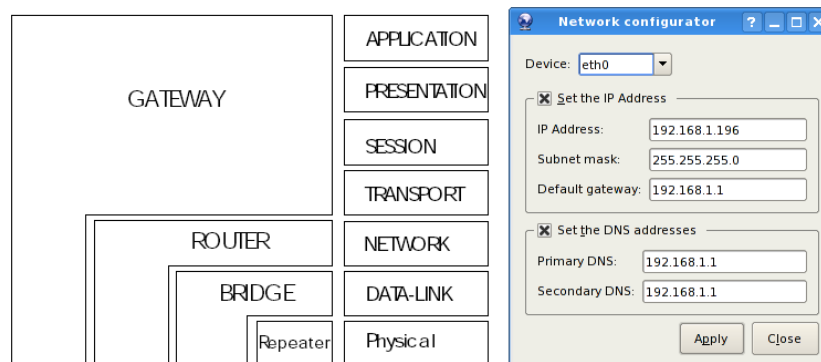


© 1998; InetDaemon Enterprises

Abbildung 3.43: Wirkungsweise von Distance Vector und Link State.

3.3.2.3 Default Gateway

Bei diesem Gerät (manchmal auch Standard Gateway genannt) handelt es sich um einen Router über den das lokale Netzwerk verlassen werden kann. Dieser Router weiß im Regelfall durch seine Routingtabelle über welches seiner Interfaces das Zielnetzwerk zu erreichen ist. Gibt es keinen direkten Eintrag zum Zielnetzwerk und ist das *gateway of last resort*⁸ nicht definiert, dann verwirft der Router das Paket. Der Begriff des Gateways stammt aus einer Zeit, als Router noch völlig unbekannt waren. Mit Gateway wurde, ganz allgemein gesprochen, jedwede Technik verstanden, mit der man Kommunikationskanäle verbinden konnte. Gateway ist heutzutage ein schwammiger Begriff und um eindeutig zu machen welche Art von Gateway gemeint ist, muß man immer die Ebene des OSI-Schichtenmodells angeben, auf welcher dieser Gateway arbeitet. Heutzutage werden Router und Gateway oftmals synonym verwendet. [14, Grundlagen zur Technologie des Internets Teil 2].



Ein Router ist immer Gateway, während ein Gateway nicht immer ein Router sein muß!

Abbildung 3.44: Die Stellung des Gateway im OSI-Schichtenmodell.

⁸z.B. `default via 192.168.1.1 dev eth0 metric 100`

3.4 Layer 4: Transportschicht

Die Transportschicht folgende Aufgaben:

- Aufbau und Wartung einer Verbindung zwischen zwei Endgeräten
- Kann für zuverlässigen (TCP) und nicht zuverlässigen (UDP) Datentransport eingesetzt werden
- Kann per flow control (z.B. Windowing) den Netzwerkdurchsatz steuern
- Kann mehrere Verbindungen gleichzeitig verwalten (multiplexing)

Zuverlässiger Datentransport bedeutet, daß Fehler erkannt und korrigiert werden. Bei nicht zuverlässigem Datentransport wird der Fehler zwar erkannt, die Korrektur jedoch höheren Schichten überlassen. TCP bedeutet, daß eine Verbindung erst nach einem three way handshake zustande kommt. Deshalb nennt man es auch verbindungsorientiert. Wenn große Datenmengen sicher übertragen werden sollen ist TCP erste Wahl. Bei UDP fehlt die einleitende Handshakeprozedur. Man spricht deshalb auch von einer verbindungslosen Übertragung. DNS-Anfragen sind z.B. wegen der minimalen Informationsmenge in UDP realisiert. Der three way handshake würde länger dauern als das Übertragen der DNS-Anfrage und der Antwort.

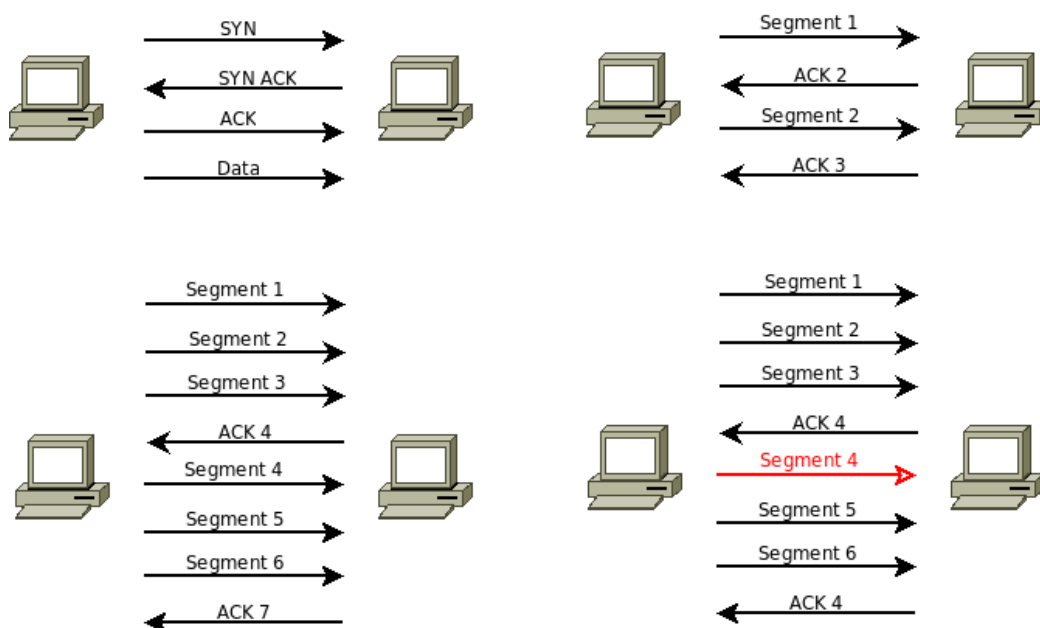


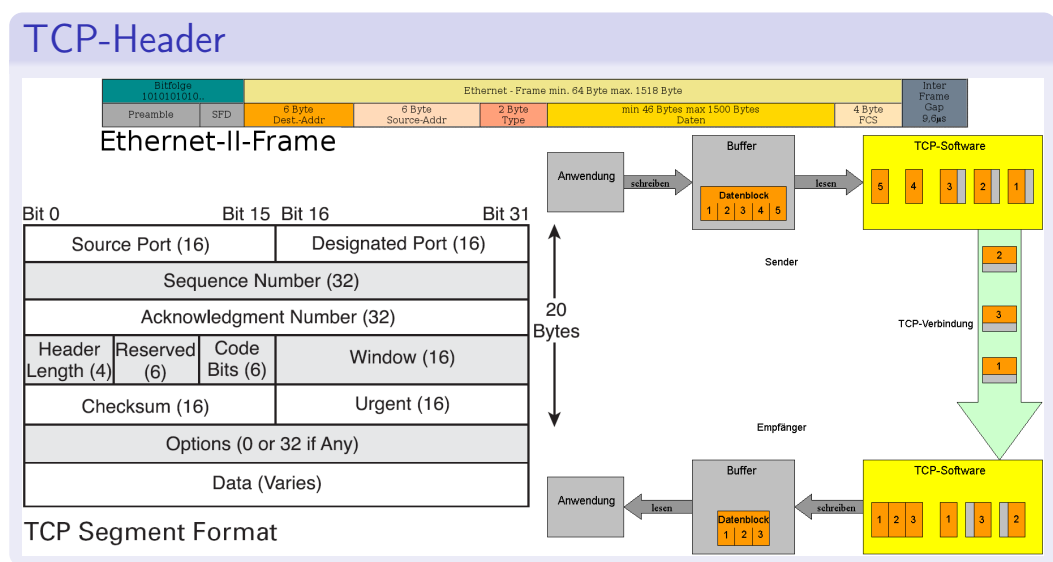
Abbildung 3.45: TCP ist verbindungsorientiert und zuverlässig.

3.4.1 TCP

Allgemeines: TCP ist eine Software-Funktionssammlung, die bei Linux z.B. im Kernel integriert ist. Ein TCP-Segment besteht immer aus zwei Teilen - dem Header und der Nutzlast (Payload). Die Nutzlast enthält die zu übertragenden Daten.

TCP/IP-Paket-Größe: Ein TCP-Segment hat typischerweise eine Größe von 1500 Bytes. Diese Größe wird durch die maximale Rahmengröße eines Ethernet-Frames von 1518 (bzw. 1522) Bytes vorgegeben. Das TCP und IP Protokoll definieren jeweils einen Header von 20 Bytes Größe. Für die Nutzdaten bleiben in einem TCP/IP-Paket also 1460 Bytes übrig⁹.

MTU: DSL-Anschlüsse tauschen ihre Daten über PPPOE (Point-to-Point-over-Ethernet) aus. Dazu müssen für den PPPOE-Rahmen nochmals 8 Bytes reserviert werden. Dem TCP/IP-Paket verbleiben dann nur noch 1492 Bytes. Die Nutzdaten reduzieren sich also auf insgesamt 1452 Bytes. Diese MTU muß man ggf. von Hand nachtragen damit die Datenübertragung fehlerfrei funktioniert.



TCP ist eine Ende-zu-Ende-Verbindung. Sie wird z.B. von Protokollen wie http://, ftp://, smtp:// oder telnet:// verwendet. Das Management dieser Verbindung sowie die Datenübertragung wird von der TCP-Software übernommen.

Im Gegensatz dazu wird auf Layer-2-Ebene eine Punkt-zu-Punkt-Verbindung aufgebaut!

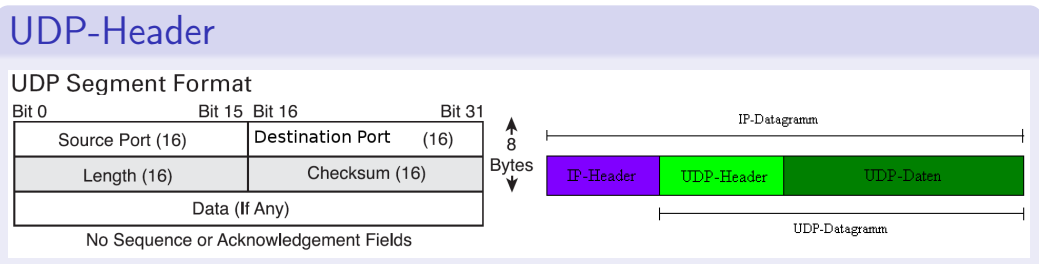
Abbildung 3.46: TCP ist ein verbindungsorientiertes Protokoll.

⁹Bei einer Mindestgröße von 6 Bytes!

3.4.2 UDP

3.4.2.0.1 Allgemein: UDP ist ein *verbindungsloses* Miniatur-Netzwerkprotokoll das in RFC 768 definiert wurde. Es bietet (durch Fehlererkennungsfelder im Header) optional eine Datenintegritätsüberprüfung an.

3.4.2.0.2 Funktionsweise: Verbindungslos bedeutet, daß nicht erst (mittels Handshake wie bei TCP) eine Verbindung zur Gegenstelle aufgebaut wird, sondern daß sofort die Daten losgeschickt werden. Es wird nicht garantiert, daß ein einmal gesendetes Paket ankommt oder daß Pakete in der gleichen Reihenfolge ankommen, mit der sie gesendet wurden: eine Quittierung ist nicht vorgesehen. Die Kommunikationspartner können also nicht feststellen, ob Pakete verloren gegangen sind oder wie lange sie verzögert wurden (Delay). Auch eine Vervielfältigung von Paketen kann auftreten. Eine Anwendung, die UDP nutzt, muß daher gegenüber verloren gegangenen und umsortierten Paketen unempfindlich sein oder selbst entsprechende Korrekturmaßnahmen beinhalten.



Der UDP-Header besteht aus vier Feldern, von denen zwei optional sind. Die Quell- und Ziel-Port Felder sind 16 Bit groß und identifizieren den sendenden und den empfangenden Prozess. Da UDP verbindungslos ist, ist der Quell-Port optional und wird ggf. auf 0 gesetzt. Den Portfeldern folgt das verbindliche Längenfeld, das die Größe der Daten und des Headers des UDP-Datagramms in Oktetten enthält. Der kleinstmögliche Wert sind 8 Oktette (was bedeutet, daß das Datenfeld leer ist). Das letzte Headerfeld ist eine 16 Bit große Prüfsumme über den Teile des IP- und UDP-Headers, und die Daten. Die Prüfsumme ist auch optional, wird aber in der Praxis fast immer benutzt (falls nicht, wird sie ebenfalls auf 0 gesetzt).

Abbildung 3.47: UDP ist ein verbindungsloses Protokoll.

3.4.3 Vergleich TCP/UDP

3.4.3.0.3 Vor- und Nachteile: Da vor Übertragungsbeginn nicht erst eine Verbindung aufgebaut werden muß, können die Hosts schneller mit dem Datenaustausch beginnen. Dies fällt vor allem bei Anwendungen ins Gewicht, die nur kleine Datenmengen ausgetauscht müssen. Ein Beispiel dafür ist das DNS, das auf UDP aufsetzt: Hier werden üblicherweise nur zwei Pakete ausgetauscht, ein Paket mit der Anfrage und eines mit der Antwort. Dazu erst noch per Drei-Wege-Handshake eine Verbindung aufzubauen würde den ganzen Prozess nur unnötig in die Länge ziehen.

Daneben bietet die ungesicherte Übertragung auch den Vorteil von geringeren Verzögerungszeiten: Geht bei einer TCP-Verbindung ein Paket verloren, so muss es erneut angefordert werden. Dies braucht Zeit, die Übertragungsdauer kann daher schwanken, was für Multimediaanwendungen schlecht ist. Bei VoIP würde es dann z.B. zu plötzlichen Aussetzern kommen (ein Lösung dazu wäre die Wiedergabepuffer zu vergrößern, was die Geräte aber teurer machen würde). Solche Anwendungen setzen daher auf UDP, verloren gegangene Pakete bringen hier nicht die gesamte Übertragung ins Stocken (und können im Extremfall in kleinen Mengen auch verloren gehen).

Übersicht über UDP und TCP

	UDP	TCP
Verbindungsaufbau	verbindungslos	verbindungsorientiert
Zuverlässigkeit und Bestätigung	Unzuverlässig, weil "best effort" Zustellung der Daten	Zuverlässig, weil bestätigend
Sendewiederholung von Daten	Nein, muß die Applikation machen	Ja, geschieht automatisch
Steuerung des Datenflusses	gibt es nicht	z.B. Flow Control (sliding windows), Stau-Erkennung
Protokoll-Overhead	sehr gering	etwas höher als bei UDP
Geschwindigkeit	hoch	etwas niedriger als bei UDP
Einsatzbereich	max. einige MB und nur in LAN's	universell

Fazit: In geschalteten Netzwerken und ohne Router zu passieren ist UDP als praktisch gleichwertig gegenüber TCP anzusehen. Aus diesem Grund verwendet auch nfs (das *network file system*) standardmäßig UDP als Transportprotokoll!

Abbildung 3.48: Subnetmasken werden verwendet, um IP-Adressklassen zu trennen.

3.5 Layer 5: Kommunikationssteuerungsschicht

Diese Schicht baut eine Verbindung zwischen zwei Softwareanwendungen auf, pflegt und beendet sie wieder. Wenn auf einem PC mehrere solcher Anwendungen laufen (also z.B. telnet, ftp, Webbrowser usw.), dann muß diese Schicht die eintreffenden Daten den jeweils richtigen Applikationen zuordnen. Dazu arbeitet die Kommunikationssteuerungsschicht sehr eng mit der tieferliegenden Transportschicht zusammen.

Weiterhin findet auf dieser Schicht eine Synchronisation statt: Wenn z.B. bei einem längeren Download die Übertragung (ev. sogar mehrmals) abbricht, dann kann diese Schicht mit Hilfe sogen. Synchronisationspunkte (*checkpoints*) die in den Datenstrom eingebaut wurden, erreichen, daß nicht noch einmal alles, sondern nur noch der fehlende Teil der Daten übertragen wird.

