



**SLEZSKÁ
UNIVERZITA**

FILOZOFICKO-
PŘÍRODOVĚDECKÁ
FAKULTA V OPAVĚ

Šárka Vavrečková

Skripta do předmětu

Počítačové sítě a decentralizované systémy

Ústav informatiky
Filozoficko-přírodovědecká fakulta v Opavě
Slezská univerzita v Opavě

Opava
7. dubna 2025

Anotace: Tato skripta jsou určena pro studenty předmětu *Počítačové sítě a decentralizované systémy* pro navazující magisterské studium na Ústavu informatiky Slezské univerzity v Opavě. V předmětu se zabýváme prostředky a postupy používanými v počítačových sítích, navazujeme na obdobný předmět z bakalářského stupně studia.

Počítačové sítě a decentralizované systémy

RNDr. Šárka Vavrečková, Ph.D.

Dostupné na: <http://vavreckova.zam.slu.cz/sitedec.html>

Ústav informatiky
Filozoficko-přírodovědecká fakulta v Opavě
Slezská univerzita v Opavě
Bezručovo nám. 13, Opava

Sázeno v systému L^AT_EX

Předmluva

Co najdeme v těchto skriptech

Tato skripta jsou určena pro studenty Ústavu informatiky Slezské univerzity v Opavě. Obsahují látku vyučovanou na přednáškách předmětu *Počítačové sítě a decentralizované systémy*, ve kterém se zabýváme (jak název napovídá) počítačovými sítěmi.

Některé oblasti jsou „navíc“ (jsou označeny ikonami fialové barvy), ty nejsou probírány a ani se neobjeví na zkoušce – jejich úkolem je motivovat k dalšímu samostatnému studiu či pokusům nebo pomáhat v budoucnu při získávání dalších informací. Pokud je fialová ikona před názvem kapitoly (sekce), platí pro vše, co se v dané kapitole či sekci nachází.

Značení

Ve skriptech se používají následující barevné ikony:

-  Nové *pojmy*, značení apod. jsou značeny modrým symbolem, který vidíme zde vlevo.
-  Konkrétní *postupy* a nástroje, způsoby řešení různých situací, do kterých se může správce počítačového vybavení dostat, atd. jsou značeny také modrou ikonou.
-  Některé části textu jsou označeny fialovou ikonou, což znamená, že jde o *nepovinné úseky*, které nejsou probírány (většinou; studenti si je mohou podle zájmu vyžádat nebo sami prostudovat). Jejich účelem je dobrovolné rozšíření znalostí studentů o pokročilá témata, na která obvykle při výuce nezbývá moc času.
-  Žlutou ikonou jsou označeny odkazy, na kterých lze získat *další informace* o tématu. Nejčastěji u této ikony najdeme webové odkazy na stránky, kde se dané tematice jejich autoři věnují podrobněji.
-  Červená je ikona pro *upozornění* a poznámky.

Pokud je množství textu patřícího k určité ikoně větší, je celý blok ohraničen prostředím s ikonami na začátku i konci, například pro definování nového pojmu:



Definice

V takovém prostředí definujeme pojem či vysvětlujeme sice relativně známý, ale komplexní pojem s více významy či vlastnostmi.



Podobně může vypadat prostředí pro delší postup nebo delší poznámku či více odkazů na další informace. Mohou být použita také jiná prostředí:



Příklad

Takto vypadá prostředí s příkladem, obvykle nějakého postupu. Příklady jsou obvykle komentovány, aby byl jasný postup jejich řešení.



Úkol

Otázky a úkoly, náměty na vyzkoušení, které se doporučuje při procvičování učiva provádět, jsou uzavřeny v tomto prostředí. Pokud je v prostředí více úkolů, jsou číslovány.



Obsah

Předmluva	iii
1 Pojmy a standardy k počítačovým sítím	1
1.1 Přehled základních pojmů a postupů	1
1.2 Aktivní síťové prvky	2
1.3 Co a jak přenášíme po počítačové síti	3
1.3.1 Protokoly	3
1.3.2 Vlastnosti protokolů	4
1.4 Standardy a standardizační instituce	5
1.4.1 Standardizační instituce pro počítačové sítě	5
1.4.2 Volně dostupné standardy	8
1.5 Referenční model ISO/OSI	10
1.5.1 Přehled modelu	10
1.5.2 Spolupráce protokolů	14
1.5.3 Protokolové zásobníky	18
1.6 Síťový model TCP/IP	19
1.7 Charakteristiky přenosu a přenosových cest	20
1.7.1 Přenos v základním a přeloženém pásmu	20
1.7.2 Multiplex	24
2 Lokální sítě – Ethernet	28
2.1 Co je to Ethernet	28
2.2 Komunikace v Ethernetu	29
2.3 Ethernet na linkové vrstvě	31
2.3.1 Adresace na vrstvě L2	31
2.3.2 Podvrstvy L2	32
2.3.3 EtherType	33
2.3.4 Formát ethernetového rámce	34
2.3.5 Tabulka MAC adres na switchi	39
2.4 Ethernet na fyzické vrstvě	40
2.4.1 Křížení a krimpování	41

2.4.2	Parametry	42
2.4.3	Implementace fyzické vrstvy	44
2.5	Průběh přenosu	46
2.5.1	Přenos v polovičním duplexu	46
2.5.2	Přenos v plném duplexu	49
2.6	Technologie související s Ethernetem	50
2.6.1	Autonegociace	50
2.6.2	Napájení přes Ethernet	51
3	Další témata k lokálním sítím	53
3.1	Adresy v rámcích a paketech	53
3.2	Přehled rodiny standardů IEEE 802	55
3.3	Přístupové metody	56
3.3.1	Token Ring	57
3.3.2	100VG-AnyLAN	58
3.4	VLAN	59
3.4.1	VLAN na jednom switchi	60
3.4.2	VLAN rámce na cestě přes trunk	62
3.4.3	Komunikace mezi různými VLAN sítěmi	65
3.5	STP – kostra v grafu	67
3.5.1	Protokol STP	67
3.5.2	Konvergence sítě switchů	70
3.5.3	Variety protokolu STP	73
3.6	EtherChannel	74
3.6.1	Vlastnosti	74
3.6.2	Řízení toku	75
3.6.3	Vytvoření kanálu	75
3.7	SAN sítě	78
3.7.1	Fibre Channel	78
3.7.2	iSCSI	79
4	Rozlehlé sítě a telekomunikace	81
4.1	WAN sítě	81
4.1.1	Struktura WAN sítí	81
4.1.2	Protokoly vrstvy L2 pro WAN sítě	82
4.2	Nejznámější WAN sítě	84
4.2.1	X.25 – přepínání paketů	85
4.2.2	Frame Relay – přepínání rámců	85
4.2.3	ATM – přepínání buněk	88
4.2.4	MPLS – přepínání značek	89
4.3	Infrastruktura optických sítí	94
4.3.1	SONET/SDH	94
4.3.2	WDM	94
4.4	Telekomunikační sítě	96

4.4.1	Využití telekomunikačních sítí pro přenos dat	96
4.4.2	Slučování linek	97
4.5	Přístupové sítě	99
4.5.1	Protokol PPP	99
4.5.2	Digitální přístupové telekomunikační sítě	101
4.5.3	DOCSIS	106
4.5.4	FFTx	107
4.5.5	Obecně k zařízením v přístupových sítích	109
4.5.6	EFM	109
4.6	Pobočkové ústředny	110
5	Bezdrátové a mobilní sítě	111
5.1	Bezdrátové technologie	111
5.2	Wi-fi	112
5.2.1	Access point a WDS	112
5.2.2	Fyzická vrstva	112
5.2.3	Interference	113
5.2.4	Diagnostické nástroje pro fyzickou vrstvu	115
5.2.5	Signál a antény	116
5.2.6	Identifikátory a služby	118
5.2.7	Přístupová metoda	119
5.2.8	Wi-fi rámce	119
5.3	Zabezpečení bezdrátové komunikace	124
5.3.1	AAA	124
5.3.2	WPS	127
5.3.3	Jak tedy zabezpečit bezdrátovou síť	128
5.4	Centrálně řízená bezdrátová síť	129
5.5	Mobilní technologie	130
5.5.1	Struktura mobilní sítě	130
5.5.2	První generace (1G)	131
5.5.3	Druhá generace (2G)	132
5.5.4	Přechodová generace (2.5G)	133
5.5.5	Třetí generace (3G)	133
5.5.6	Čtvrtá generace	135
5.5.7	Pátá generace	135
5.6	Další bezdrátové sítě	137
5.6.1	Zajištění provozu sítě	137
5.6.2	Komunikace na logické úrovni	137
6	Síťová a transportní vrstva	138
6.1	IP adresy	138
6.1.1	Hierarchie adres a sítí	138
6.1.2	CIDR	139
6.2	Překlad adres	141