Die Kommunikationssteuerungsschicht und alle höheren Schichten kümmern sich nur noch um anwendungsspezifische Aspekte. Im Gegensatz dazu müssen sich die tieferliegenden Schichten (also 4, 3, 2 und 1) mit dem Transport der Daten beschäftigen.

Beispiel: Network File System (NFS) ist eine Anwendung auf Ebene der Kommunikationssteuerungsschicht.

Beispiel

Ein Web-Browser öffnet eine Webseite. Auf dieser Seite befindet sich ein Java-Applet, ein Macromedia Flash Objekt, Grafiken und Text. Alle diese Objekte sind in eigenen Dateien auf dem Web-Server gespeichert. Damit sie dargestellt werden können müssen sie zuerst heruntergeladen werden. Der Download findet in eigenen Sessions statt. Der Session Layer koordiniert nun welche Pakete zu welcher Datei gehören und sammelt alles im Web-Browser.

Bei heutigen Internetanwendungen sind meist alle drei oberen Schichten (5, 6 und 7) bereits in einer Anwendung zusammengefaßt.

Abbildung 3.49: Daten den richtigen Anwendungen zugeordnen.

3.5.1 Ports und Portnummern

TCP ist für den korrekten Transport der Daten von einem, an ein anderes Endgerät zuständig. Zur Zuordnung der Daten an die richtige Software-Anwendungen werden, ebenfalls auf der Transportebene, Portnummern (16 Bit breite Integerwerte) verwendet. Theoretisch kann ein Host somit bis zu 65.535 verschiedene TCP-Verbindungen aufbauen. Auch UDP verwendet Portnummern zur Adressierung.

Ein gute Übersicht dazu findet man unter

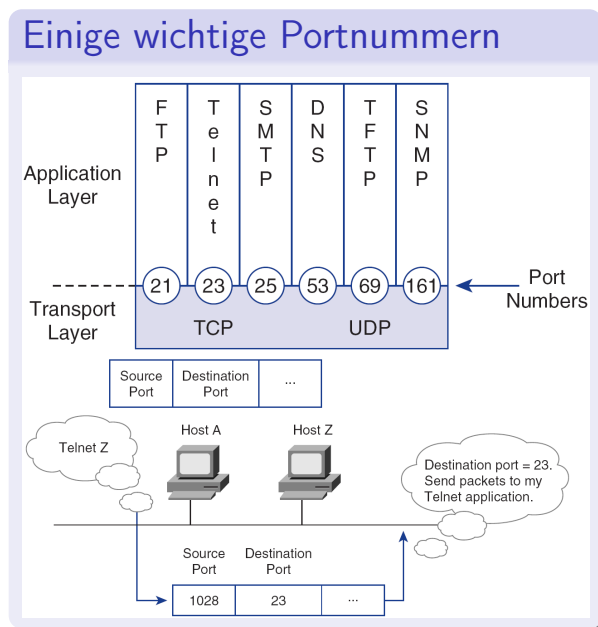
<http://www.iana.org/assignments/port-numbers>

Eine IP-Adresse zusammen mit der Portnummer spezifiziert einen Kommunikationsendpunkt, den sogenannten Socket.

Auf UNIX-Systemen sind Portnummern in der Datei `/etc/services` definiert.

Es gibt drei Bereiche in die die Portnummern eingeteilt werden:

- Well Known Ports: 0-1023 werden/wurden von der IANA vergeben
- Registered Ports: 1024-49151 werden von Applikationen benutzt
- Dynamic/Private ports: 49152-65535 sind frei verwendbar



In diesem Beispiel öffnet Host A eine telnet Verbindung zu Host Z. Host A wählt eine beliebige Quell- portnummer > 1023 aus und nimmt als Zielport 23. Öffnet Host A eine zweite telnet-Verbindung, weist er dieser eine andere Quellportnummer zu.

Weitere wichtige Portnummern sind 80 für http und 22 für ssh!

Well Known Ports: 0-1023

Registered Ports: 1024-49151

Dynamic/Private ports: 49152-65535

Abbildung 3.50: Über Ports finden Daten den Weg zur richtigen Applikation.

3.6 Layer 6: Darstellungsschicht

Auf dieser Schicht wird sichergestellt, daß die Daten von der Anwendungsschicht des einen Systems auch von der Anwendungsschicht eines anderen Systems wiedergegeben werden kann. Auch Aspekte der Datenkompression bzw. der Einsatz kryptographischer Methoden können Aufgabe der Schicht 6 sein.

Beispiele

- Grafik
 - TIFF, JPEG GIF, BMP, PNG und viele andere
- Video
 - MPEG, AVI, WMV
- Sound
 - WAV, MIDI
- Text
 - ASCII (heute Standard)
 - EBCDIC (IBM, wird in Mainframes zur Zeichendarstellung verwendet)

Abbildung 3.51: Daten für die Darstellung aufbereiten.

3.7 Layer 7: Anwendungsschicht

“Normale” Anwender haben es im Regelfall nur mit dieser Schicht zu tun. Sie stellt die Schnittstelle Mensch/Software dar und unterscheidet sich von tieferliegenden Schichten dadurch, daß sie diesen keine Dienste zur Verfügung stellt.

Beispiele für Anwendungen auf dieser Schicht sind z.B. (heutzutage auch) Textverarbeitungs- und Tabellenkalkulationsprogramme, vor allem aber Webbrowser, E-Mail- und Chat-Programme und alles was sonst noch irgendwie interaktiv über das Netz kommunizieren kann.

Beispiele

- ftp
- telnet
- E-Mail (smtp)
- Web-Browser (http)

Abbildung 3.52: Die Ebene der Benutzerprogramme.

3.7.1 Network File System - NFS

NFS ist ein von Sun Microsystems entwickeltes Protokoll mit dem es möglich ist, Daten über ein Netzwerk auszutauschen. Es ist vor allem in der Unix-Welt weit verbreitetes. Man kann mit ihm bequem auf Verzeichnisse entfernter Rechner zugreifen und sie in das eigene Dateisystem einbinden.

In RFC 1813 wird das zur Zeit weit verbreitete NFSv3 dokumentiert. Mit dieser NFS-Version ist es zum ersten Mal möglich, NFS wahlweise auch über TCP zu benutzen. NFSv2 hingegen wurde komplett über UDP (Port 2049) abgewickelt. Da NFS eigentlich für lokale Ethernet-Netzwerke entwickelt wurde stellte UDP, bei kollisionsfreien Netzwerken auch kein Problem dar. Soll der Datentransport allerdings über Router erfolgen oder arbeitet man in einem nicht geschwächten Netzwerk bzw. in einem Netzwerk das häufig stark ausgelastet ist, dann sollte man besser auf TCP zurückgreifen weil diese Protokoll automatisch dafür sorgt, fehlerhafte oder fehlende Pakete erneut zuzustellen. Heutzutage wird üblicherweise NFS over TCP eingesetzt.

Der neue Standard NFSv4 bringt auch endlich eine Benutzerauthentifizierung mit. Bei NFSv3 wird lediglich der Name des Client-Rechners überprüft.

```
Die Datei /etc/exports des NFS-Servers 192.168.1.145:

# /etc/exports: the access control list for filesystems which may be exported
#               to NFS clients.  See exports(5).
#
# Example for NFSv2 and NFSv3:
# /srv/homes      hostname1(rw,sync) hostname2(ro,sync)
#
# Example for NFSv4:
# /srv/nfs4       gss/krb5i(rw,sync,fsid=0,crossmnt)
# /srv/nfs4/homes gss/krb5i(rw,sync)
#
/tmp 192.168.59.129(rw,no_root_squash,sync)
/home/herzer/2006/testing 192.168.59.129(rw,no_root_squash,sync)

Ausschnitt der Datei /etc/fstab des NFS-Clients 192.168.59.145:

192.168.59.132:/home/herzer /home/herzer/AP059132 nfs defaults,auto 0 0
192.168.59.158:/home/ima /home/herzer/2006/IMA/ima nfs defaults,rw,auto 0 0

Auszugsweise Wiedergabe des df Kommandos auf dem NFS-Client 192.168.59.145:

herzer@herzer-desktop:~$ df
Filesystem            1K-blocks      Used Available Use% Mounted on
/dev/sdc1              73939452    27036644  43146816   39% /
...
192.168.59.132:/home/herzer
38101024    23238784    12926816    65% /home/herzer/AP059132
192.168.59.158:/home/ima
5075840     1284288     3533696    27% /home/herzer/2006/IMA/ima
/dev/sdd1     150877984  40635676  102578100   29% /media/disk
/dev/sda6      76390460   67297308   5212656    93% /home/herzer/media
```

NFS im OSI-Schichtenmodell	
Anwendung	NFS
Darstellung	XDR
Sitzung	(Sun-) RPC
Transport	(UDP) TCP
Netzwerk	IP
Netzzugang	Ethernet Token Ring FDDI ...

Abbildung 3.53: Anwendung des NFS-Protokolls.

3.7.2 Server Message Block - SMB (Samba)

SMB ist ein von Microsoft und IBM entwickeltes Protokoll mit dem man unter Windows Datei- und Druckerfreigaben realisieren kann. Linux kann diese Freigaben mit Samba entweder als Client nutzen oder sogar selbst als SMB-Server konfiguriert werden. Samba ist eine freie SMB-Implementierung.

Microsoft bezeichnet das SMB-Protokoll seit einiger Zeit auch als CIFS (Common Internet File System), was etwas moderner klingt. Die von Microsoft zur Verfügung gestellte Dokumentation ist allerdings mager und unvollständig - es gibt keinen offiziellen Standard. Die Samba-Entwickler mußten ihre Protokollkenntnisse teilweise durch Analysen der Kommunikation ergänzen¹⁰.

Traditionell nutzt SMB den

- UDP-Port 137 (NetBIOS name service oder netbios-ns)
- UDP-Port 138 (NetBIOS datagram service oder netbios-dgm)
- TCP-Port 139 (NetBIOS session service oder netbios-ssn)

Jeder, der mit einem guten Netzwerk-Packet-Analyzer sein Netzwerk beobachtet, wird über die Menge an Verkehr amüsiert sein, der ausgelöst wird, wenn man nur eine einzige Datei öffnet.

Quelle: <http://gertranssmb3.berlios.de/output/IntroSMB.html>

Group	#	Layer Name	Key Responsibilities	Data Type Handled	Scope	Common Protocols and Technologies
Lower Layers	1	Physical	Encoding and Signaling; Physical Data Transmission; Hardware Specifications; Topology and Design	Bits	Electrical or light signals sent between local devices	(Physical layers of most of the technologies listed for the data link layer)
	2	Data Link	Logical Link Control; Media Access Control; Data Framing; Addressing; Error Detection and Handling; Defining Requirements of Physical Layer	Frames	Low-level data messages between local devices	IEEE 802.2 LLC, Ethernet Family; Token Ring; FDDI and CDDI; IEEE 802.11 (WLAN, Wi-Fi); HomePNA; HomeRF; ATM; SLIP and PPP
	3	Network	Logical Addressing; Routing; Datagram Encapsulation; Fragmentation and Reassembly; Error Handling and Diagnostics	Datagrams / Packets	Messages between local or remote devices	IP; IPv6; IP NAT; IPsec; Mobile IP; ICMP; IPX; DLC; PLP; Routing protocols such as RIP and BGP
	4	Transport	Process-Level Addressing; Multiplexing/Demultiplexing; Connections; Segmentation and Reassembly; Acknowledgments and Retransmissions; Flow Control	Datagrams / Segments	Communication between software processes	TCP and UDP; SPX; NetBEUI/NBF
Upper Layers	5	Session	Session Establishment, Management and Termination	Sessions	Sessions between local or remote devices	NetBIOS, Sockets, Named Pipes, RPC
	6	Presentation	Data Translation; Compression and Encryption	Encoded User Data	Application data representations	SSL; Shells and Redirectors; MIME
	7	Application	User Application Services	User Data	Application data	DNS; NFS; BOOTP; DHCP; SNMP; RMON; FTP; TFTP; SMTP; POP3; IMAP; NNTP; HTTP; Telnet

http://www.tcpipguide.com/free/t_OSIRferenceModelLayerSummary.htm

Abbildung 3.54: Einordnung des SMB-Protokolls und seiner Komponenten Netbios bzw. NetBEUI.

¹⁰<http://gertranssmb3.berlios.de/output/IntroSMB.html>

Kapitel 4

Netzwerkanwendungen

4.1 VLAN

Definition: Ein VLAN ist ein virtuelles Netzwerk in einem physikalischen Netzwerk. Switches können, je nach hard- bzw. softwaretechnischer Ausbaustufe mehrere VLANs parallel betreiben. Ein weit verbreiteter VLAN-Standard ist **Frame Tagging** nach IEEE 802.1q.

Vorteile: VLAN's ermöglichen eine logische Zuordnung, von Switch-Ports zu z.B. Funktion im Betrieb. VLAN-fähige Switches sind meist in größerer Anzahl installiert.

Beispiel: Bevor es VLAN-fähige Switches gab, waren z.B. alle User auf der dritten Etage mit einem Switch verbunden (also in der gleichen Broadcast-Domäne). Man konnte die Benutzer zwar durch eigene Subnetze logisch voneinander trennen, Broadcast-Anfragen erreichten aber alle am Switch angeschlossenen User. Die Alternative wäre gewesen, auf jeder Etage drei Switches (mit entsprechend hohem Verkabelungsaufwand) zu betreiben.

Mit VLAN's ist es möglich, logische Benutzergruppen (also kleine Broadcast-Domänen **innerhalb des Switches** zu bilden. Durch diese Isolierung des Netzwerkverkehrs, wird zum einen die Performance, zum anderen auch die Sicherheit des Netzwerks erhöht. Um Daten untereinander auszutauschen benötigen VLANs einen Router oder einen Layer-3-Switch.

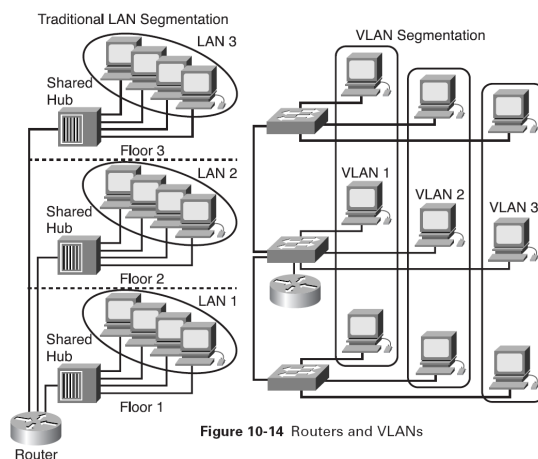


Figure 10-14 Routers and VLANs

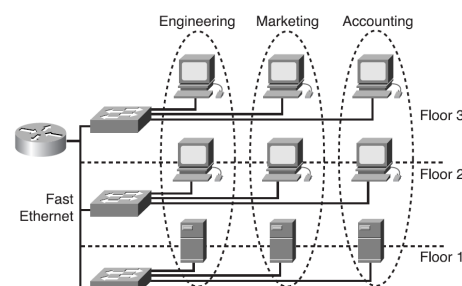


Figure 10-15 Removing Physical Boundaries

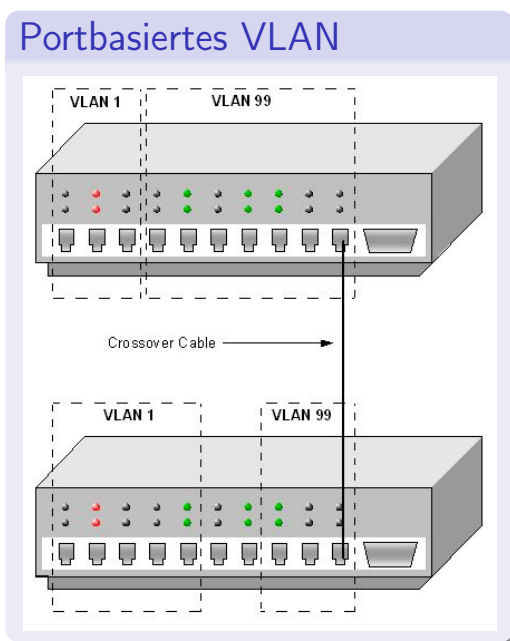
Will ein Host in einer Broadcast-Domäne mit einem Host in einer anderen Broadcast-Domäne eine Verbindung aufnehmen, muß ein Router verwendet werden (Fig. 10-14). Dasselbe gilt auch für VLAN's.