6.2.1	NAT na hranici klientské sítě	141
6.2.2	Carrier-Grade NAT se sdíleným adresovým prostorem	143
6.2.3	Koexistence IPv4 a IPv6	144
6.3	Protokol ICMP a zprávy	145
6.3.1	Účel protokolu ICMP	145
6.3.2	ICMP verze 6	147
6.4	Správa skupin	148
6.5	Objevování sousedů	150
6.5.1	Tabulky sousedů	150
6.5.2	K protokolům	151
6.5.3	Bezpečnost	153
6.6	Jak funguje router	154
6.7	Protokoly transportní vrstvy	154
7	Decentralizované a distribuované systémy	156
7.1	Typy systémů	156
7.2	Bridging, switching, routing	157
7.3	Směrování	159
7.3.1	Jak směřujeme	159
7.3.2	Autonomní systém	162
7.3.3	Směrovací algoritmy	164
7.3.4	Směrovací protokol RIP	166
7.3.5	Směrovací protokoly IGRP a EIGRP	168
7.3.6	Směrovací protokol OSPF	169
7.3.7	Směrovací protokol IS-IS	172
7.3.8	Směrovací protokol BGP	172
7.3.9	Softwarový router	174
7.4	Mechanismus DNS	175
7.4.1	Domény a jmenné adresy	175
7.4.2	Zóny a DNS servery	176
7.4.3	DNS dotazy	177
7.4.4	Tabulka hostitelů	179
7.4.5	Bezpečnost v DNS	181
7.4.6	Využití DNS při zabezpečení e-mailové komunikace	182
7.5	QoS	183
7.6	VoIP a videotelefonie	185
7.6.1	Princip	185
7.6.2	Protokoly	185
7.6.3	Videotelefonie a videokonference	187
8	Management sítě a monitoring	189
8.1	Správa sítě	189
8.1.1	NMS	189
8.1.2	ISO model řízení sítě	190

8.1.3	Management v ISO/OSI	191
8.2	Management v TCP/IP	191
8.2.1	SNMP	191
8.2.2	MIB-II	193
8.2.3	Používání protokolu SNMP	196
8.3	Syslog	197
8.4	Monitorování sítě	200
8.4.1	Protokol RMON	200
8.4.2	Snort	201
8.4.3	NetFlow	202
8.5	VPN	204
8.5.1	IPSec VPN – na síťové vrstvě	206
8.5.2	GRE tunely	207
8.5.3	L2TP tunely	208
8.5.4	OpenVPN	209
8.5.5	SSH	209
8.5.6	MPLS VPN	210
8.5.7	Další možnosti řešení VPN	210
9	Bezpečnost a správa	212
9.1	Pojmy z oblasti kyberbezpečnosti	212
9.2	Kybernetické útoky	213
9.2.1	Průzkumné útoky	213
9.2.2	Přístupové útoky	214
9.2.3	Útoky způsobující odmítnutí služby	216
9.3	České organizace	217
9.3.1	CZ.NIC	217
9.3.2	CESNET2	217
9.3.3	Bezpečnostní týmy	218
9.4	Bezpečnostní zařízení v síti	220
9.4.1	Monitorování provozu	220
9.4.2	Firewall	222
9.4.3	Filtrování	224
9.4.4	IDS/IPS a hloubková inspekce paketů	225
9.5	Dohledové systémy	228
9.5.1	Nagios	228
9.5.2	Zabbix	229
9.5.3	OpenNMS	229
9.5.4	Zenoss	229
9.5.5	Cacti	230
9.6	SIEM	231
	Seznam doporučené literatury	234

Přílohy	236
A Práce s adresami a sítěmi ve Windows	237
A.1 Problémy při používání příkazů ve Windows	237
A.2 Soubory související se sítěmi	237
A.3 Práce s adresami a síťovými kartami	238
A.3.1 Základní práce s IP adresami <code>--ipconfig</code>	238
A.3.2 Překlad adres <code>--arp</code> a <code>nslookup</code>	240
A.3.3 Směrování – <code>route</code>	242
A.3.4 Malá síť a skupina (workgroup)	244
A.4 Testování a statistiky	246
A.4.1 Testování cesty a průchodnosti	246
A.4.2 Statistiky v <code>netstatu</code>	247
A.5 Služba WMI a program <code>wmic</code>	249
A.6 Firewall ve Windows	252
A.7 Další příkazy	253
B Práce s adresami a sítěmi v Linuxu	255
B.1 Specializované distribuce	255
B.2 Soubory související se sítěmi	257
B.3 Starší příkazy pro práci s adresami	259
B.4 Mechanismus <code>iproute2</code> (příkaz <code>ip</code>) – adresy, síť, směrování	261
B.4.1 Konfigurace síťového rozhraní a adres	261
B.4.2 Směrování a filtrování	263
B.4.3 Objevování sítě	268
B.4.4 Tunely	270
B.5 Překlad názvů	273
B.5.1 Název počítače	273
B.5.2 Resolver a soubor <code>resolv.conf</code>	273
B.5.3 Testování DNS	274
B.5.4 Zjistění informací o doméně	277
B.6 Testování a statistiky	278
B.6.1 Základní nástroje	278
B.6.2 Pokročilé testování	280
B.7 Firewall v Linuxu	280
B.7.1 Princip firewallu	281
B.7.2 Základní vnitřní příkazy a parametry pro filtrování	282
B.7.3 SPI	287
B.7.4 NAT a maškaráda	289
B.7.5 Označování paketů	291
B.7.6 Skripty	293
B.7.7 Uložení změn	293
B.8 Vzdálený přístup	294
B.9 Další příkazy	296

Pojmy a standardy k počítačovým sítím

 **Rychlý náhled:** V této kapitole si pouze ujasníme či připomeneme pojmy, které bychom už měli znát z bakalářského studia. Připomeneme si základní koncepty a pojmy, jak probíhá komunikace v referenčním modelu ISO/OSI a TCP/IP, jak jsou zpracovávána data na různých úrovních této komunikace.

Dále budeme používat pojem počítačová síť, ale to je nepřesné, ve skutečnosti zde ani zdaleka nepůjde pouze o počítače (nicméně ten pojem je poměrně vžitý).

 **Klíčová slova:** Počítačová síť, zařízení, aktivní a pasivní prvek, repeater, hub, bridge, switch, router, gateway, protokol, standard, RFC dokument, PDU, ISO/OSI, TCP/IP, SDU, SAP, entita, komunikační primitivum, soket, protokolová sada, protokolový zásobník.

 **Cíle studia:** Cílem této kapitoly je připomenout si základní koncepty a postupy v počítačových sítích, včetně nejdůležitějších protokolů, síťových zařízení, referenčních, síťových a protokolových modelů. Cílem je porozumět principům standardizace počítačových sítí a porozumět veřejně dostupným dokumentům popisujícím standardy.

1.1 Přehled základních pojmů a postupů

 V síti máme samozřejmě zařízení nacházející se na začátku nebo konci komunikační cesty – koncová zařízení. Jsou to počítače, servery, tablety, mobily, síťové tiskárny, zařízení internetu věcí a další. Dále do počítačové sítě patří síťové prvky, přes které komunikace probíhá. Rozlišujeme

- *aktivní síťové prvky* – aktivně ovlivňují komunikaci, směřují, zesilují signál apod.,
- *pasivní síťové prvky* – pouze „pasivně přenesou“ data, například kabeláž.

Definice (Počítačová síť)

Počítačová síť je soustava zařízení – *uzlů* (node) vzájemně propojených *přenosovými cestami* (transmission path). Uzel v počítačové síti je buď koncové zařízení nebo aktivní síťový prvek.





Následuje přehled pojmů, o kterých bude v dalším textu předpokládáno, že je čtenář zná a chápe.

- spoj (link), přenosový kanál (channel), přenosový okruh (circuit), pronajatý okruh (leased circuit),
- poskytovatel síťové konektivity (ISP – Internet Service Provider),
- přenosový signál, vlnová délka, frekvence, amplituda, šířka pásma,
- fyzický/pevný okruh (physical circuit), virtuální okruh (virtual circuit),
 - pevný virtuální okruh (PVC, permanent virtual circuit)
 - přepínaný virtuální okruh (SVC, switched virtual circuit)
- sítě PAN, LAN, MAN, WAN,
- přenos simplexní, plně duplexní (full duplex), poloduplexní (half duplex),
- přenos proudový (stream) a paketový (blokový, block, packet),
- přepojování okruhů (circuit switching), přepojování paketů (packet switching),
- přenos spojovaný (služba se spojením – connection-oriented) nebo nespojovaný (connectionless),
- datagramová služba,
- přenos spolehlivý (reliable) nebo nespolehlivý (unreliable), best effort (nejlepší snaha),
- komunikace typu unicast, multicast, broadcast, anycast,
- fyzická a logická topologie, druhy: sběrnice, kruh, hvězda, strom, mesh, point-to-point,
- páteřní vedení,
- síťové rozhraní,
- referenční model ISO/OSI, síťový model TCP/IP (obojí bude ještě dále připomenuto),
- jednotlivé typy kabelů (twisted pair, optika, koaxiál, twin-ax) – typické vlastnosti, použití,
- u twisted pair kategorie, nestíněný (UTP), různé druhy stíněného, kabel typu lanko (licna) vs. drát.



Taktéž se předpokládá, že čtenář zvládá minimálně následující činnosti:

- jak se pracuje v binární a hexadecimální číselné soustavě,
- jak zjistit svou MAC a IP adresu v grafickém i textovém režimu,
- jak nastavit svou IP adresu, masku, adresu brány, adresu DNS serveru,
- subnetting – počítání rozsahů adres, alespoň pro IPv4,
- jak se používá Wireshark (<https://www.wireshark.org/download.html>),
- jak se krimpuje kabel typu twisted pair na RJ-45 (8p8c) včetně testování,
- jak se používá buď Packet Tracer nebo jiný obdobný program pro simulaci sítě a její konfigurace,
- jak se zjišťuje dostupnost zařízení na síti a cesta k němu (ping, atd.),
- jak se dá zjistit, které procesy naslouchají na některém portu, a jak se vypisují různé statistiky související se sítí,
- jak u Wi-fi zjistíte obsazenost kanálů a informace o AP.

1.2 Aktivní síťové prvky

Následuje přehled běžných aktivních síťových zařízení. Ke každému je taktéž přidána informace o tom, na kterou vrstvu ISO/OSI patří.



Repeater (opakovač) je jednoduchý aktivní síťový prvek se dvěma porty. Příchozí signál zesílí (případně znovu vygeneruje) a pošle dál. Pracuje na vrstvě L1.

Dnes se s repeaterem (příp. extenderem) setkáváme spíše u bezdrátových lokálních sítí, ale také např. HDMI, kde slouží ke zvýšení dosahu signálu.



Hub (rozbočovač) je jednoduchý aktivní síťový prvek, který má typicky více než dva porty.

Cokoliv přijde na některý port, jen přepoše na všechny ostatní porty (tento signál zesílí nebo znovu vygeneruje). Pracuje na vrstvě L1.

Jeho výhodou je rychlost, nevýhodou je, že zbytečně zahlcuje síť (paket pošle všem kromě odesílatele). V běžných počítačových sítích se už dnes s huby moc nesetkáváme, používají se ale například v některých PAN sítích nebo u USB (ale s trochu jinými charakteristikami).



Switch (přepínač) je již inteligentnější aktivní síťový prvek, který si vede tabulku adres uzlů.

U příchozího paketu zjistí adresáta, podle tabulky určí port, který k němu vede, a paket odešle jen na tento port. Pokud adresáta nemá v tabulce nebo když jde o broadcast, odešle ho na všechny porty. Switch rozděluje síť na segmenty – pokud mezi sebou komunikují uzly ze stejného segmentu, pak taková komunikace neprotrvá do jiného segmentu. Pracuje na vrstvě L2, ale ve skutečnosti může obsahovat i funkcionality vyšších vrstev (multilayer switch).

Výhodou je, že tolik nezahluje síť, také je možné implementovat správní a bezpečnostní funkce. Nevýhodou je výpočetně náročnější provoz (což je technicky řešitelné, velká část funkčnosti je implementována hardwarově). Se switchi se běžně setkáváme v LAN, MAN i WAN sítích.



Bridge (most) je zařízení navzájem oddělující dva segmenty sítě, jako switch. Na rozdíl od switche má méně portů, jeho funkčnost je implementována softwarově a je standardizována (IEEE 802.1).

V praxi se přímo s bridgi nesetkáváme. Most pracuje na vrstvě L2.



Router (směrovač) si vede směrovací tabulku, v ní však nemá adresy konkrétních uzlů, ale L3 adresy sítí a podsítí. Příchozí paket „převrhne“ o něco důkladněji než switch, přičemž podle cílové

adresy zjistí, do které sítě patří adresát, ve směrovací tabulce ověří, který port je cestou do této sítě, a na ten port paket odešle. Broadcasty zahazuje.

Zatímco switch odděluje segmenty sítě, router odděluje různé sítě. Routery pracují na vrstvě L3.

Výhodou routeru je možnost implementace pokročilých správních a bezpečnostních funkcí a schopnost oddělit komunikaci v různých sítích. Nevýhodou je výpočetně náročnější provoz, proto routery mají méně portů než switch, výkonnější procesor a více paměti, aby vyšší výpočetní zátěž unesly.



Brána (gateway) slouží k propojení dvou *různých typů* sítí, resp. slouží jako spojovací bod a zároveň „překladač“. Například pokud máme v domácnosti VDSL router, pak toto zařízení má vestavěnou bránu pro komunikaci mezi naší lokální sítí a VDSL sítí poskytovatele internetu.

Brána typicky pracuje na vyšší vrstvě než protokoly, které má propojit.

1.3 Co a jak přenášíme po počítačové síti

1.3.1 Protokoly

Protokol určuje, jakým způsobem má komunikace začít, případně jak se dohodnout na určitých parametrech komunikace, jak sdělit druhé straně určitý typ informace, jak má druhá strana potvrdit příjem nebo sdělit, že je třeba přenos opakovat, jak se komunikace ukončuje, atd.



Definice (Protokol a jeho implementace)

Protokol je konvence, podle které probíhá určitý typ (většinou elektronické) komunikace. Definuje pravidla určující syntaxi (jak jsou které signály poskládány), sémantiku (význam) a synchronizaci komunikace.

Protokol je pouze předpis, který je třeba implementovat (naprogramovat, aby mohl být v praxi používán). Implementace může být softwarová, hardwarová (v obvodech) nebo kombinace softwarové a hardwarové.



Například protokol HTTP je implementován (naprogramován) v operačním systému běžícím na počítači, u kterého sedíte, a taky na druhé straně – v operačním systému WWW serveru, se kterým komunikujete.



Poznámka:

Pro návrh protokolu platí jedno důležité pravidlo (vlastně je zachováváno i jinde, například u systému UNIX): protokol by měl být jednoduchý, krátký, jednoznačně implementovatelný. Neměl by být moc složitý, protože čím větší složitost, tím větší pravděpodobnost chyb. Proto žádný protokol není moc univerzální – umí jednu konkrétní věc, a umí ji dobře. V angličtině se to vyjadřuje zkratkou KISS (Keep it Simple, Stupid – ať je to jednoduché, hloupé).

Aby tento princip mohl být dodržován, existují určité konvence pro spolupráci protokolů: to, co protokol neumí, předá ke zpracování jinému protokolu.



Princip protokolů není ani zdaleka jen záležitostí počítačových sítí – protokoly se používají v telekomunikacích, spotřební elektronice, strojírenství, ale například i „uvnitř“ počítače. Každý se setkal třeba s protokolem USB, SATA,...

1.3.2 Vlastnosti protokolů

Běžně používané protokoly mají obvykle (kromě jiných) tyto tři důležité vlastnosti:

- jsou všeobecně známé a (téměř) kdokoli je může implementovat (kromě proprietárních),
- jsou spíše jednoduché, nepříliš komplexní,
- mohou spolupracovat s jinými protokoly.

Nejdřív se zaměříme na první vlastnost. Proč je důležitá? Představme si situaci, kdy výrobce síťového hardwaru přijde s novým zařízením a rozhodne se, že toto zařízení bude komunikovat podle nového protokolu, jehož specifikaci nikomu jinému neposkytne. Toto zařízení si ovšem bude „rozumět“ jen s takovými zařízeními, která budou podporovat tentýž protokol, takže nedokáže komunikovat se zařízeními od jiných výrobců. Pro dotyčného výrobce by to snad mohlo být výhodné, pokud by dokázal své potenciální zákazníky přesvědčit, aby kupovali jenom od něj, ale realita bývá jiná – zákazníci by radši kupovali taková zařízení, která by mohli bez problémů zařadit do své sítě, ve které už pravděpodobně mají nějaká zařízení od jiných výrobců.

Proto je většina protokolů buď volně dostupná (specifikace je zcela volně k nahlédnutí, implementovat může kdokoli) nebo alespoň dostupná jiným způsobem (za poplatek, udělení licence).

Otevřený protokol je protokol volně dostupný, naopak proprietární protokol je takový protokol, jehož specifikace není nikde zveřejněna a jeho tvůrce si ji buď nechává jen pro sebe nebo licencuje za poplatek.