Abbildung 4.1: VLAN's und traditionelle Verkabelung

4.1.1 Funktionsweise

VLANs können auf verschiedenen Schichten des OSI-Modells realisiert werden:

- Schicht 1: Die einfachste Art der Zuordnung ist die feste Definition einer VLAN ID auf einen bestimmten Port des Switches. Dadurch entstehen statische VLAN's, mit denen man z.B. die Anzahl der PC's auf einer Etage kontrollieren kann.
- Schicht 2: VLAN's können auch auf der Sicherungsschicht angelegt werden. Diese dynamischen VLAN's ordnen MAC-Adressen (anhand einer Datenbank im Switch) VLAN's zu. Taucht die MAC-Adresse eines Rechners in einem anderen Switch auf, wird er dem richtigen VLAN zugeordnet.
- VLAN's sind auch auf Schicht 3 und Schicht 4 realisierbar
- nur auf managbaren Switches können VLAN's eingerichtet werden



VLAN's können z.B. über die Management-Konsole oder das Web-Interface eingerichtet werden. Die Anzahl der Ports, die einem VLAN zugeordnet werden, kann dabei von Switch zu Switch variieren. In diesem Beispiel sind beide VLAN's über ein Netzkabel miteinander verbunden und können Daten austauschen. Ein Datenaustausch zwischen VLAN 1 und VLAN 99 ist nicht möglich. Statische VLAN's sind einfach zu konfigurieren, aber unflexibel und können u.U. leicht missbraucht werden. Dynamische VLAN's sind aufwendiger zu konfigurieren, bieten aber eine höhere Sicherheit.

Abbildung 4.2: Portbasierte VLAN's.

4.1.2 VLAN Tagging

Portbasiertes VLAN stößt bei mehreren VLANs auf einem Gerät schnell an seine Grenzen. Der Verkabelungsaufwand wird sehr hoch und immer mehr Ports gehen an Uplinks verloren.

Eine elegante Lösung bietet hier das **Frame Tagging** nach 802.1q. Dabei wird jeder Frame mit einer VLAN-ID markiert. Damit eindeutig gekennzeichnete Frames können jetzt über eine, für alle VLANs gemeinsame(!), Datenleitung zwischen Switches ausgetauscht werden.

Diese gemeinsame Datenleitung wird häufig als Trunk bezeichnet, der entsprechende Port heißt (wegen seiner besonderen Funktion im Gegensatz zum Switched Port) **Trunked Port**.

Frames die auf einem Trunked Port ankommen, werden anhand der VLAN-ID an den entsprechenden VLAN-Port weitergeleitet. Ports an denen sich Endgeräte befinden, löschen vor der Weiterleitung des Frames an das Endgerät den Tag. Umgekehrt wird beim Empfangen eines Frames von einem Endgerät auch das entsprechende VLAN-Tag vom Switch hinzugefügt, falls der Frame den Switch über den Trunked Port verlassen muß.

Achtung: Wenn ein Switch als VLAN-fähig angeboten wird, bedeutet das nicht automatisch, daß er auch Frame Tagging nach IEEE 802.1q beherrscht (und damit auch trunking kann).

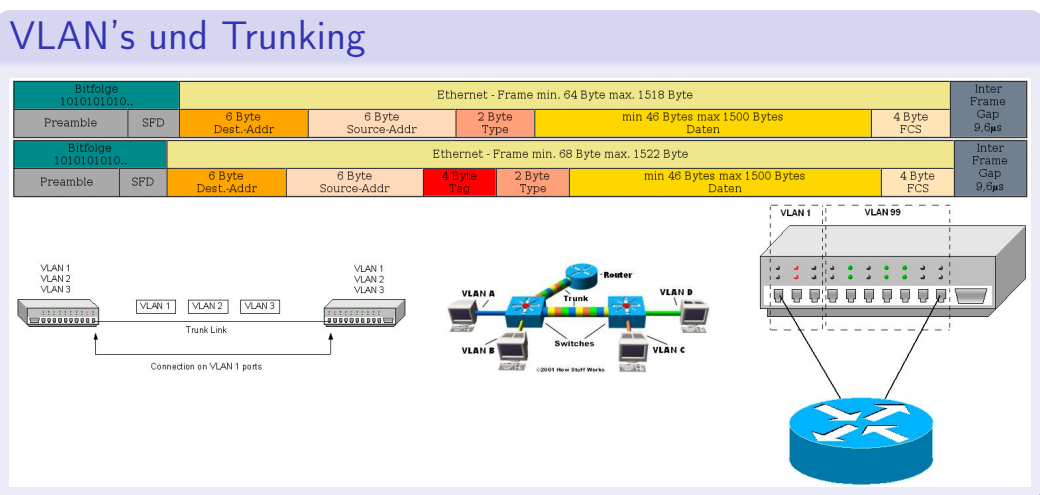


Abbildung 4.3: Trunking erlaubt ein wesentlich flexibleres VLAN-Management.

4.2 Firewall

Eine Firewall ist ein Netzwerkgerät, das zwischen einem LAN und dem WAN installiert ist. Sie soll das LAN vor Angriffen aus dem WAN schützen. Erreicht wird dieser Schutz durch die Analyse des Netzwerkverkehrs bis auf Layer 7 des OSI-Schichtenmodells. Gegen Angriffe aus dem LAN ist eine solche WAN-Firewall machtlos, weil die Filterung nur in einer Richtung stattfindet.

Software-Firewalls: Auch Personal Firewall genannt. Sie sollen den Rechner auf dem sie installiert sind schützen. Das gelingt in den meist Fällen mehr schlecht als recht.

Hardware-Firewalls: Geräte mit mehreren Netzwerkschnittstellen. Sicherheitstechnisch der Software-Firewall weit überlegen.

Figure 20-15 Firewall Architecture

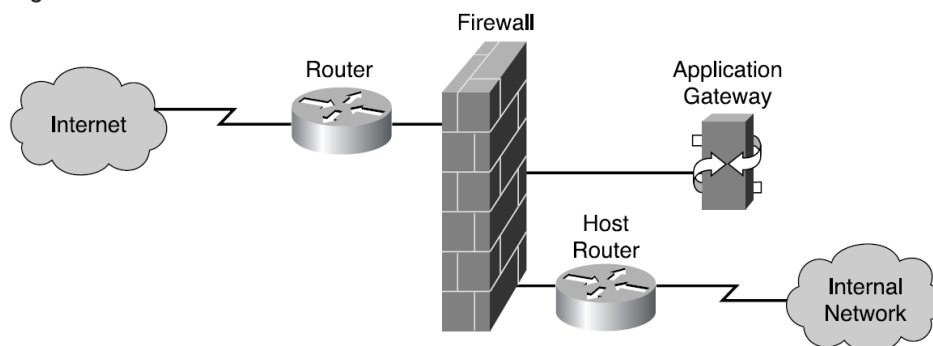


Abbildung 4.4: Firewalls erhöhen die Sicherheit von LAN's gegenüber WAN's.

4.3 DMZ

Die Demilitarisierte Zone ist ein besonderer Netzwerkbereich. Wie sein militärischer Name schon andeutet, handelt es sich dabei um eine Art Niemandsland. In der demilitarisierten Zone haben externe Nutzer (also solche die über das WAN kommunizieren) und lokale Nutzer (aus dem LAN) eine Plattform auf der sie relativ frei (aber trotzdem kontrolliert) bestimmte Ressourcen gemeinsam nutzen können. In der DMZ werden deshalb häufig Dienste wie ftp-, Mail oder auch ssh-Server angeboten.

Host DMZ: Auch unechte DMZ genannt. Alle externen Anfragen werden an eine IP-Adresse im LAN weitergeleitet. Diese Funktion findet man oft bei Billig-Routern. Der Schutz des LANs ist hier aber ähnlich schwach wie bei einer Software-Firewall.

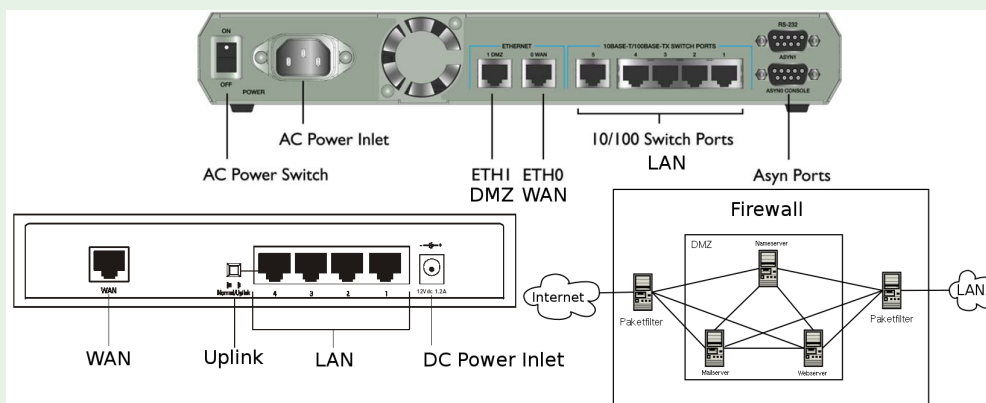
DMZ: Bei der echten DMZ hat man ein Netzwerkgerät mit zwei LAN-Schnittstellen (im wesentlichen also einen Router). Aus diesem Grund kann man sich auch aus zwei Routern die jeweils eine LAN-Schnittstelle haben ebenfalls eine DMZ aufbauen.

Ein weiteres Thema sind LAN-Angriffe von innen, also den lokalen LAN-Usern. Durch VLAN-Segmentierung und Verwendung von geschwichten Netzwerken sind solche Probleme aber meist relativ leicht in den Griff zu kriegen.

Demilitarisierte Zone:

Aufbau einer DMZ

Die DMZ ist ein isolierter Bereich, der sich zwischen zwei IP-Netzen befindet. Man kann nun den Zugriff auf öffentliche Dienste (z.B. WWW-Server) innerhalb der DMZ anbieten und gleichzeitig das LAN vor unberechtigten Zugriffen aus dem WAN schützen.



Zwei Router: Der eine mit HOST DMZ, der andere mit echter DMZ.

Abbildung 4.5: DMZs sind besonders gesicherte Netzwerkbereiche.

4.4 DHCP

Durch DHCP ist die Einbindung eines neuen Computers in ein bestehendes Netzwerk ohne weitere Konfiguration möglich. Lediglich der automatische Bezug der IP-Adresse muß in der Netzwerkeinstellung des Client eingestellt sein. Ohne DHCP muß man den Netzwerkzugang "von Hand" einrichten. Dazu müssen Angaben über IP-Adresse, Netzmaske, Gateway und DNS-Server vorhanden sein. Ein DHCP-Server vergibt diese Parameter beim Starten eines neuen Rechners (DHCP-Client) im Netzwerk automatisch.

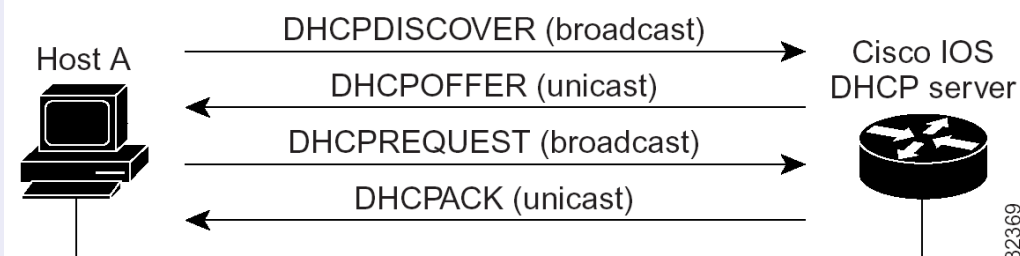
DHCP wurde im Hinblick auf zwei Einsatzszenarien entwickelt:

1. Große Netzwerke mit häufig wechselnder Topologie
2. Anwender, die "einfach nur eine Netzwerkverbindung" haben wollen und nicht mit Netzwerkkonfigurationen belastet werden sollen

Der DHCP-Server ist in der Regel als Dämon (z. B. `dhcpcd`) implementiert und benutzt UDP-Port 67 für Client-Anfragen. In der Konfigurationsdatei des DHCP-Servers befinden sich u.a. Informationen über den zu vergebenden Adresspool sowie Angaben über die Subnetzmaske, die lokale DNS-Domäne und den Default Gateway.

DHCP - Funktionsweise

Figure 14 *DHCP Request for an IP Address from a DHCP Server*



Ein DHCP-Client kann u.U. auch von mehreren DHCP-Servern einen Vorschlag erhalten. Der Client wählt dann im Regelfall das Angebot aus, welches zuerst bei ihm eintraf.

► Praktisches Beispiel

Abbildung 4.6: DHCP läuft in mehreren Stufen ab.

4.5 PAT

Mehrere Rechner in einem lokalen Netzwerk sollen sich einen gemeinsamen Internetzugang teilen! Häufig stehen aber nicht so viele öffentliche IP-Adressen zur Verfügung, wie Hosts, die ins Netz wollen.

PAT (RFC 3022, oft auch Masquerading genannt) ist eine Technik mit der genau dieses Problem gelöst werden kann. Bei dieser spezielle Form von NAT wird sowohl die IP-Adresse als auch die Port-Nummer jedes einzelnen Pakets umgeschrieben. Masquerading hat die Eigenschaft, daß auf alle Rechner des LANs nicht von außen zugegriffen werden kann. Das hat Vorteile (Angreifer wissen nicht wie das Netz hinter einem PAT-Router aussieht) und Nachteile (Rechner im LAN sind nicht mehr direkt, also per IP-Adresse, von außen sichtbar). Um diesem Nachteil Abhilfe zu verschaffen, wurde Port Forwarding entwickelt. Damit kann z.B. ein ssh-Zugriff ganz gezielt an einen Rechner im LAN weitergeleitet werden, während ein WEB-Zugriff an einen anderen Rechner geht.

Im täglichen Leben wird PAT oft fälschlicherweise als NAT beschrieben. NAT ist eine ähnliche Technik wie PAT, es wird aber nur die IP-Adresse umgeschrieben!