**Poznámka:**

O rozšířenosti volných specifikací mluví například to, že momentálně je na routerech nejběžnějším směrovacím protokolem otevřený protokol OSPF. Naopak IGRP (proprietární směrovací protokol společnosti Cisco) se téměř nepoužívá dokonce ani na zařízeních od samotného Cisca.



Pokud síťové zařízení podporuje některý proprietární protokol, obvykle pro danou funkci taky podporuje některý alternativní protokol s dostupnější specifikací, aby si mohlo „popovídat“ se zařízeními jiných výrobců.

Dále se k principu otevřenosti, jednoduchosti a spolupráce budeme často vracet.

1.4 Standardy a standardizační instituce

**Definice (Standard, norma a standardizace)**

Standard (v oblasti technologií) je požadavek na splnění určitých konkrétních vlastností pro určitý typ hardwaru, softwaru apod. Je to dokument popisující požadavky, specifikace, návody a popisy pro daný produkt, proces či službu. Rozlišujeme standardy

- *normativní (de iure, podle zákona)* – jejich dodržování je vyžadováno, obvykle je stanoví příslušná státní instituce,
- *popisné (de facto)* – jejich dodržování sice není striktně vyžadováno, ale je v zájmu výrobce.

V oblasti technologií obvykle používáme pojem *norma* pro normativní standardy, jinak prostě použijeme pojem *standard*, třebaže se taky často mluví o *doporučeních* (recommendation).

Standardizace je proces sjednocení vlastností a funkcí daného produktu pomocí stanovení standardu.



1.4.1 Standardizační instituce pro počítačové sítě

Následuje přehled nejdůležitějších standardizačních organizací a institucí, které vydávají standardy související s počítačovými sítěmi a technologiemi obecně. Většinou se jedná o standardy *de facto* (nejsou závazné, ale běžně bývají dodržovány), samozřejmě až náš vlastní (a částečně i evropský) normalizační institut.

Situace se mírně komplikuje tím, že některé standardy vydané jednou organizací bývají často adaptovány jednou či dokonce více dalšími organizacemi (například je běžná adaptace standardů českým normalizačním úřadem). Původce standardu obvykle poznáme podle začátku označení tohoto standardu, kombinací více organizací pak díky kombinaci jejich označení.



Úřad pro technickou normalizaci, metrologii a státní zkušebnictví je národním standardizačním orgánem pro Českou republiku, zabývá se tvorbou českých norem. Normy vydané tímto ústavem poznáme podle toho, že začínají písmeny ČSN, za nimi následuje číslo normy. V současné době ve značné míře zabývá adaptací (přejmutím a přizpůsobením) norem vydaných Evropskou unií nebo některou standardizační organizací. Tyto normy poznáme podle toho, že za ČSN následuje zkratka původní organizace či instituce, například ČSN EN... jsou normy přejaté od Evropské unie.

**Další informace:**

Normy vydané Úřadem pro technickou normalizaci, metrologii a státní zkušebnictví jsou k nahlédnutí na <http://www.unmz.cz/urad/csn-online> (jen náhledy, za celé znění se platí).



ETSI (European Telecommunications Standard Institute) je sice původně evropská organizace, ale ve skutečnosti jsou její členové ze všech obydlených kontinentů (státy, významní výrobci komunikačních zařízení, poskytovatelé síťových služeb, výzkumné organizace, atd.). Náměty na nové standardy nebo změny stávajících vycházejí ze tří zdrojů – buď se domluví nejméně čtyři členové ETSI, nebo je přijat podnět od Evropské komise nebo od EFTA (European Free Trade Association).

Rozlišuje se několik typů ETSI standardů – například EN (European Standard), ES (ETSI Standard), TS (ETSI Technical Specification) a další. K nejznámějším patří standardy související s mobilními sítěmi (především GSM, 3G, 4G sítě), chytrými sítěmi, komunikací machine-to-machine, ICT ve zdravotnictví, atd. ETSI standardy jsou dostupné volně bez poplatků.



<https://portal.etsi.org>



ITU (International Telecommunications Union) je celosvětová organizace pro informační a komunikační technologie. Je součástí OSN a sídlí ve Švýcarsku. Má tři části:

- ITU-R: radiokomunikační sektor (původně CCIR),
- ITU-T: sektor pro standardizaci telekomunikací (původně CCITT),
- ITU-D: sektor rozvoje telekomunikací.

ITU-T je součástí ITU pro standardizaci telekomunikací, její činnost souvisí i s počítačovými sítěmi. Členy s hlasovacím právem jsou zainteresované země, členem bez hlasovacího práva může být kterákoliv organizace. Členství je placené.

ITU-T se dále člení do *studijních skupin* (SG, Study Group), které mají každá své vlastní zaměření. Například SG13 v poslední době vydala několik standardů o cloud computingu, SG16 se zabývá multimédií (přenos videa, obrázky, zvuk apod.), SG17 bezpečností, SG20 chytrými sítěmi (Internet of Things, ...). Standardy vytvořené ITU-T jsou většinou otevřené, volně dostupné (kromě některých nejnovějších verzí).

Se standardy ITU-T se setkáváme poměrně běžně: například má „na svědomí“ protokol H.323 používaný v multimediálních přenosech a internetové telefonii, protokoly G.992.1 and G.992.2 pro přístupové sítě ADSL, standard X.509 používaný při šifrování a další.



<http://www.itu.int/en/ITU-T/Pages/default.aspx>

ITU-R je součástí ITU pro radiokomunikace, a tedy se zabývá standardy pro bezdrátové vysílání (také datové). Momentálně se hodně probírají například standardy pro LTE Advanced (formálně IMT Advanced), jako je M.2012 – *Detailed specifications of the terrestrial radio interfaces of International Mobile Telecommunications Advanced (IMT-Advanced)*.

Standardy začínající písmenem „M“ jsou pro mobilní radiokomunikační sítě, „RS“ pro bezdrátové sítě senzorů a jejich vzdálené ovládání, „SM“ pro správu frekvenčního spektra, atd.



<http://www.itu.int/pub/R-REC>



ISO (International Organization for Standardization) je nezávislá organizace vydávající standardy pro komunikační technologie. Jejími členy jsou standardizační instituce z různých zemí (každá země tam má jedno zastoupení), naopak organizace ISO je členem ITU. Sídlo ISO je ve Švýcarsku.

ISO má hierarchickou strukturu. Člení se na 200 TC (Technical Committee, technický výbor), každý TC má svou oblast působnosti. Například ISO/TC 47 se zabývá chemií, ISO/TC 20 letadly a vesmírnými vozidly, ISO/TC 272 forenzními vědami, ISO/TC 299 robotikou, ISO/IEC JTC 1 informačními technologiemi.

Každý TC se člení na subkomise (SC, Subcommittee) a každá subkomise se člení na pracovní skupiny (WG, Working Group). Návrhy na standardy nebo jejich změnu jdou vždy zespoda, od pracovních skupin, postupně se přepracovávají a „probojovávají“ nahoru až ke schválení standardu. Pracovní skupina vytvoří *pracovní verzi* (Working Paper), ten je přepracován na *koncept* výboru (Committee Draft), následuje *koncept mezinárodního standardu* (Draft International Standard), posledním stupněm je *mezinárodní standard* (International Standard).

Nejznámějším standardem, se kterým se seznámíme už na následujících stranách, je referenční model ISO/OSI (standard ISO 7498 a hodně dalších s ním souvisejících), dále se ISO zabývá sítěmi senzorů, vzdáleným přístupem, UPnP, webovými službami a dalšími tématy.

 http://www.iso.org/iso/home/standards_development/list_of_iso_technical_committees.htm

 **IEEE** (Institute of Electrical and Electronics Engineers, čteme anglicky „Eye-triple-E“, tedy [aj tripli i:]) je největší světová organizace sdružující odborníky z oblasti elektroniky, elektrotechniky, informatiky a souvisejících oborů, je to tedy profesní sdružení. Nezabývá se jen standardy, ale pořádá různá setkání, konference, vzdělávací akce, vydává odborné časopisy a vyvíjí další aktivity. Sídlo IEEE je v USA. Jednou z aktivních sekcí IEEE je i ta česká.

Sdružení IEEE stojí především za skupinou standardů IEEE 802 pro lokální a částečně i rozlehlé sítě, do které patří například IEEE 802.11 pro bezdrátové lokální sítě (Wi-fi) nebo IEEE 802.3 pro Ethernet. Přístup k plnému znění standardů je placený.

 <http://www.ieee.org>

 **IETF** (The Internet Engineering Task Force) se zabývá především standardy souvisejícími s Internetem, včetně například protokolů TCP/IP. Velmi úzce spolupracuje s dalšími organizacemi, například IANA (The Internet Assigned Numbers Authority), ISOC (Internet Society), IAB (Internet Architecture Board, Rada pro architekturu Internetu), W3C (World Wide Web Consortium, standardy pro web) a dalšími.