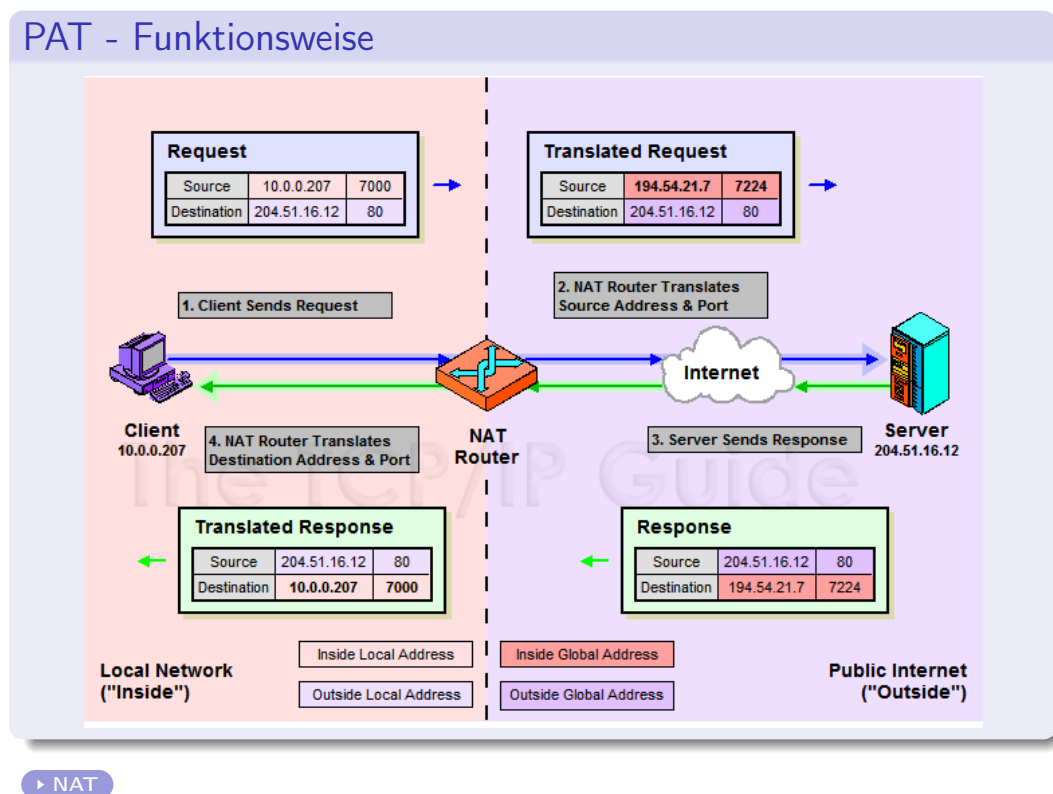


Abbildung 4.7: DSL-Router machen PAT. In den Manuals wird es aber meist NAT genannt.

4.6 MD5

Ein weiterer Sicherheitsaspekt in Netzwerken ist Authentizität:

1. Sind Daten, die ich erhalten habe auch unverfälscht vom Sender bei mir angekommen?
Verfälschungen können z.B. durch fehlerhafte Übertragung oder absichtliche Manipulation entstehen.
2. Ist der Sender der Daten wirklich der, für den er sich ausgibt?

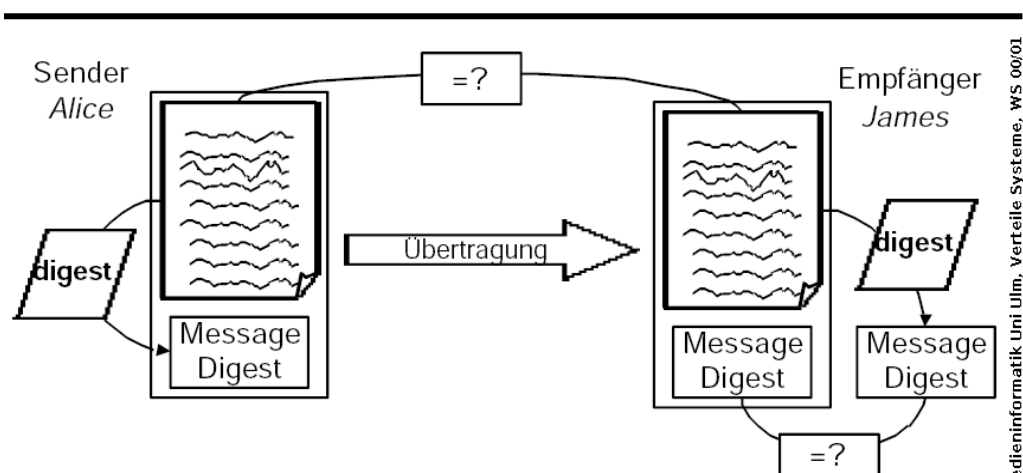
Die erste Frage kann man mit dem Verfahren *Message Digest*, ein von Ronald L. Rivest entwickeltes Hashverfahren (RFC 1321, MD5), beantworten.

Hash-Verfahren (dt. gestreute Speicherung) erzeugen einen sogn. Hashwert (Prüfsumme). Dabei wird eine umfangreiche Quellmenge auf eine wesentlich kleinere Zielmenge abgebildet, welche immer eine bestimmte Länge hat. Im Fall von MD5 wird eine beliebig lange Eingabedatei auf 128 Bit abgebildet. Der dazu verwendete Algorithmus arbeitet so eindeutig, daß es unmöglich ist, zwei verschiedene Eingabedateien zu erzeugen, die den gleichen Hashwert haben. Und noch besser: Selbst ein einziges verändertes Bit produziert einen völlig anderen Hashwert. Damit ist eine Veränderung in der Datei sofort eindeutig auf den ersten Blick zu erkennen!

MD5

72 / 110

Ist die Nachricht unverändert ?



- Sender berechnet den Message Digest vor Übertragung
- Empfänger berechnet den Message Digest nach Übertragung
- Wenn beide Digests gleich sind, wurde die Nachricht nicht verändert
- Aber! - ist Alice wirklich Alice?

Bei Hash-Verfahren handelt es sich also nicht direkt um Verschlüsselung wie z.B. bei Blowfish oder AES. Vielmehr sind Hashwerte mit Fingerabdrücken zu vergleichen. Auf diese Weise lassen sich z.B. Dokumente signieren. Eine weitere sehr bekannte Prüfsummenteknik ist CRC.

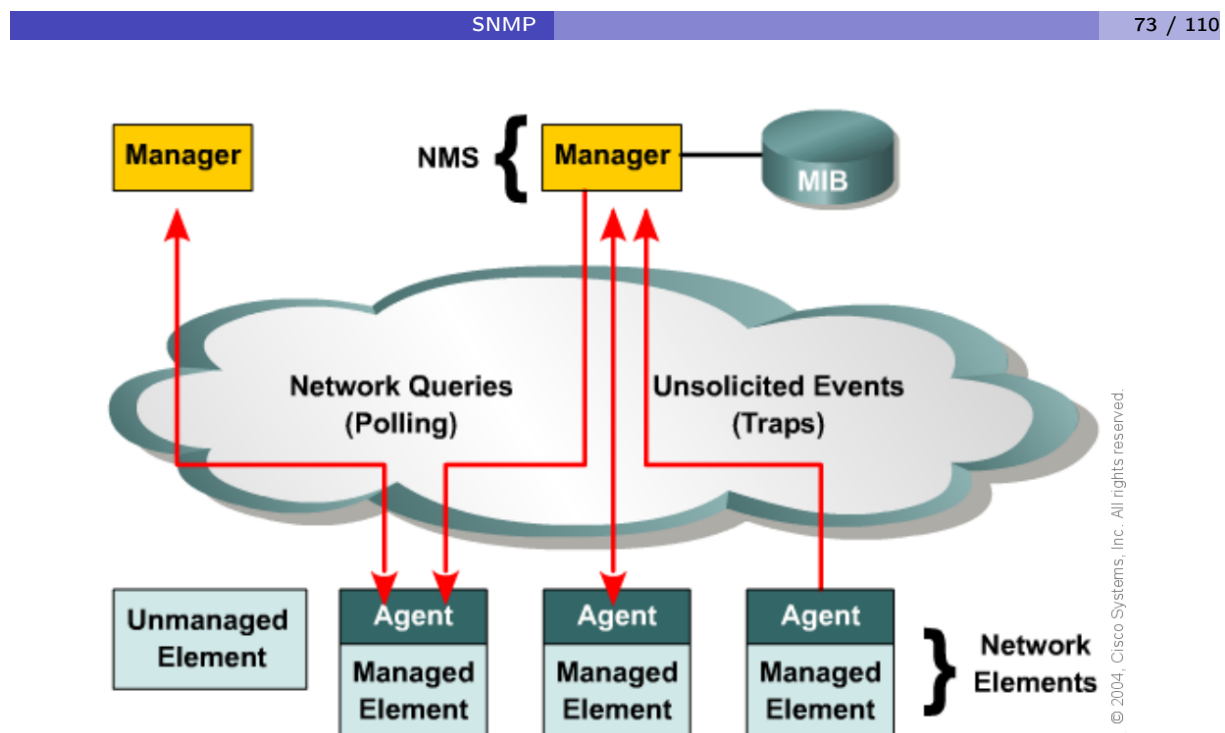
Prof. M. Weber, Medieninformatik Uni Ulm, Verteilte Systeme, WS 00/01

Abbildung 4.8: Prinzipielle Anwendung der md5-Summe.

4.7 SNMP

Je größer ein Netzwerk wird, umso schwieriger wird es seine Funktion zu überwachen. Aus diesem Grund wurde bereits 1989 mit SNMP ein Protokollstandard eingeführt, mit dem es möglich ist, Informationen in einem Netzwerk zu sammeln. SNMP ist ein allgemein anerkannter Standard der IETF und wird von allen führenden Netzwerkausrüstern in den Netzwerkgeräten implementiert. 1991 wurde mit SNMPv2c eine erweiterte Spezifikation der Öffentlichkeit vorgestellt. Die zur Zeit letzte Überarbeitung des SNMP-Standards, SNMPv3, u.a. einen Verschlüsselungsmechanismus mit, der es Datenbanken jetzt erlaubt SNMP-Pakete nur nach erfolgreicher Authentifizierung zu akzeptieren. SNMP arbeitet auf den UDP-Ports 161/162 und kennt folgende Kommandoarten:

- **Get:** Damit kann ganz gezielt ein Parameter eines Netzwerkgerätes abgefragt werden (z.B. wieviele Bytes auf einem Interface empfangen oder versendet wurden)
- **Set:** Damit kann ganz gezielt ein Parameter im Netzwerkgerät verändert werden (z.B. Zähler löschen)
- **Trap:** Vom Netzwerkgerät selbstständig versendete Meldungen (z.B. Alarmer o.ä.)



Die Toolsammlung net-snmp ermöglicht es, die einzelnen von SNMP-fähigen Netzwerkgeräten, anzusprechen. [► Objekte](#)

Abbildung 4.9: Agents (z.B. Router oder Switches), stellen Managern (meist Server-Rechnern) Informationen zur Verfügung.

4.8 Kryptographische Verfahren

4.8.1 Einleitung

Bei der Verschlüsselung geht es darum, Nachrichten so zu transportieren, daß es außer dem Sender und dem Empfänger niemandem möglich ist, die Nachricht zu lesen. Neben der "klassischen" Variante (symmetrische Verschlüsselung), gibt es seit Ende der 70er Jahre auch zwei weitere Varianten: die asymmetrische Verschlüsselung und die hybride Verschlüsselung.

Zur Verschlüsselung benötigt man in jedem Fall eine Verschlüsselungsvorschrift, den sogn. Algorithmus. Diese wird durch einen Schlüssel (eine Folge von z.B. zufällig generierten Zeichen) ergänzt. Sowohl von der Qualität des Algorithmus, als auch von der Anzahl der Möglichkeiten den Schlüssel zu bilden (also der Schlüssellänge) hängt die Güte ("Knackbarkeit") der verschlüsselten Daten ab.

Historisches

Schon Julius Cäsar (100-44 v. Chr.) soll angeblich bei der Kommunikation auf Feldzügen Verschlüsselung eingesetzt haben. Seine Methode war folgende: Aus dem Text wurden alle Leerstellen und Satzzeichen entfernt und alle Buchstaben des Textes in Großbuchstaben umgewandelt. Danach wurde im Alphabet (A, ..., Z) jedes Zeichen zyklisch um einen festen Wert verschoben, also zum Beispiel A auf B, B auf C ... und Z auf A. Alle Zeichen des Textes wurden auf diese Weise ersetzt.



Beispiel zu Cäsar-Code

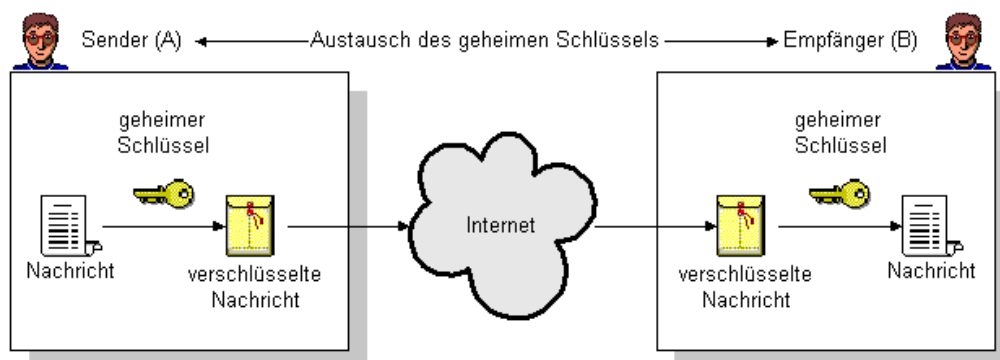
A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	Klar
D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	A	B	C	Geheim

Nach dieser Verschlüsselung wird aus **JULIUS CAESAR** dann **MXOLXV FDHVDU** !

Abbildung 4.10: Einfache Verschlüsselungstechniken sind schon seit langer Zeit bekannt.

4.8.2 Symmetrische Verschlüsselung

Als symmetrische Verschlüsselung wird grundsätzlich jedes Verfahren bezeichnet, das zum Ver- und Entschlüsseln den gleichen Schlüssel benutzt. Die einfachste und zugleich sicherste Methode ist ein Wegwerfsschlüssel der z.B. mit einem Zufallsgenerator erzeugt werden kann. Dieser Schlüssel wird dann mit der Botschaft "verrechnet". Der Empfänger muß dann mit dem gleichen Schlüssel "zurückrechnen". Dabei gibt es aber folgendes Problem: Der Empfänger muß ja irgendwie in den Besitz des Schlüssels kommen (aber wie kann er ihn "sicher" erhalten)? Zur geheimen Aufbewahrung von Informationen ist dieses Verfahren aber sehr gut geeignet.



Beispiele für bekannte symmetrischer Verschlüsselungsalgorithmen sind der Data Encryption Standard (DES), welcher von IBM Anfang der siebziger Jahre entwickelt wurde und mit einer
► Schlüssellänge von 56 Bit arbeitet. Später wurde er von 3DES abgelöst. Heute wird vorzugsweise AES eingesetzt!

Abbildung 4.11: Funktionsweise der symmetrischen Verschlüsselung.

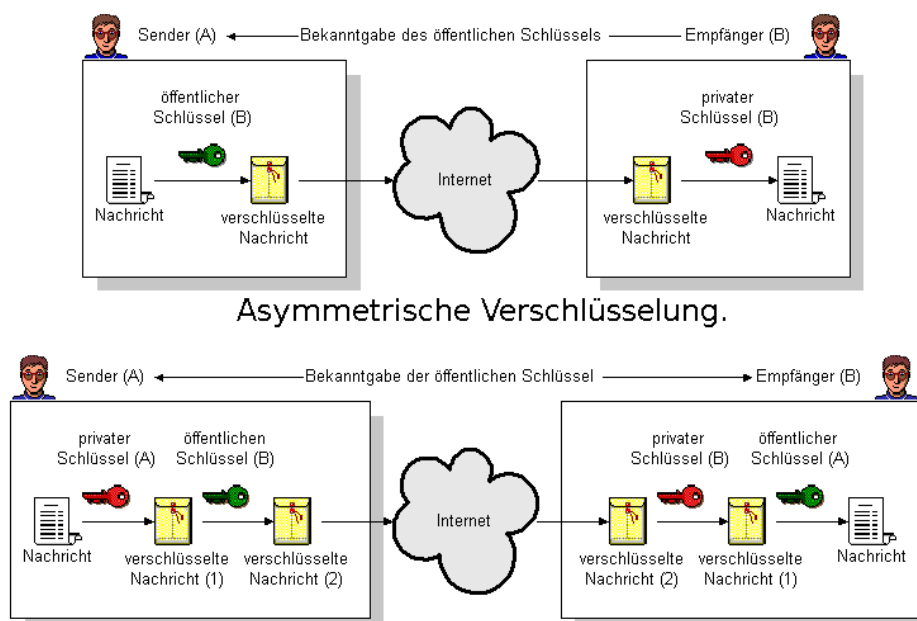
4.8.3 Asymmetrische Verschlüsselung

Zwei Mathematiker, Diffie und Hellmann, entwickelten 1976 ein Verfahren, bei dem der Schlüssel aus zwei Teilen besteht: Alles was mit einem öffentlichen Schlüssel (*public key*) kodiert wird, kann nur mit dem privaten Schlüssel (*private key*) wieder gelesen werden. Das DHA-Protokoll diente lediglich dem Schlüsselaustausch über einen unsicheren Kommunikationskanal und war als U.S. Patent 4,200,770 bis zum 06.09.1997 geschützt. Sein Schwäche war (wegen fehlender Authentifizierung der Kommunikationspartner) ein möglicher Man-in-the-Middle Angriff.