Standardy vydané IETF obvykle začínají zkratkou RFC (Request for Comments) a následuje číslo. Pokud je třeba standard aktualizovat, nemění se původní znění, ale vytvoří se RFC dokument (obsahující popis standardu) s novým číslem. Například pro výše zmíněný otevřený směrovací protokol OSPF bylo takto postupně vytvořeno několik dokumentů (verzí), z nichž jsou momentálně aktuální RFC 2328 (OSPF verze 2) a RFC 5340 (OSPF verze 3).

Další informace:

Oficiální stránky organizace jsou na <https://www.ietf.org/>. Všechny dokumenty RFC jsou volně dostupné, dokonce existuje řada webových stránek, které je zpřístupňují, například:

- <https://tools.ietf.org/html/>
- <https://datatracker.ietf.org/doc/>
- <http://www.rfc-base.org/>
- <https://www.rfc-editor.org/>



 **IEC** (International Electrotechnical Commission) se z oblasti počítačových sítí zabývá především standardy pro elektrická a elektronická zařízení, spolupracuje především s organizací ISO a hodně standardů nese označení „ISO IEC“.

 <http://www.iec.ch/>

 **ANSI** (American National Standards Institute) je normalizační úřad USA. Sdružuje státní instituce, komerční společnosti, členy z akademické sféry a další. Kromě vyvíjení vlastních standardů také reprezentuje USA v ISO, IEC a dalších nadnárodních organizacích. Standardy jsou dostupné za poplatek.

ANSI stojí například za standardem pro znakovou sadu ASCII, známý je také standard pro programovací jazyk C (ANSI C), v oblasti sítí například FDDI, CDMA, Frame Relay.

 <https://www.ansi.org/>

 **TIA a EIA** (Telecommunication Industries Association, Electronic Industries Alliance) jsou americké organizace zaměřující se především na fyzickou úroveň komunikace mezi zařízeními, v čemž hodně spolupracují také s ANSI.

TIA se zabývá standardy pro nejrůznější typy kabelů a konektorů, připojení antén, mobilní sítě, ICT ve zdravotnictví, chytré sítě. EIA stojí například za starým známým konektorem SCART, dále v sítích se setkáváme s konektorem RS-232. Existují standardy vzniklé při spolupráci těchto společností, například TIA/EIA 568, se kterým se setkáme v kapitole o Ethernetu.

 **Další informace:**

- <http://www.tiaonline.org/>
- https://www.eia.gov/about/eia_standards.cfm



1.4.2 Volně dostupné standardy

 **RFC dokumenty.** Pro fungování (nejen) Internetu jsou důležité také RFC dokumenty, ve kterých jsou popsány nejdůležitější protokoly a způsob jejich spolupráce. Připomeňme si základní vlastnosti těchto dokumentů:

- Každý RFC má jednoznačné číslo a také svůj název.
- Neexistují žádné aktualizace RFC dokumentů, aktualizace obsahu je vydána pod novým číslem (ale název bývá stejný).
- Jedno téma (případně protokol) může být popsáno ve více než jednom RFC dokumentu.

Výhodou druhé vlastnosti je, že se nemusíme ztrácet v různých verzích téhož dokumentu, nevýhodou je, že hledání aktuálního znění je o něco těžší. Výhodou třetí vlastnosti je, že RFC dokumenty nejsou obvykle až tak dlouhé, aby ztrácely na přehlednosti, nevýhodou je, že někdy musíme dohledávat související informace.

 **Příklad**

Podíváme se na některé RFC dokumenty související s protokolem TCP. Tento protokol pracuje na transportní vrstvě a jeho úkolem je navázat a spravovat spojení se strojem, s nímž komunikuje náš stroj.

Na adrese <https://tools.ietf.org/html/> nejdřív zadáme do vyhledávání číslo RFC dokumentu 7414. Tento dokument nepopisuje přímo protokol TCP, ale je jakýmsi rozcestníkem k RFC dokumentům, které s TCP mají něco společného.

Aby se předešlo problémům s formáty, jsou RFC dokumenty čistě textové. Všimněte si struktury dokumentu:

- Na začátku je krátké informační záhlaví dokumentu, pak název („A Roadmap for Transmission. . .“), následuje abstrakt se stručným popisem, stav dokumentu a licenční ujednání.
- Dále tu máme obsah a za ním kapitoly dokumentu. Poslední části jsou odkazy na další zdroje (většinou další RFC) a kontakty na autory.
- Třebaže si prohlédneme „dlouhou“ webovou stránku, je dokument rozčleněn na stejně dlouhé stránky a každá má své záhlaví a zápatí (mezi zápatím předchozí strany a záhlavím následující je čára).

Zaměřme se na informační záhlaví dokumentu. Vlevo se dozvíme, že se jedná o produkt IETF a je to RFC číslo 7414, na dalším řádku stojí „**Obsoletes: 4614**“. To znamená, že předchozí varianta tohoto dokumentu (zastaralý, obsolete) má číslo 4614. Všimněte si, že čím novější, tím vyšší číslo. Ve sloupci vpravo zjistíme, kdy tento RFC dokument vznikl: v únoru 2015.

V úvodní kapitole dokumentu je obvykle povídání o obsahu dokumentu a jakýsi souhrn. V další kapitole – Core Functionality – najdeme odkaz na hlavní dokument obsahující popis protokolu TCP, tedy RFC 793. Všimněte si nízkého čísla – tento dokument je platný už dlouho (od roku 1981), ale další RFC přidávají novou funkcionalitu.

Na to, že jde vlastně jen o komentovaný seznam RFC dokumentů souvisejících s protokolem TCP, je dokument opravdu hodně dlouhý. Ovšem celý dokument studovat nebudeme.



Příklad

V minulém příkladu jsme se dozvěděli, že RFC 7414 je novější variantou nahrazující zastaralý (obsolete) dokument RFC 4614. Podívejme se na tento starší dokument (stejným způsobem – na adrese <https://tools.ietf.org/html/> zadáme do vyhledávacího okna číslo 4614, nebo prostě využijeme odkaz z novějšího dokumentu).

Vidíme, že název dokumentu („A Roadmap for Transmission. . .“) je stejný jako u novějšího, jen číslo je jiné. V levém sloupci záhlaví dokumentu je řádek „**Obsoleted by: 7414**“. Tento řádek je důležitý, protože pokud narazíme na RFC dokument, o jehož platnosti nic nevíme, tento řádek nám řekne, že je zastaralý a který dokument je jeho náhradou.

Dále je řádek „**Updated by: 6247**“. Nejedná se o novou verzi, pouze se mění „vztah k okolí“, v tomto konkrétním případě jsou některá nepoužívaná rozšíření protokolu TCP „odsunuta do historie“.



Příklad

Teď se podívejme na „Core“ RFC dokument o protokolu TCP. Z prvního příkladu víme, že jeho číslo je 793. Začátek dokumentu vypadá trochu jinak: struktura RFC dokumentů se postupně obohacovala.

Všimněme si nákresů – všechny jsou čistě textové, nicméně pro tyto účely to dostačuje. V kapitole 1.1 je hrubý nákres vzhledem k okolním vrstvám (jak vidíme, pod TCP má být internetový protokol, tedy IP), a dále v kapitole 2.5 na straně 9 je podrobnější nákres zahrnující konkrétní spolupracující protokoly

(samozřejmě ne všechny). V kapitole 3.1 je náčrtek záhlaví TCP segmentu, za nímž jsou všechny součásti vysvětleny. Na straně 23 najdeme stavový diagram popisující komunikaci s využitím protokolu TCP.



Úkol

Prohlédněte si RFC 2460 a odpovězte na tyto otázky:

- Co tento dokument popisuje? Kdy byl publikován?
- Nahrazuje některý zastaralý (obsolete) dokument?
- Najděte kapitolu 2 o terminologii. Které pojmy vám něco říkají?
- Prohlédněte si náčrty v dokumentu uvedené.



 **ITU-T.** Další organizací, která své standardy celé zveřejňuje, je ITU-T. Informace jsou dostupné na webu <http://www.itu.int/en/ITU-T/Pages/default.aspx>.



Příklad

Jedním ze standardů, které vznikly v ITU-T, je X.500. Je věnován adresářovým službám, což jsou síťové služby, které mají zjednodušovat správu sítí – v distribuované (rozkládající se na více zařízeních) databázi jsou evidovány jak jednotlivé objekty v síti (počítače, servery, prostředky na nich dostupné, aktivní prvky, atd.), tak i uživatelé a jejich přístupová práva k těmto objektům. Odlehčenou variantou protokolu X.500 je protokol LDAP, který je implementován jak ve Windows (jako Active Directory) tak i v Linuxu a dalších operačních systémech.