Der ElGamal-Algorithmus kann Verschlüsselungs- und Signaturverfahren sein! Er ist eine Abänderung des DHA-Protokolls und gegen den Man-in-the-Middle Angriff immun. Taher ElGamal veröffentlichte sein Verfahren 1985 ohne jegliche Patentansprüche!

Ein weitere Vertreter dieser Algorithmen ist RSA, benannt nach seinen Entwicklern Ron Rivest, Adi Shamir, Leonard Adleman. Er wurde 1977 in den USA als U.S. Patent 4,405,829 entwickelt und bis September 2000 patentiert.

Die Basis aller Verfahren sind mathematische Berechnungen mit großen Primzahlen die mit den o.g. Algorithmen sehr großen Zahlen (1000 Bit, also z.B. 2^{1024} Stellen) ergeben, Solche Berechnungen lassen sich nur in einer Richtung schnell und effektiv durchführen.



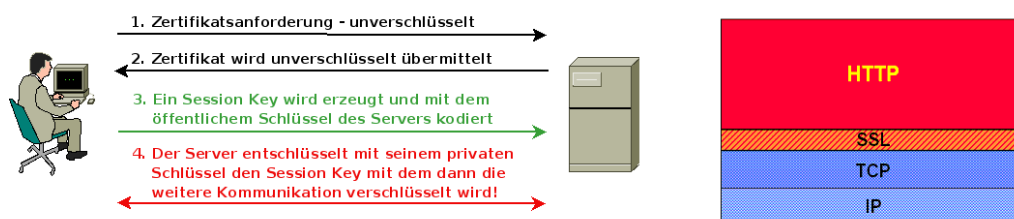
Die asymmetrische Verschlüsselung kann auch genutzt werden um das Problem der Authentifizierung zu lösen. Zu diesem Zweck werden die öffentlichen Schlüssel von Sender und Empfänger gegenseitig bekannt gemacht.

Abbildung 4.12: Funktionsweise der asymmetrischen Verschlüsselung.

4.8.4 Hybride Verschlüsselung

Ein Problem bei der asymmetrischen Verschlüsselung ist der hohe Rechenaufwand. Deswegen wird oftmals eine Kombination aus symmetrischer und asymmetrischer Verschlüsselung genutzt. Da jetzt nur noch der symmetrische Schlüssel asymmetrisch verschlüsselt wird, die Nachricht hinterher aber symmetrisch, bleibt der Rechenaufwand bei der hybriden Verschlüsselung relativ gering. Die Sicherheit leidet darunter nicht!

Ein Anwendungsfall hybrider Verschlüsselungstechniken wurde mit SSL realisiert. SSL ist im OSI-Schichtenmodell über der Transportschicht angesiedelt (Schicht 6 oder 7 je nach Anwendung). SSL bietet die Möglichkeit, Übertragungen über das Internet, unabhängig vom verwendeten Protokoll, gegen fremdes Abhören zu sichern. Vertrauliche Daten können so aus einem WEB Formular sicher an den **HTTPS**-Server gesendet werden. Auch E-Mails können über SSL beim POP3-Server abgerufen bzw. an einen SMTP-Server zu gesendet werden. Weiterhin kann die Echtheit des kontaktierten Servers mit Zertifikaten garantiert. Die Verbindung wird ständig überprüft.



Beim Aufbau einer SSL-Verbindung generiert der WWW-Browser einen zufälligen Schlüssel (Session Key), der für die Dauer der Verbindung zur Verschlüsselung genutzt wird. Damit die SSL-Verbindung nicht abgehört werden kann, muss zunächst dieser Session Key auf einem sicheren Weg zum WWW-Server übertragen werden. Um dies zu gewährleisten, wird der Session Key seinerseits mit einem Public Key Verfahren, meist RSA, verschlüsselt. Dazu präsentiert der WWW-Server seinen öffentlichen RSA-Schlüssel; der WWW-Browser verschlüsselt damit den Session Key und übermittelt das Ergebnis wieder dem WWW-Server. Erst danach wird die eigentliche Datenkommunikation aufgenommen.

... over SSL-Dienste:

Dienst	Port
https	443
ftps/ftps-data	990/989
telnets	992
pop3s	995

Weiter ausgeführtes Beispiel:

- 1 A generiert Session Key
- 2 A verschlüsselt Session Key mit Public Key von B
- 3 B entschlüsselt Session Key mit seinem Private Key
- 4 Datenübertragung zwischen A und B wird mit (symmetrischem) Session Key verschlüsselt und entschlüsselt

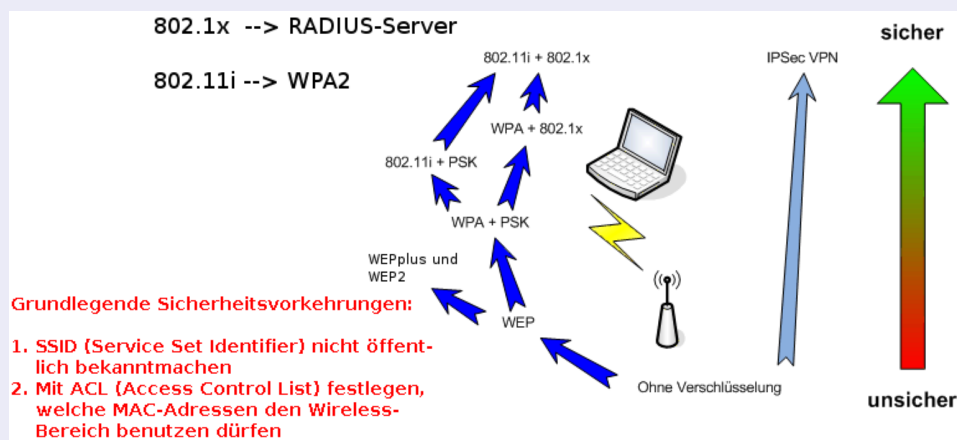
Abbildung 4.13: Ein Anwendungsbeispiel für hybride Verschlüsselung ist SSL.

4.9 Datensicherheit und Verschlüsselung im WLAN

Jeder Benutzer des Funknetzwerks könnte den gesamten Verkehr mithören. Deshalb muß bei WLANs zwingend auf eine verschlüsselte Datenübertragung geachtet werden:

- WEP (Wired Equivalent Privacy)
 - Verschlüsselt mit 64 bzw. 128 Bit wobei aber nur die ersten 24 Bit von der Hardware für jedes Datenpaket zufällig erzeugt werden
- TKIP(war WEP2)
 - Hier wird pro Datenpaket mit 128 Zufallsbits verschlüsselt
- WPA (Wi-Fi Protected Access)
 - aus Hardware-Kompatibilitätsgründen sehr stark an WEP2 angelegt
- WPA2 (802.11i)
 - WPA2 setzt als Verschlüsselung den AES-Algorithmus (Advanced Encryption Standard) ein, wie er auch z.B. von ssh verwendet wird

Sicherheit im WLAN



Der Anwender kann mit der Wahl des Verfahrens die Sicherheit und den Aufwand an seine Bedürfnisse anpassen.

http://www.tomsnetworking.de/content/reports/j2004a/neue_wlan_sicherheit/page10.html

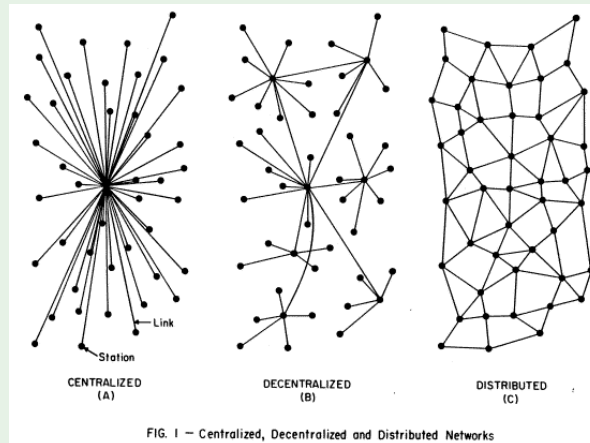
WEP eignet sich (bei wöchentlicher Erneuerung des Schlüssels) bestenfalls noch als Notlösung für ein Heimnetzwerk. WPA-PSK ist, bei genügend langen Schlüsseln, wesentlich sicherer und genau so leicht wie WEP einzurichten.

Abbildung 4.14: WLAN-Netze müssen besonders gut abgesichert werden.

Anhang A

Zusätzliche Folien

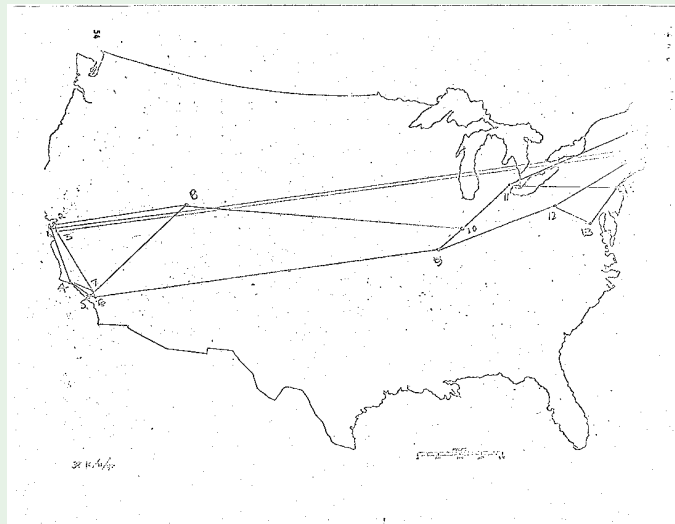
Klassisches und neuartiges Vermittlungssystem:



Um ca. 1960 schlug Paul Baran als Lösung ein verteiltes Netzwerk vor. Dem Pentagon gefiel das Konzept, AT&T ließ die Idee abblitzen.

← zurück

So was wollen wir haben. . . :



Unter <http://www.zakon.org/robert/internet/timeline/> kann man sich das gesamte Dokument ansehen!

← zurück

Der IMP und das Team von BBN 1968/69:



Der IMP war der erste Paketrouter. Er basierte auf einem Honeywell DPD-516 Rechner. US Senator Edward Kennedy gratulierte BBN in einem Telegram zum Bau des *“Interfaith” Message Processor*, und dankte für die *“ecumenical” efforts*.
Quellennachweise:

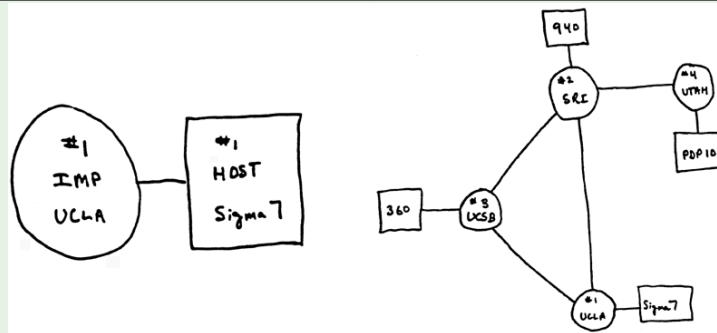
Linkes Bild: <http://www.computerhistory.org>

Rechtes Bild: <http://www.dirac.org/hackers/pictures/pic002/imp2.jpg>

◀ zurück

Ende 1969: Das ARPAnet besteht aus vier Knoten

Monat	Name	Standort	Hostrechner
09/69	UCLA	University of California Los Angeles	Sigma 7 ^a
10/69	SRI	Stanford Research Institute	SDS 940 ^b
11/69	UCSB	University of California Santa Barbara	IBM 360/75 ^c
12/69	UTAH	University of Utah	PDP 10 ^d



◀ zurück

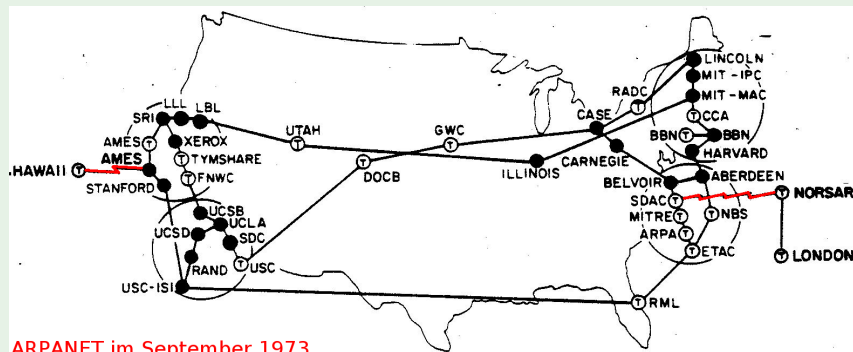
^aScientific Data Systems (SEX), für kommerziellen Einsatz

^bScientific Data Systems (Genie), für wissenschaftlichen Einsatz

^cIBM (OS/MVT)

^dDEC (Tenex)

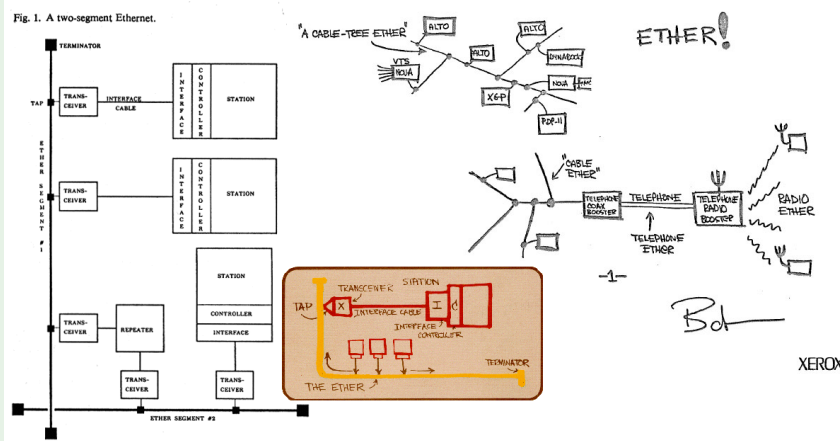
September 1973: Die Ozeane wurden überquert!



ARPANET im September 1973

← zurück

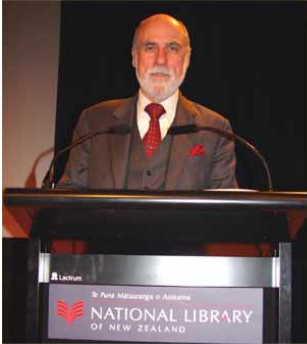
Ethernet:



← zurück

Das Bilder links und in der Mitte stammen aus einem Vortrag von Robert Metcalfe von der "National Computer Conference" im Juni 1976:
<http://www.acm.org/classics/apr96/>

1974: V. Cerf, der "Vater des Internet", publiziert TCP/IP



Format. An IP address is made up of four bytes of information (totaling 32 bits) expressed as four numbers between 0 and 255 shown separated by periods. For example, your computer's IP address might be 238.17.159.4, which is shown below in human-readable decimal form and in the binary form used on the Internet.