Na adrese <http://www.itu.int/itu-t/recommendations/rec.aspx?rec=X.500> najdeme základní informace o protokolu a také odkaz na PDF soubor s celým zněním (stačí klepnout na červenou ikonu vpravo).



Úkol

Prohlédněte si obsah PDF souboru se specifikací X.500, o který se jednalo v předchozím příkladu, zejména:

- Srovnajte strukturu dokumentu s RFC dokumenty.
- Na začátku kapitoly 6 (Overview of the Directory) si přečtěte, co je to Directory (adresář). Všimněte si, že zde je tento pojem chápán trochu jinak než v běžných operačních systémech (obdobu složky).
- Všimněte si náčrsků. Většinou jde o náčrty komunikace s databází nebo vztahové diagramy (stromy). Na obrázku 3 (tisknutá strana 7, podle pořadí 13) je ukázková struktura adresářového stromu podle X.500, pokuste se pochopit vztahy mezi uzly tohoto stromu.

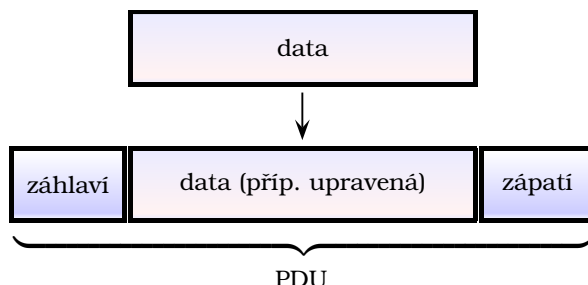


1.5 Referenční model ISO/OSI

1.5.1 Přehled modelu

 *Protokolová datová jednotka* (PDU, Protocol Data Unit) je sekvence dat opatřená metadaty (informacemi o datech) vztahující se ke konkrétnímu protokolu. Protokoly obvykle obdrží data, podle potřeby

je určitým způsobem zpracují (strukturují, rozdělí na menší části, zašifrují, komprimují, přeloží, určí adresu příjemce apod.) a přidají před ně *záhlaví* (header) se souvztažnými informacemi (délka dat, použitý šifrovací algoritmus, adresa odesílatele a příjemce, atd.). Některé protokoly také přidávají za data *zápatí* (trailer) obsahující například kontrolní součet. Data přenášená v datové jednotce taky nazýváme *payload*.



Obrázek 1.1: Datová jednotka (PDU)

 K nejdůležitějším standardům z oblasti počítačových sítí patří skupina standardů popisujících *referenční model OSI* (Open Systems Interconnection) publikovaných organizací ISO, proto se označuje RM ISO/OSI. Původní označení standardu je ISO/IEC 7498-1, dnes je dostupný jako ITU-T X.200 na webu <http://www.itu.int/rec/T-REC-X.200-199407-1>.

OSI definuje sedm vrstev. Každá vrstva plní v komunikaci přes počítačovou síť konkrétní funkci a je přesně stanoveno, co se na dané vrstvě může dít, jedná se tedy o *konceptuální model* (tj. popisuje logiku návrhu, vztahy mezi součástmi).

	Referenční model ISO/OSI	Význam:	PDU:	Adresování:
L7	Aplikační vrstva	Poskytuje služby aplikacím – manipule s daty, určení jejich struktury, sémantické překlady, bezpečnost.	data, zprávy	
L6	Prezentační vrstva	Provádí konverze dat jako je šifrování, (de)komprese, konverze do/z jiného datového formátu, ...		
L5	Relační vrstva	Otevírá, řídí and ukončuje konverzace mezi dvěma vzdálenými aplikacemi, odděluje data různých aplikací.		
L4	Transportní vrstva	Zabezpečuje spojení mezi dvěma koncovými body; segmentace proudu dat před přenosem, skládání po přenosu.	segmenty	porty
L3	Síťová vrstva	Přeposílá data k danému cíli, provádí směrování, pracuje s logickou topologií sítě.	pakety, datagramy	IP adresy
L2	Linková vrstva	Pracuje s fyzickou topologií sítě, synchronizuje přenos, zde obvykle pracují LAN řešení.	rámce	MAC adresy
L1	Fyzická vrstva	Řídí proces posílání a přijímání proudu dat, definuje fyzikální a elektrické specifikace zařízení (rozhnutí).	bity	

Obrázek 1.2: referenční model ISO/OSI

Pořadí vrstev a stručný popis najdeme na obrázku 1.2. Vrstvy zobrazené zeleně (L1 až L3) jsou závislé na přenosovém médiu, vrstva L4 (modrá) je přechodová, vrstvy zobrazené žlutě (L5 až L7) jsou nezávislé na přenosovém médiu a na samotném přenosu se podílejí jen nepřímo (přípravou dat a komunikací s aplikacemi).

 **Fyzická vrstva (L1, Physical Layer).** Vrstva L1 je odpovědná za samotný fyzický přenos dat. Definuje přenosové médium (kabely, rádiové vlny apod.), jak je reprezentován bit o hodnotě 0 nebo 1 (kódování), síťové rozhraní (porty, konektory), k čemu konkrétně je používán který vodič v kabelu, synchronizace a obecně vše, co je potřeba pro konverzi sekvence bitů do formy přenášeného signálu. Fyzická vrstva „vidí“ pouze své přímé sousedy (na druhém konci kabelu, v dosahu bezdrátového signálu).

Fyzickou vrstvu mají implementovanou všechna zařízení v počítačové síti, tedy všechna, která jsou opatřena nějakým síťovým rozhraním.

Zařízení, která mají implementovanou pouze fyzickou vrstvu: hub a repeater. Výhodou těchto zařízení je jednoduchost a rychlost, nevýhodou je nemožnost implementovat pokročilejší funkce.

 **Linková vrstva (L2, Data-link Layer, spojová).** Na této vrstvě se určuje vztah mezi přichozím proudem bitů (který vidí fyzická vrstva) a konkrétním uzlem v síti. Zařízení s implementovanou vrstvou L2 má přehled o zařízeních připojených do místní sítě (minimálně „vidí“ ta zařízení, ke kterým je přímo připojeno), vede si tabulku fyzických adres těchto zařízení (nazývá se obvykle MAC tabulka, CAM tabulka apod., podle konkrétního protokolu a konkrétního zařízení). Ke každé adrese v tabulce máme kromě jiného i port, přes který je dotyčné zařízení dosažitelné.

Na linkové vrstvě jsou také zajištěny funkce, které se sice vztahují k přenosovému médiu, ale zároveň vyžadují práci s datovými jednotkami. Například se zde určuje rychlost přenosu, protože tu je třeba řídit i podle toho, v jakém stavu je příjem datových jednotek. Pokud se sem tam nějaká datová jednotka ztratí (což by fyzická vrstva nepoznala), pak zřejmě celý mechanismus „nestíhá“ a je třeba snížit rychlost přenosu. K funkcím vrstvy L2 tedy patří i detekce přenosových chyb.

Datové jednotky, se kterými pracují protokoly této vrstvy, se nazývají *rámce*. Každý rámec jednoznačně označuje začátek a konec dat a obsahuje (kromě jiného) fyzickou adresu příjemce a fyzickou adresu odesílatele v rámci sítě.

Vrstvu L2 implementují ta zařízení v síti, která potřebují přehled o uzlech v místní síti a pracují s adresami těchto uzlů (tj. tato vrstva „vidí“ různé uzly v síti, nejen své bezprostřední sousedy). Switche a bridge implementují vrstvu L2, přičemž zde pracuje vždy některý protokol, který si vede tabulku fyzických adres. Ovšem existují i switche implementující vrstvu L3, ale to už je dodatečná funkcionalita.

Pokud se jedná o aktivní síťový prvek, pak takové zařízení dokáže *oddělovat segmenty v síti*.

 **Síťová vrstva (L3, Network Layer).** Zatímco protokoly linkové vrstvy mají přehled o fyzické topologii sítě, protokoly síťové vrstvy pracují s logickou topologií sítě (udržují si přehled o hierarchické struktuře sítě) a „vidí“ i za hranice lokální sítě. Úkolem síťových protokolů je stanovit skutečnou cestu (nebo úsek cesty, za který je dotyčné zařízení zodpovědné), tedy určit adresu vyšší úrovně a *směřovat*.

Datové jednotky na síťové vrstvě jsou označovány jako *pakety* nebo *datagramy* (záleží na konkrétním protokolu, a taky na dotyčném zdroji informace). Jaký je mezi nimi rozdíl?

- *Datagram* je datová jednotka posílaná v rámci nespojované (datagramové) služby. Problém je v tom, že datagramové služby existují i na jiných vrstvách a tento pojem lze používat třeba i na vrstvě L4.
- *Paket* je datová jednotka posílaná (nejen) v rámci spojované služby, ale často se používá i v obecnějším významu (prostě jako datová jednotka). Nicméně se tento pojem vžil především pro PDU na vrstvě L3.