Example IP Address			
Decimal:	238 . 17 . 159 . 4		
Binary:	11101110	00010001	10011111 00000100

Each of the four numbers uses eight bits of storage, and so can represent any of the 256 numbers in the range between zero (binary 00000000) and 255 (binary 11111111). Therefore, there are more than 4 billion possible different IP addresses in all:

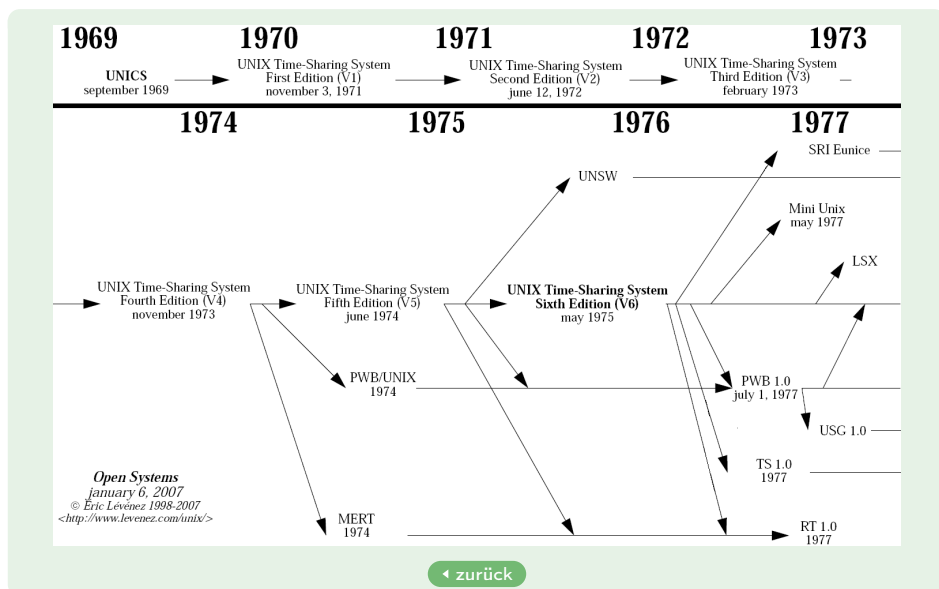
$$4,294,967,296 = 256 * 256 * 256 * 256$$

← zurück

Das Design der 32 Bit breiten IP-Adresse, bietet Platz für einen 8 Bit breiten Netzwerkanteil und einen 24 Bit breiten Hostanteil. Damit lassen sich 256 Netzwerke, jedes mit 16,777,216 eindeutigen Host-Adressen, bilden.

http://www.livinginternet.com/i/iw_ip.htm und

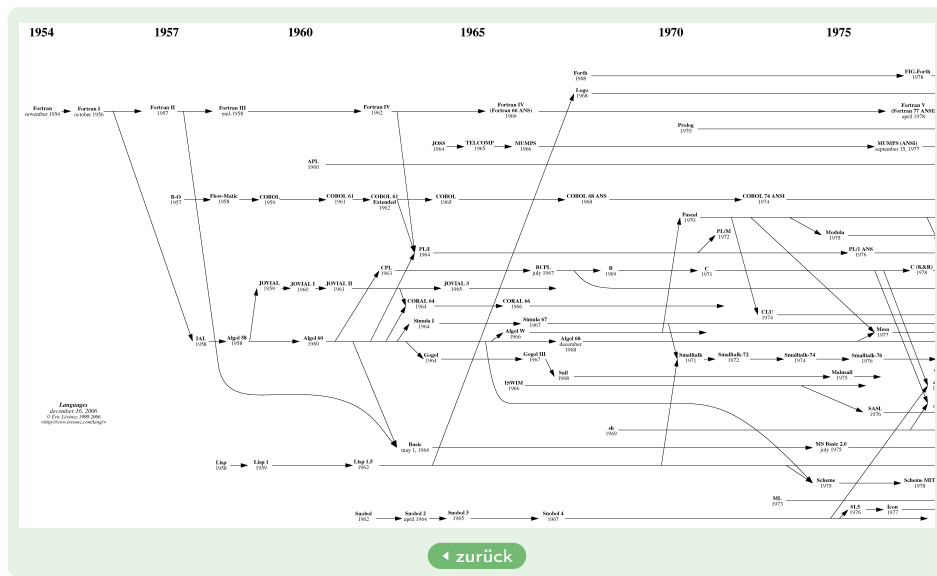
Ende der 70er Jahre: Unix wird Standard auf vernetzten Rechner



http://www.livinginternet.com/i/iw_ip.htm und

<http://www.levenez.com/unix/>

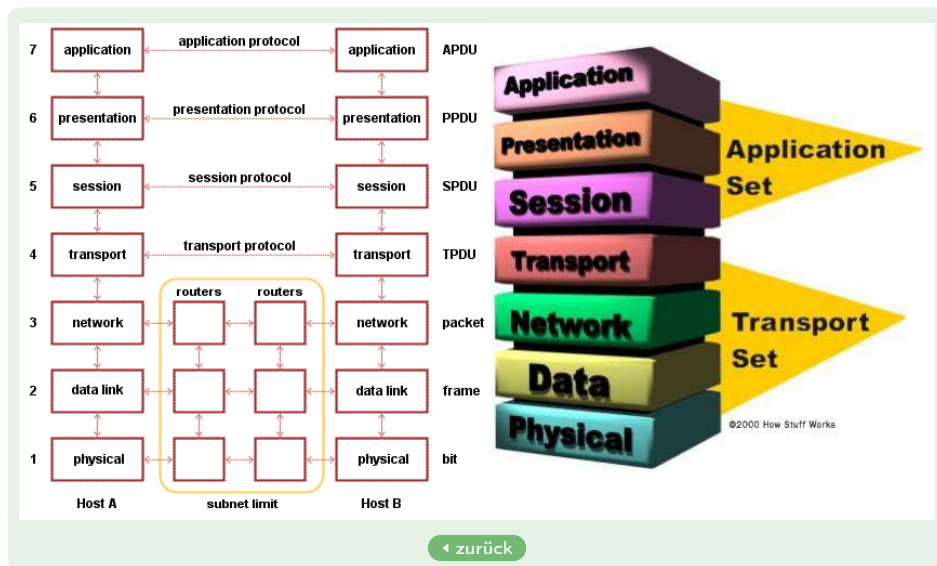
Ende der 70er Jahre: C wird zur Standardprogrammiersprache unter Unix



<http://www.levenez.com/lang/>

<http://www.isi.edu/in-notes/ien/ien98.txt>

1984: Das OSI-Schichtenmodell



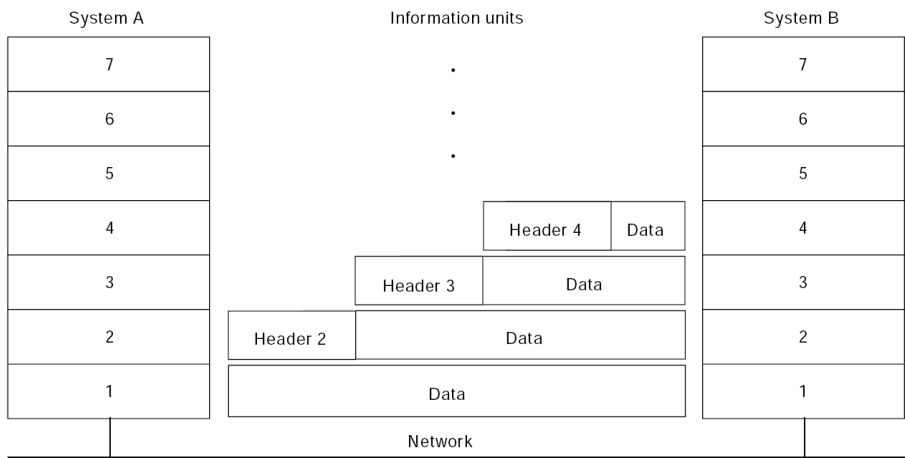
http://searchnetworking.techtarget.com/sDefinition/0,sid7_gci523729,00.html

<http://computer.howstuffworks.com/osi1.htm>

[Details](#)

Datenfluß im OSI-Schichtenmodell

Figure 1-6 Headers and Data Can Be Encapsulated During Information Exchange



◀ vorherige Seite

http://www.cisco.com/univercd/cc/td/doc/cisintwk/ito_doc/introint.htm

◀ zurück

1984/85: IEEE 802 Standards

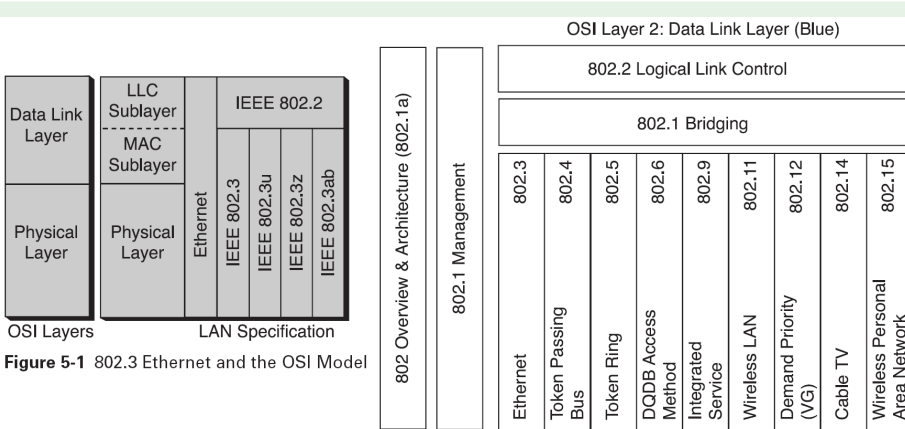
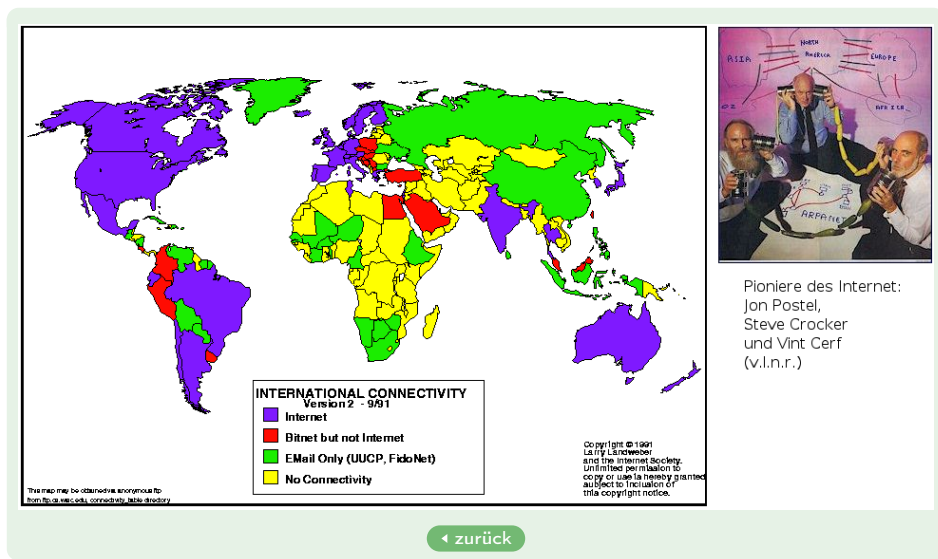


Figure 5-1 802.3 Ethernet and the OSI Model

Figure 5-2 LAN Specifications and OSI Model

◀ zurück

1990: Das Internet ist da



<http://www.netzmafia.de/skripten/internet/inetein1.html#1.1>

<http://www.cs.wisc.edu/~lhl/maps/>

1991: Der WWW wurde erfunden

Tim Berners-Lee
followed his dream of a better, easier way to communicate via computers on a global scale, which led him to create the World-Wide Web.

Robert Cailliau,
collaborator on the World-Wide Web project and first Web surfer.

The historic **NeXT computer** used by Tim Berners-Lee in 1990, on display in the Microcosm exhibition at CERN. It was the first web server, hypermedia browser and web editor.

← zurück

Die erste Web-Adresse hieß

<http://info.cern.ch/hypertext/WWW/TheProject.html>.

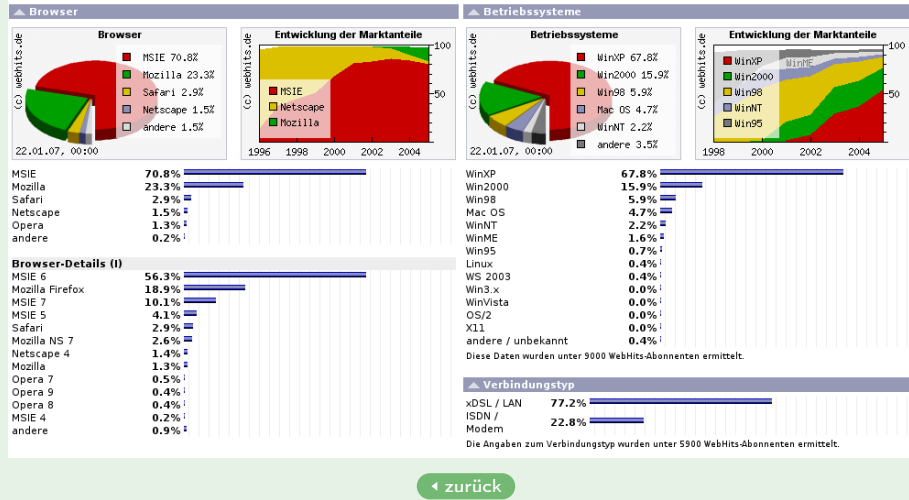
<http://info.cern.ch/>

1996: Browserschlacht

Web-Barometer

Aktuelle Daten zur allgemeinen Nutzung von Suchmaschinen, Browsern und Betriebssystemen sowie zu den Verbreitungszahlen der wichtigsten Plug-ins.
Historische Entwicklung der Marktanteile.

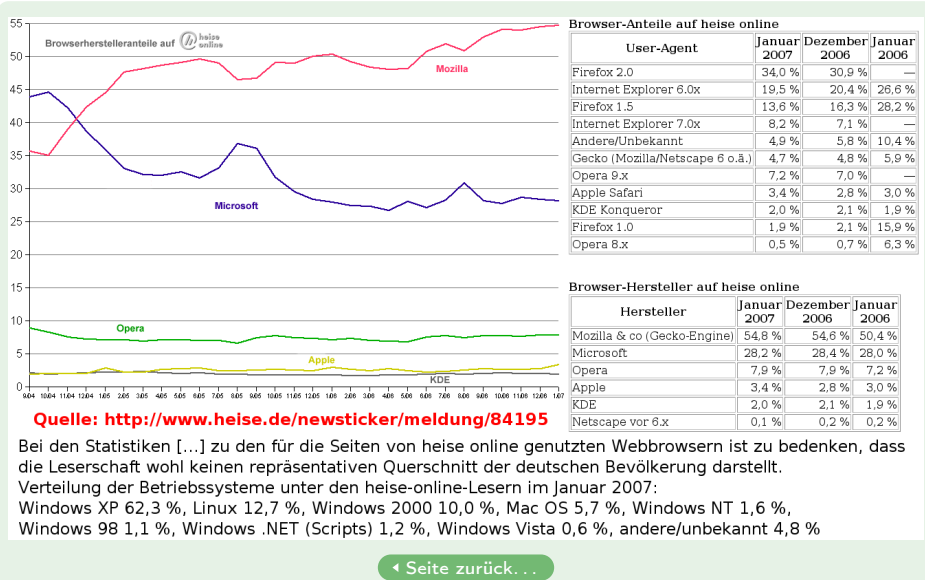
Aktualisiert am 22. Januar 2007



<http://http://www.webhits.de/deutsch/index.shtml?/deutsch/webstats.html>

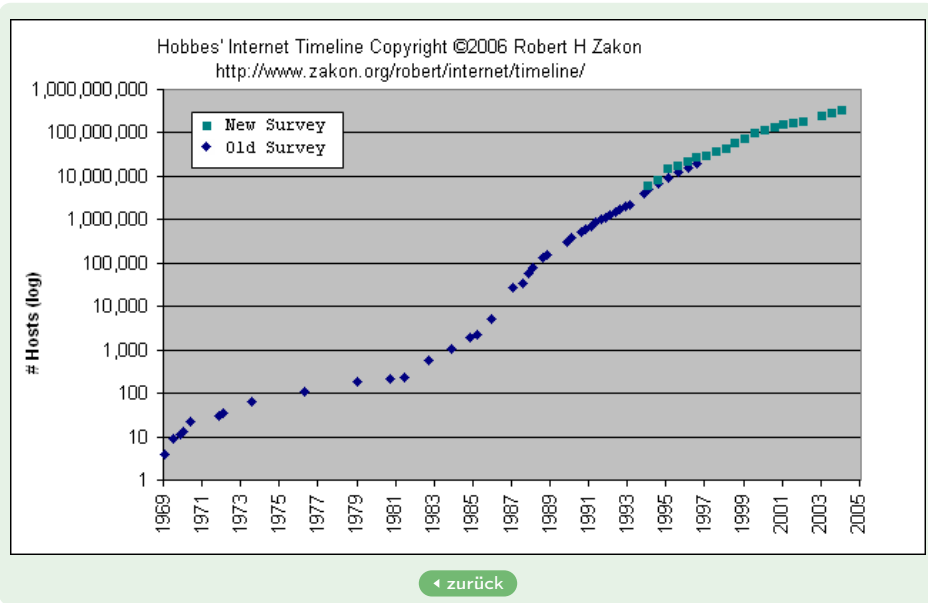
<http://www.metager.de/browser.html> [mehr...](#)

1996: Browserschlacht



← zurück

Anzahl der Hosts im Internet: Logarithmische Darstellung



<http://www.zakon.org/robert/internet/timeline/>

IPv6: Vor- und Nachteile im Vergleich zu IPv4

Mit IPv6 können für jeden Quadratmillimeter Erdoberfläche ca. 689 Billionen Adressen ($6.89 \cdot 10^{17}$) bereitgestellt werden (Erdradius von ca. 6267 km).

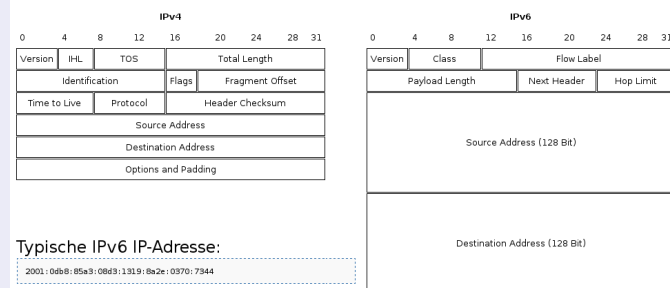
Vorteile

- Größerer Adressraum
- Einfacherer Aufbau des Headers
- IPsec Verschlüsselung im Protokoll

Nachteile

- Wenig praktische Erfahrung
- Unerwartete Detailprobleme
- Häufige Nachbesserungen

Aufbau des Headers



zurück

mehr...

IPv4 vs. IPv6: Unterschiede und Gemeinsamkeiten. . .

IP ist bei jeder Internetverbindung beteiligt und ist das grundlegende Protokoll in TCP/IP Netzwerken. IPv6 ist eine Weiterentwicklung von IPv4 und bringt einige Neuerungen und Verbesserungen mit.

IPv6 und IPv4 können gemeinsam betrieben werden. IPv6 sieht eine Abwärtskompatibilität zu IPv4 vor, um die Umstellung Schritt für Schritt vornehmen zu können. Wann diese Umstellung abgeschlossen sein wird ist weder terminlich festgelegt noch vorhersehbar. IPv4 funktioniert im Moment noch erstaunlich problemlos. Ein Umstellung der Infrastruktur (Hard- und Softwareseitig) ist mit großem zeitlichen und finanziellen Aufwand verbunden, der im Moment noch äußerst kritisch betrachtet wird.

Merkmal	IPv4	IPv6
Notation für <i>localhost</i>	127.0.0.1	0:0:0:0:0:0:0:1 bzw. ::1
Addressraum	32 Bit	128 Bit
verfügbare IP-Adressen	$2^{32} = 4.3 \cdot 10^9$	$2^{128} = 3.4 \cdot 10^{38}$

Zum Vergleich: Die Anzahl der Atome in unserem Universum wird auf ca. 10^{78} geschätzt!

[← zurück](#)

[Ein Seite zurück. . .](#)

OSI-Schichtenmodell: Übersicht 1

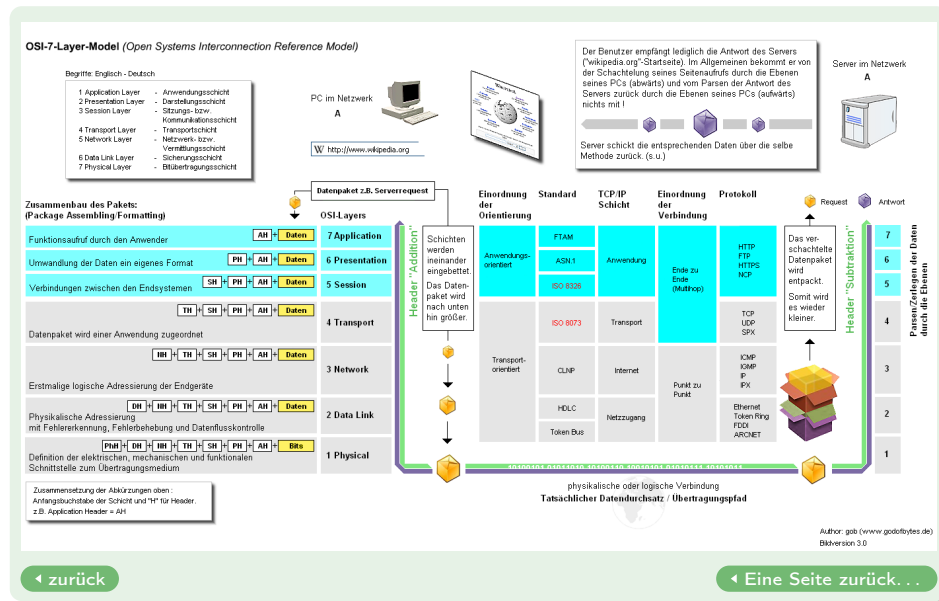
OSI-Schicht	Einordnung	Standard	TCP/IP-Schicht	Einordnung	Protokollbeispiel	Einheiten	Kopplungselemente	
7	Anwendung (Application)	FTAM	Anwendung	Ende zu Ende (Multihop)	HTTP FTP HTTPS NCP	Daten	Layer 4-7 Switch, Content Switch, Gateway	
6	Darstellung (Presentation)	ASN.1						
5	Sitzung (Session)	ISO 8326						
4	Transport (Transport)	ISO 8073	Transport	Punkt zu Punkt	TCP UDP SPX	Segmente	Router, Layer-3 Switch	
3	Vermittlung (Network)	CLNP	Internet		ICMP IGMP IP IPX	Pakete		
2	Sicherung (Data Link)	HDLC	Netzzugang		Ethernet Token Ring FDDI ARCNET	Rahmen (Frames)	WLAN Access Point, Switch, Bridge (Netzwerk)	
1	Datenübertragung (Physical)	Token Bus				Bits		Hub, Repeater

[◀ zurück](#)[▶ OSI-Überblick 2](#)

[← zurück](#)

[OSI-Überblick 2](#)

OSI-Schichtenmodell: Übersicht 2



Verschiedene Router

Beschreibung

Je nach Einsatzgebiet und Aufgabenstellung wählt man einen entsprechenden Router aus:

- Router für DSL-Verbindungen
 - z.T. extrem günstig
 - geringer Funktionsumfang
 - meist Plug&Play
- Router für kleine bis mittlere Unternehmen
 - ISDN-Fallback
 - komplett frei konfigurierbar
 - modular erweiterbar
 - teuer

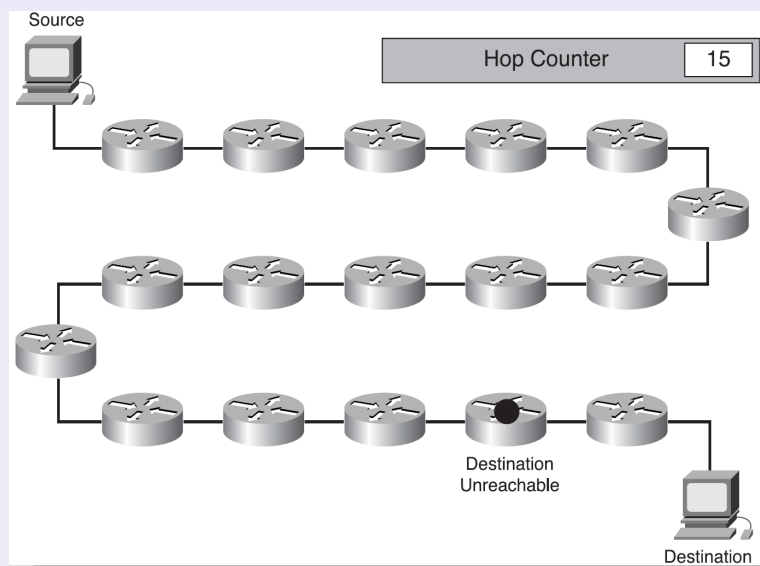
Beispiele



zurück

Destination Unreachable:

Konvergiertes RIPv1-Netzwerk



◀ zurück

Der DHCP-Server dokumentiert seine Arbeit in /var/log/syslog.
Eine erfolgreiche IP-Adressvergabe sieht wie folgt aus:

Feb 27 12:53:23 AP059132 dhcpd: **DHCPDISCOVER** from
00:0c:29:99:c9:cc via eth0

Feb 27 12:53:24 AP059132 dhcpd: **DHCPOFFER** on
192.168.59.138 to 00:0c:29:99:c9:cc via eth0

Feb 27 12:53:24 AP059132 dhcpd: **DHCPREQUEST** for
192.168.59.138 (192.168.59.132) from 00:0c:29:99:c9:cc via eth0

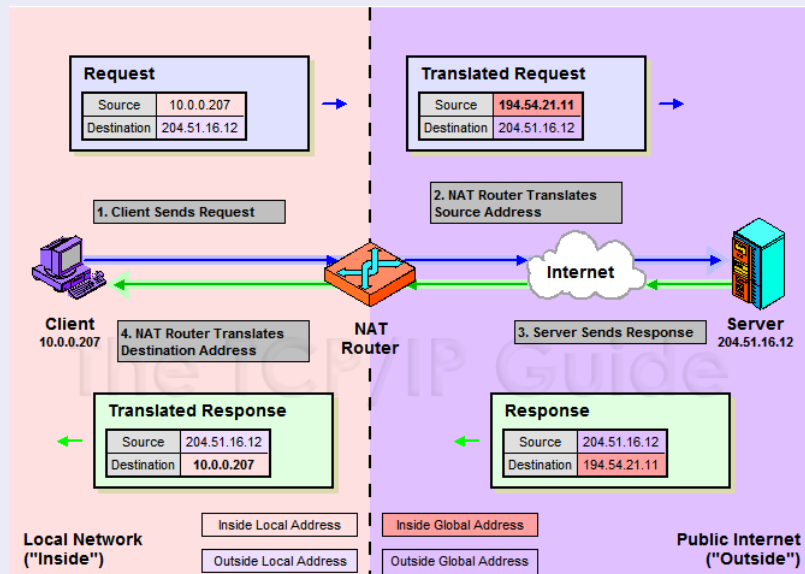
Feb 27 12:53:24 AP059132 dhcpd: **DHCPACK** on 192.168.59.138
to 00:0c:29:99:c9:cc via eth0

Versucht dagegen ein nicht berechtigter Client sich eine IP-Adresse
zu besorgen, dann sieht man:

Feb 28 09:50:29 AP059132 dhcpd: **DHCPDISCOVER** from
00:1a:4b:46:97:00 via eth0: unknown client

◀ zurück

NAT - Funktionsweise



◀ zurück

Systematik der IP-Klassen

Klasse	1. Oktett	Binärdarstellung
A	1...	00000001
	...	...
	126...	01111110
B	128...	10000001
	...	...
	191...	10111111
C	192...	11000000
	...	...
	223...	11011111
D	224...	11100000
	...	...
	239...	11101111

Ausnahme: localhost

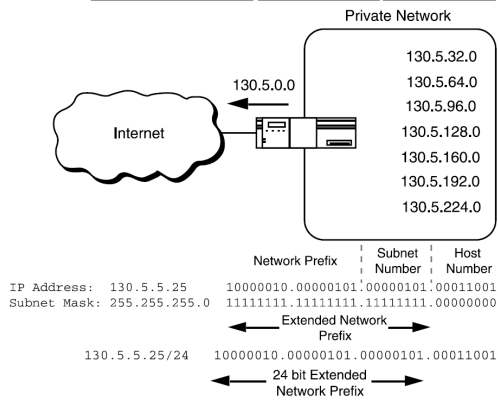
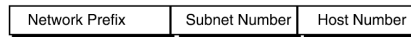
127.0.0.1

◀ zurück

Two-Level Classful Hierarchy



Three-Level Subnet Hierarchy

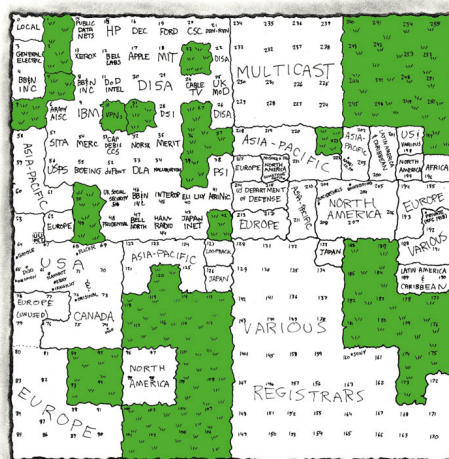


← zurück

Klasse	CIDR
A	/8
B	/16
C	/24

Daneben sind auch alle Möglichkeiten zw. /9 bis zu /32 (ein Host, bzw. Punkt-zu-Punkt-Verbindungen oder Hostroute) sowie /0 (alle Adressen oder auch Default Route) erlaubt!