Síťová vrstva je implementovaná na koncových zařízeních a dále na těch aktivních síťových prvcích, které zajišťují směrování, pracují s logickou topologií sítě, propojují nejen segmenty, ale také celé sítě, což jsou routery (směrovače). Router tedy implementuje vrstvy L1, částečně L2 (jen to, co je nezbytné k propojení s vyšší vrstvou) a L3.

Na síťové vrstvě je vedena (minimálně jedna) *směrovací tabulka* (routing table). Ve směrovací tabulce je informace o tom, kam poslat datovou jednotku patřící do určité (pod)sítě – adresa cílové (pod)sítě, brána (zde ve smyslu blízkého zařízení, přes které se dá do té (pod)sítě dostat, tedy určení dalšího kroku na cestě), síťové rozhraní, přes které mají data vyjít z tohoto zařízení a další informace.

 **Transportní vrstva (L4, Transport Layer).** Tato vrstva je jakýmsi přechodem mezi vrstvami orientovanými na proces přenosu (L1–L3) a vrstvami orientovanými na aplikace a tedy nezávislými na procesu přenosu (L5–L7). Takže směrem nahoru potřebuje vazbu na konkrétní protokol vyšší vrstvy (této vazbě, číslu, které je obdobou adresy, říkáme *port*) a směrem dolů uplatňuje funkce související s přenosem dat.



Poznámka:

Pozor – pojem „port“ se v počítačových sítích používá ve dvou velmi odlišných významech:

- (fyzický) port jako součást síťového rozhraní, tam zapojujeme kabel (nebo jde o anténu u bezdrátu),
- port (určený číslem) v záhlaví datové jednotky transportní vrstvy určující, se kterým protokolem vyšší vrstvy se právě komunikuje.

Například číslo 80 znamená při použití protokolů TCP nebo UDP komunikaci s protokolem HTTP.



Datovou jednotkou transportní vrstvy je *segment*, a pokud jde o nespojovanou službu, může se použít i pojem *datagram*. Hlavním úkolem protokolů transportní vrstvy (tedy kromě práce s číslem portu) je *segmentace dat* z vyšší vrstvy na dostatečně malé úseky, které bude možné přes síť přepřavit. Data jsou rozdělena na úseky o stanovené délce, ke každému je přidáno záhlaví s údaji transportní vrstvy, čímž je vytvořen segment.

Na transportní vrstvě se rozhoduje, zda bude přenos realizován formou služby se spojením nebo služby bez spojením. U služby se spojením zajišťuje vrstva L4 navázání spojení (handshake), řídí průběh spojení, potvrzování doručených segmentů, v případě ztráty či poškození dat zajistí opakování přenosu segmentu (tedy potvrzovaná, spolehlivá služba) a podle potřeby zajistí úpravu parametrů spojení, a pak ukončí spojení.

Transportní vrstva je implementovaná obvykle jen na koncových zařízeních, nicméně některým údajům v PDU transportní vrstvy mohou rozumět i jiná zařízení.

 **Relační vrstva (L5, Session Layer).** Na této vrstvě jsou oddělena data patřící různým aplikacím. Každá aplikace komunikující se sítí přes tuto vrstvu navazuje *relaci* (session) s nějakou aplikací na jiném systému, a v rámci této relace se například přenášejí data. Z toho vyplývá, že relace je logické spojení mezi dvěma aplikacemi navázané (většinou) za účelem výměny dat, které může vést přes síť.

Jaký je rozdíl mezi spojením na transportní vrstvě a relací na relační vrstvě? Zatímco L4 spojení se navazuje mezi dvěma zařízeními, L5 relace se navazuje mezi dvěma aplikacemi na těchto zařízeních. Jak bylo výše uvedeno – spodní vrstva poskytuje služby vyšší vrstvě, a tedy pokud relační vrstva chce pro určitou aplikaci navázat relaci s aplikací na jiném systému, potřebuje k tomu kromě jiného také spojení, které jí zajistí transportní vrstva.

Na relační vrstvě je implementován například koncept socketů. Také zde bývá implementováno to, co nutně potřebujeme k navázání relace – autentizace a autorizace.

Relační vrstva, stejně jako všechny vyšší vrstvy, je implementovaná na koncových zařízeních.

 **Prezentační vrstva (L6, Presentation Layer).** Vrstva L6 je zodpovědná za provádění konverzí dat (žádná z nižších vrstev do dat nezasahuje), například úpravy kódování textu (ASCII, EBCDIC, BASE64 apod.), komprese a dekomprese, šifrování a dešifrování, některé úkoly související se zpracováním multimédií.

Proč se tato vrstva vlastně nazývá prezentační? Protože má za úkol data prezentovat vyšším vrstvám v takové formě, které tyto vrstvy rozumí.

 **Aplikační vrstva (L7, Application Layer).** Na této vrstvě pracují protokoly, které jsou využívány aplikacemi. Například aplikace webový prohlížeč využívá (kromě jiného) aplikační protokol HTTP. Funkcí této vrstvy je při odesílání dat převzat data od aplikace a předat nižším vrstvám; naopak při přijímání dat jsou tato data přijata od nižších vrstev a předána příslušné aplikaci.



Poznámka:

Pozor, na aplikační vrstvě jsou aplikační protokoly, nikoliv aplikace!



Úkol

Zjistěte, co je to BASE64 a k čemu se dnes používá (soustřeďte se na přenosy binárních souborů).



1.5.2 Spolupráce protokolů

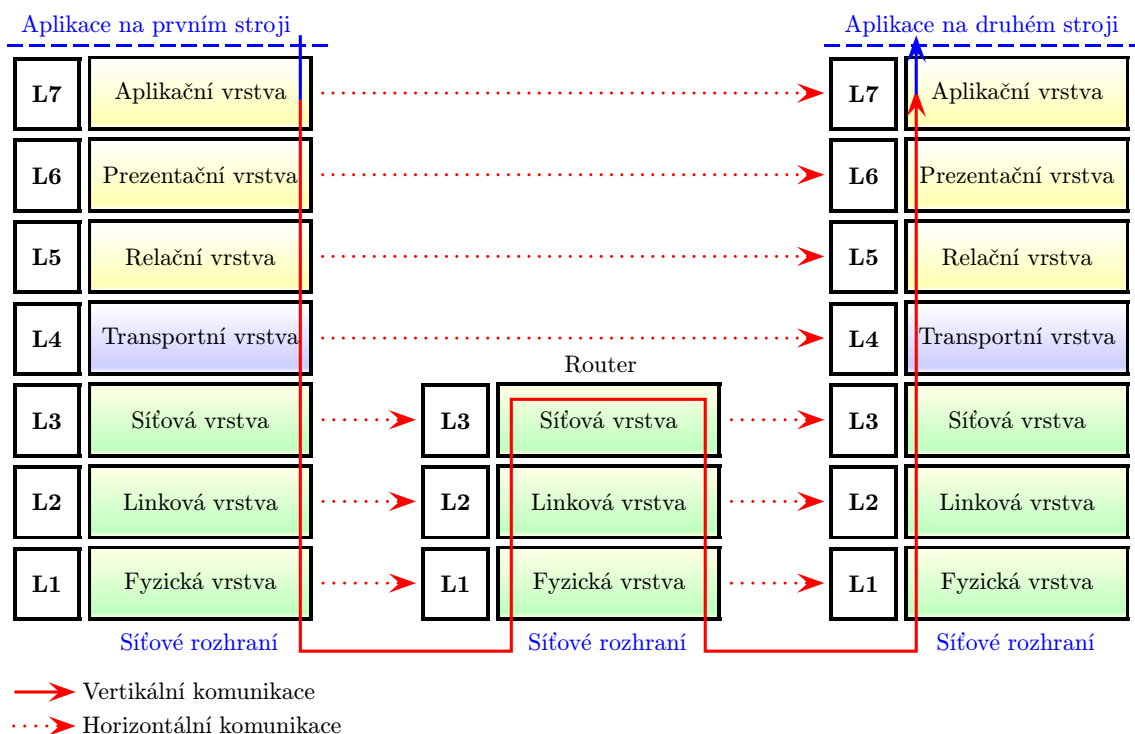
 Nadřazená vrstva vytvoří svou PDU (například paket) a odešle podřízené vrstvě. Pro podřízenou vrstvu je to, co takto obdržela, označováno jako *SDU* (Service Data Unit). Tedy celý postup je takový, že každá vrstva obdrží od nadřazené vrstvy SDU a přidáním záhlaví (a případně i zápatí) z ní vytvoří PDU a předá nižší vrstvě. To platí i pro aplikační vrstvu – nad ní sice žádná vrstva není, ale sekvenci dat, kterou obdrží od aplikace, která požaduje její služby, také označujeme SDU.