MAP OF THE INTERNET
THE IPv4 SPACE, 2006



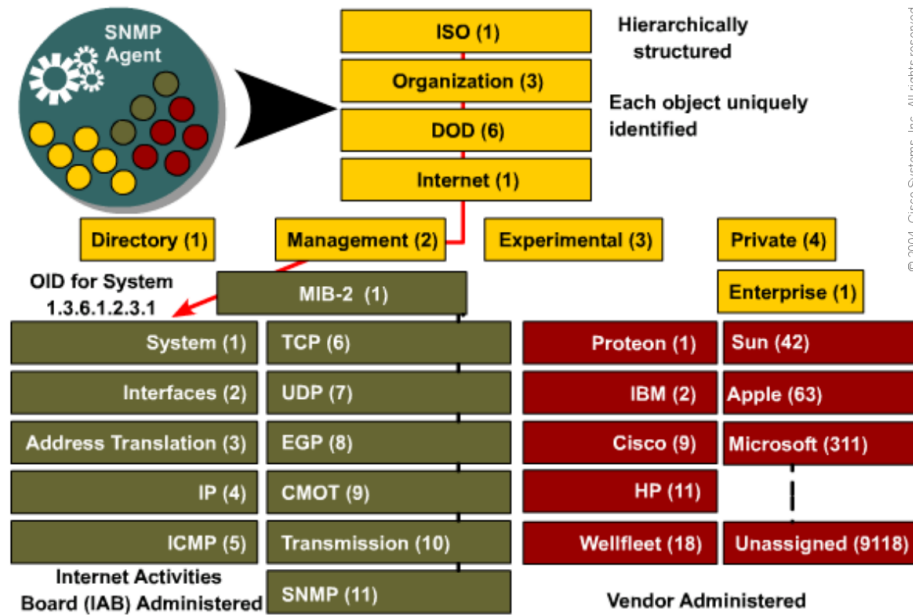
THIS CHART SHOWS THE IP ADDRESS SPACE ON A PLANE USING A FRACTAL MAPPING WHICH PRESERVES GROUPING--ANY CONSECUTIVE STRING OF IPs WILL TRANSLATE TO A SINGLE COMPACT, CONTIGUOUS REGION ON THE MAP. EACH OF THE 256 NUMBERED BLOCKS REPRESENTS ONE /8 SUBNET (CONTAINING ALL IPs THAT START WITH THAT NUMBER). THE UPPER LEFT SECTION SHOWS THE BLOCKS SOLD DIRECTLY TO CORPORATIONS AND GOVERNMENTS IN THE 1990s BEFORE THE RIRs TOOK OVER ALLOCATION.

0 1 14 15 16 19 →
3 2 13 12 17 18
4 7 8 11
5 6 9 10



■ UNALLOCATED BLOCK

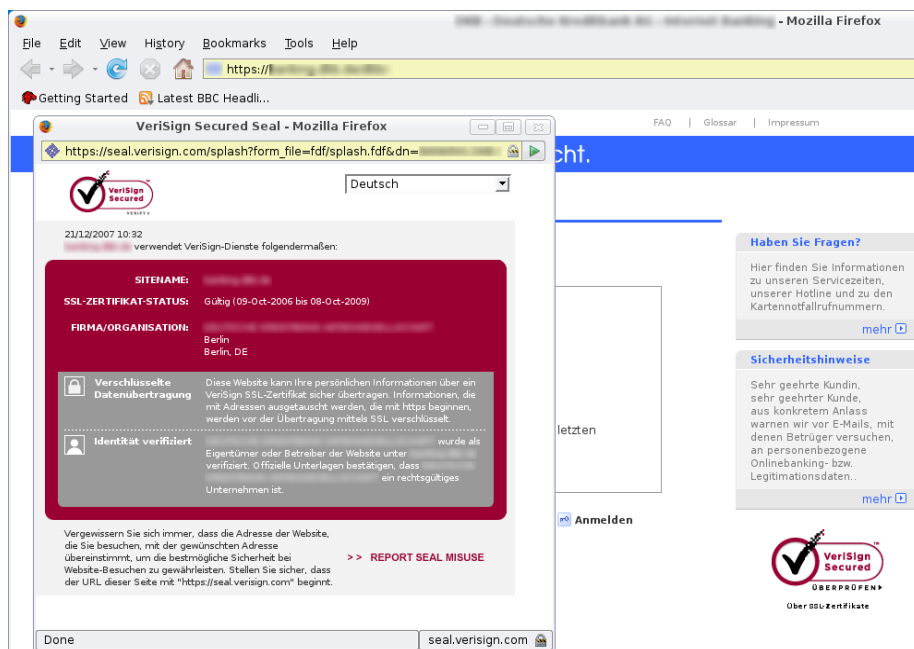
← zurück



© 2004, Cisco Systems, Inc. All rights reserved

Die Objekte eines SNMP-fähigen Geräts sind hierarchisch aufgebaut und können entweder über eine Nummer oder ihren Namen angesprochen werden.

◀ zurück



◀ zurück

Schlüssellänge bei symmetrischer Verschlüsselung	Schlüssellänge bei asymmetrischer Verschlüsselung
56 Bit	384 Bit
64 Bit	512 Bit
80 Bit	768 Bit
112 Bit	1792 Bit
128 Bit	2304 Bit

Aufgrund der Arbeitsweise müssen bei asymmetrischer Verschlüsselung (z.B. GPG) wesentlich größere Schlüssellängen verwendet werden, als bei symmetrischer Verschlüsselung (z.B. DES), um die selbe Sicherheit zu erreichen.

Vgl. dazu auch:

http://www.tcp-ip-info.de/security/key_comparison.htm [← zurück](#)

Kosten in Dollar	http://www.tcp-ip-info.de/security/key_length.htm (RSA 1996)					
	40 Bit	56 Bit	64 Bit	80 Bit	112 Bit	128 Bit
1.000	4 min	145 Tg.	100 J.	7 Mio. J.	10^{16} J.	10^{21} J.
10.000	20 sec	14 Tg.	10 J.	700.000 J.	10^{15} J.	10^{20} J.
100.000	2 s	35 h	1 J.	70.000 J.	10^{14} J.	10^{19} J.
1 Mio.	200 ms	3,5 h	37 Tg.	7.000 J.	10^{13} J.	10^{18} J.
10 Mio.	20 ms	21 min	4 Tg.	700 J.	10^{12} J.	10^{17} J.
100 Mio.	2 ms	2 min	9 h	70 J.	10^{11} J.	10^{16} J.
1 Mrd.	200 μ s	13 s	1 h	7 J.	10^{10} J.	10^{15} J.
10 Mrd.	20 μ s	1 s	5,4 min	245 Tg.	10^9 J.	10^{14} J.
100 Mrd.	2 μ s	100 ms	32 s	24 Tg.	100 Mio. J.	10^{13} J.
1 Bio.	200 ns	10 ms	3 s	2,4 Tg.	10 Mio. J.	10^{12} J.
10 Bio.	20 ns	1 ms	300 ms	6 h	1 Mio. J.	10^{11} J.

Der lange Zeit als sicher geltende DES-Schlüssel ist längst mit Brute-Force-Methoden knackbar. In Kryptologenkreisen geht man davon aus, daß erst Schlüssel mit einer Länge von >100 Bit (bei symmetrischer Verschlüsselung) sicher sind. Das Alter des Universums beträgt 10^{10} Jahre, die NSA verfügte damals über ca 10-100 Mrd US-\$ Jahrebudget. [← zurück](#)

Anhang B

Weitergehende Informationen

B.1 Bolt Beranek and Newman

Bolt, Beranek and Newman (BBN) in Cambridge (Massachusetts) waren ein privates Forschungsunternehmen, welches 1968 das ARPANET - den Vorgänger des Internets - aufzubauen. [8,] BBN still provides networking services, but is no longer an independent company having been taken over by GTE in 1997. [10,]

B.2 Robert Metcalfe

Metcalfe erwarb am MIT zwei Bachelor Abschlüsse in Elektrotechnik und Industrial Management. In seiner Doktorarbeit an der Harvard University beschäftigte er sich mit dem Thema Paketvermittlung in Computernetzwerken.

Während seiner Doktorarbeit (s. Grafik) arbeitete er im Xerox Palo Alto Research Center (Xerox PARC) und wurde mit der Aufgabe betraut, die Firmenrechner miteinander zu vernetzen. Auf der Basis des ersten funkbasierten Computernetzwerks Alohanet entwickelte er 1973 eine modifizierte, kabelgebundene Netzwerktechnologie, das Ethernet. 1979 verließ er Xerox, um seine eigene Firma 3Com, einen Hersteller von Komponenten für Computernetzwerke, zu gründen. Er überzeugte erfolgreich die Firmen DEC, Intel und Xerox mit ihm zusammenzuarbeiten, um Ethernet zum Standard zu machen. Damals konkurrierende Techniken waren die proprietären Systeme Token Ring und ARCNET, die beide bald in einer wahren Flut von Ethernet-Produkten untergingen. 3Com wurde dabei zu einem großen Unternehmen.

B.3 Internet

Das Internet ist ein logisches Netzwerk, bestehend aus vielen physischen Netzwerken die alle mit einem IP-Protokoll betrieben werden. [12,]

B.4 Usenet

Das Usenet ist ein weltweites, elektronisches Netzwerk, das Diskussionsforen (sogenannte Newsgroups") aller Art bereitstellt und an dem grundsätzlich jeder teilnehmen kann. Der Teilnehmer verwendet dazu üblicherweise einen Newsreader.

B.5 MILNET

Das MILNET ist ein militärisches Kommunikationsnetz der US-Streitkräfte. Es entstand 1983 aus dem ARPANET. Das ARPANET war ursprünglich ein gemeinsames Projekt von wissenschaftlichen und militärischen Einrichtungen. Ab 1983 wurde dies aus Sicherheitsgründen zweigeteilt in das militärische MILNET und das nun wissenschaftliche ARPANET, aus dem sich später das Internet entwickelte. In der Übergangsphase waren beide Netze durch Gateways miteinander verbunden.

B.6 BITNET

Das BITNET war ein kooperatives Rechnernetzwerk. Es verband Großrechner von wissenschaftlichen Institutionen und öffentlichen Forschungseinrichtungen in den USA. Die erste Verbindung wurde 1981 zwischen der City University of New York (CUNY) und der Yale University hergestellt.

BITNET benutzte ein einheitliches Kommunikationsverfahren, das auch von Rechnernetzen wie EARN in Europa und NetNorth in Kanada eingesetzt wurde. So entstand ein weltweites homogenes Rechnernetz. Auf seinem Höhepunkt zu Beginn der 1990er Jahre kommunizierten an die 3500 Rechner in über 1400 Organisationen miteinander. Danach wurden die Rechner immer kleiner und zahlreicher. Das Internet übernahm die Kommunikation. Das BITNET wurde aufgelöst.

B.7 RFC

Die Requests for Comments (zu deutsch: Aufforderung zu Kommentaren) sind eine Reihe von technischen und organisatorischen Dokumenten des RFC-Editor zum Internet (ursprünglich ARPANET), die am 7. April 1969 begonnen wurden. Bei der ersten Veröffentlichung noch im ursprünglichen Wortsinne zur Diskussion gestellt, behalten RFC auch dann ihren Namen, wenn sie sich durch allgemeine Akzeptanz und Gebrauch zum Standard entwickelt haben.

B.8 Wi-Fi

Die Wi-Fi Alliance (Wireless Fidelity) ist eine 1999 ursprünglich unter dem Namen WECA gegründete Organisation bestehend aus über 200 Unternehmen, die es sich zur Aufgabe gemacht hat, Produkte verschiedener Hersteller auf der Basis des IEEE 802.11 Standards zu zertifizieren und somit den Betrieb mit verschiedenen Wireless Geräten zu gewährleisten (Interoperabilität).

Hintergrund war, daß in vielen Produkten der Standard nicht vollständig implementiert bzw. durch proprietäre Erweiterungen aufgeweicht wurde. Somit ergaben sich häufig Inkompatibilitäten zwischen Produkten verschiedener Hersteller.

Die Wi-Fi testet entsprechende Komponenten gemäß eigener Richtlinien. Produkte, die diese Prüfung bestehen, erhalten das Wi-Fi-Zertifikat und dürfen damit das Wi-Fi-Logo tragen. Allerdings werden nur die Produkte der Wi-Fi-Mitglieder getestet. Die Mitglieder müssen zusätzlich zur Mitgliedsgebühr für jede geprüfte Komponente eine Gebühr entrichten. Ein fehlendes Wi-Fi-Logo stellt daher nicht zwingend eine Abweichung vom Standard dar.

Glossar

ASIC	Application Specific Integrated Circuit	53
AUTO-MDI	AUTO(Automatische Erkennung) MDI/MDIX MDI steht für Medium Dependent Interface (X für Crossover) und bezieht sich auf einen RJ45 Ethernet Anschluss. Auto MDI/MDIX wählt automatisch für die Sende- und Empfangsleitung eines Ports, egal ob das angeschlossene Ethernet-Kabel ein Crossover(gekreuzt) oder ein Patchkabel(nicht gekreuzt 1:1 geschaltet) ist, die jeweils richtige Konfiguration des eigenen Ports. Somit können alle Auto MDI/MDIX gekennzeichneten Ports als Uplink/Patch Port genutzt werden.	49
BBN	Bolt, Beranek and Newman	10
BITNET	Because It's Time NETwork	114
BPDU	Bridge Protocol Data Units	52
CSMA/CD	Carrier Sense Multiple Access with Collission Detection	42
DHCP	Dynamic Host Configuration Protocol	84
DMZ	Demilitarized Zone	83
DNS	Domain Name Service	19
DSL	Dual Subsriber Line. Populäre, kostenpflichtige Verbindungstechnik für Privatanwender und kleine Betrieb um über einen analogen bzw. digitalen Telefonzugang das Internet zu nutzen	57
FSO	Free Space Optics	30
GAN	Global Area Network	12
HTML	Hypertext Markup Language	11
HTML	World Wide Web	11
HTTP	Hypertext Transfer Protocol	19
ICMP	Internet Control Message Protocol	19

IETF	Internet Engineering Task Force	64
ISO	International Organization for Standardization	18
LAN	Local Area Network	12
MAN	Metropolitan Area Network	12
MILNET	MILitary NETwork	113
MTU	Maximum Transfer Unit	68
NAT	Network Address Translation	85
NWay	NWay ist ein von der IEEE standardisiertes Protokoll zum automatischen Feststellen von Verbindungscharakteristiken in Ethernet-Netzwerken. Damit kann zwischen zwei Endpunkten die höchst mögliche Datenübertragungsrate ermittelt und eingestellt werden kann. Das NWay Protokoll (manchmal auch als auto-sensing bezeichnet) wurde 1994 von National Semiconductor entwickelt.	37
OSI	Open System Interconnection	16
OSPF	Open Shortest Path First	64
PAN	Personal Area Network	12
PAT	Port Address Translation	85
PPP	Point-to-Point	19
RFC	Request For Comments.	57
RIP	Routing Information Protocol (RFC 1058)	62
TCP	Transmission Control Protocol	19
UDP	User Datagram Protocol	19, 69
URL	Uniform Resource Locator	11
USENET	Unix User Network	10
VLAN	Virtual Local Area Network	79
WAN	Wide Area Network	12
WECA	Wireless Ethernet Compatibility Alliance, s.a. Wi-Fi	114
WLAN	Wireless LAN (Wireless Local Area Network)	29

Literaturverzeichnis

- [1] Katie Hafner, Matthew Lyon: *Arpa Kadabra oder die Geschichte des Internet*. dpunkt.verlag, Heidelberg, 2. Auflage 2000.
- [2] Andrew S. Tanenbaum: *Computernetzwerke*. Pearson Studium, München, 4. überarbeitete Auflage 2003.
- [3] http://www.cisco.com/univercd/cc/td/doc/product/access/acs_mod/cis4000/4000m/4000sig/cable.
- [4] http://www.weberco.ch/t_rs232_2.htm
- [5] Linux-User 09/05, S.24: Netzwerke überwachen und kartieren (<http://netwhistler.spb.ru>)
- [6] Linux-User 09/05, S.24: Mit List und Tücke (Verzeichnisinhalt en detail mit ls)
- [7] <http://www.cybergeography.org/atlas/historical.html>, Januar 2007
- [8] <http://www.bbn.com>, Januar 2007
- [9] <http://www.zakon.org/robert/internet/timeline/>
- [10] http://www.mundi.net/maps/maps_001/
- [11] <http://www.acm.org/classics/apr96/>
- [12] <http://world.std.com/franl/worm.html>
- [13] http://www.cs.wisc.edu/lhl/maps/version_2.bmp
- [14] <http://www.extendedsystems.de>
- [15] http://www.linux-magazin.de/heft_abo/ausgaben/2007/09/verkehrsregelung
- [16] http://www.livinginternet.com/i/ii_cerf.htm