Víme, že protokoly nemají být moc komplexní, a tedy potřebují spolupracovat s jinými protokoly. Spolupráce může probíhat buď v rámci jedné vrstvy, nebo mezi sousedními vrstvami (ne na přeskáčku), přičemž platí, že spodní vrstva poskytuje služby horní vrstvě.

 *Entita* je aktivní prvek na určité vrstvě v modelu ISO/OSI, který má definováno rozhraní – sadu služeb, které může využívat entita z bezprostředně nadřazené vrstvy. Entitu vrstvy n označujeme *n-entita*. Na n -té vrstvě ISO/OSI je tedy sada n -entit.

 Rozhraní mezi komunikujícími entitami (a tedy mezi vrstvami) se nazývá *SAP* (Service Access Point, přístupový bod služby). SAP tedy propojuje dvě entity v sousedních vrstvách – uživatele služby (service user) a poskytovatele služby (service provider).

Způsob komunikace v rámci jednoho systému (řetěz střídačích se entit a SAP) se v ISO/OSI nazývá *vertikální komunikace*. *Horizontální komunikace* v ISO/OSI je komunikace mezi dvěma stejnými vrstvami umístěnými na různých strojích. Oba způsoby komunikace vidíme na obrázku 1.3.



Obrázek 1.3: Horizontální a vertikální komunikace v ISO/OSI

Pod pojmem entita si můžeme představit (spuštěnou) instanci některého konkrétního protokolu, případně na úrovni aplikací je to proces či vlákno procesu. Jeden SAP může být v jednom okamžiku využíván pouze jedním uživatelem a jedním poskytovatelem, ale jedna entita může zároveň používat více SAPů.

 Přes SAP se posílá blok *Interface Data Unit* (IDU), který se skládá ze dvou částí:

- *Service Data Unit* (SDU) = PDU některého protokolu vyšší vrstvy,
- *Interface Control Information* (ICI) = informační struktura potřebná k tomu, aby přijímající entita správně přijala SDU, tedy jakási metadata (řídící informace).

 Protokol tedy při odesílání dat (podle obrázku 1.3 na stroji vlevo)

- obdrží data přes SAP od vyšší vrstvy,
- pokud je to nutné, stanoveným způsobem je zpracuje či rozdělí na menší bloky,
- stanoví příslušné metainformace (tj. informace o informacích), například adresy, velikost dat, informace o tom, jak se má po cestě s daty zacházet, apod.,
- přidá záhlaví (header) s metainformacemi a pokud je to třeba, pak i zápatí (trailer), přidá k datům, čímž data „zabalí“ do PDU (paket, rámec, datagram apod.),
- předá PDU přes SAP nižší vrstvě.

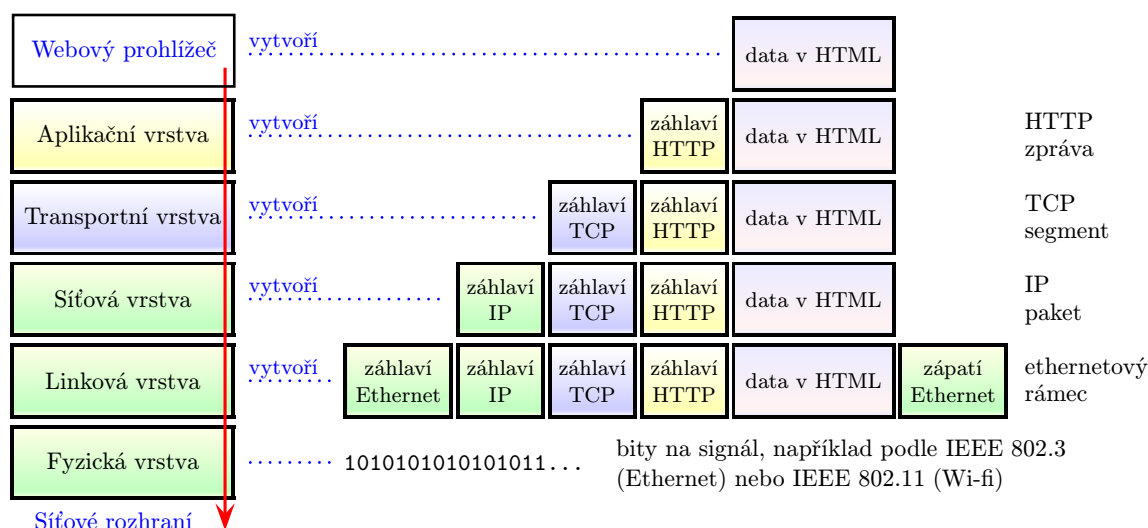
Jestliže jsou data naopak přijímána (podle obrázku 1.3 na stroji zcela vpravo), pak protokol

- přijme PDU přes SAP od nižší vrstvy,
- oddělí od PDU záhlaví a zápatí (pokud tam ovšem zápatí je),
- analyzuje metainformace v záhlaví a zápatí, stanoveným způsobem je zpracuje,
- pokud byla data před odesláním rozdělena na více bloků a jedná se o vrstvu, kde jsou bloky kompletovány (pozná se ze záhlaví), postupně shromáždí všechny bloky a zkompletuje,

- „rozbalená“ data předá přes SAP do vyšší vrstvy.

Není řečeno, že se nutně mají zúčastnit všechny vrstvy, některé nejsou pro konkrétní komunikaci potřebné.

 Při odesílání se tedy provádí *enkapsulace* (zabalení, zapouzdření, encapsulation) dat a při přijetí *dekapsulace* (rozbalení, decapsulation). To, co je pro nadřazenou vrstvu její vlastní PDU, to je pro jí podřazenou vrstvu pouze sekvence dat, ze kterých vytvoří vlastní PDU.



Obrázek 1.4: Zapouzdření PDU protokolu HTTP

Na obrázku 1.4 je naznačen proces enkapsulace. Postup při odesílání:

- data vyprodukovaná webovým prohlížečem jsou na vrstvě L7 zabalena do protokolové datové jednotky protokolu HTTP (do záhlaví se uloží například adresa serveru, informace o „žádajícím“ webovém prohlížeči, znakové sadě, časové údaje a další),
- následně se na transportní vrstvě dojedná spojení (nebo se využije už dojednané), pokud je to nutné, provede se segmentace (rozdělení na menší bloky) a následně se připojí TCP záhlaví (to obsahuje například číslo portu určující, že na nadřazené vrstvě komunikuje protokol HTTP, informace pro zajištění zkompletování celku při rozdělení na víc segmentů, kontrolní součet atd.), čímž je vytvořen TCP segment,
- na vrstvě L3 je tento segment předán entitě protokolu IP a je přidáno záhlaví protokolu IP (obsahuje například IP adresu zdroje a cíle, informaci o tom, který protokol sestavil to, co je „uvnitř“, atd.), výsledkem je IP paket,
- přes SAP mezi síťovou a linkovou vrstvou se IP paket dostane k entitě protokolu IEEE 802.3 (Ethernet), který před něj připojí ethernetové záhlaví (to obsahuje například synchronizační sekvenci bitů, aby byl rozeznatelný začátek bitů paketu, fyzické adresy apod.) a za něj ethernetové zápatí (s kontrolním součtem),
- další na řadě je fyzická vrstva, která již zajistí převod sekvence bitů na signál podle určeného protokolu a připojeného přenosového média.

Na straně adresáta proběhne opačný proces – rozbalování.

Ve skutečnosti by se tohoto přenosu účastnily i další protokoly, například protokol DNS překládající adresy nebo při zabezpečené komunikaci protokol SSL.

Všimněte si, že některé vrstvy modelu ISO/OSI se tohoto procesu neúčastní, jsou tedy výjimky z pravidla komunikují pouze bezprostředně sousedící vrstvy – protokoly aplikační vrstvy mohou používat přístupové body SAP z vrstev L6, L5 i L4.

 Jak se vlastně komunikuje přes takový SAP, co vše se na tomto přístupovém bodu děje? Ke každému SAPu jsou definována komunikační *primitiva*, což jsou jednoduché funkce, například *Request* (žádost o poskytnutí služby k nižší vrstvě, navazuje komunikaci přes SAP směrem dolů na straně odesílajícího stroje), *Indication* (upozornění na potřebu komunikace od nižší vrstvy k vyšší vrstvě, na straně přijímajícího stroje), *Response* (odpověď na Indication), *Confirm* (odpověď na Request).

Tato primitiva mohou mít, podobně jako běžné funkce, parametry – buď se jedná o SDU předávanou přes SAP z vyšší vrstvy, nebo se může jednat o dodatkové provozní informace, které nemají být součástí výsledné PDU.



Poznámka:

Zvědavého čtenáře určitě napadlo, že musí existovat způsob, jak se například síťová vrstva (L3) „dozví“, na kterou adresu má být vlastně dotyčný paket poslán. Přes záhlaví PDU to být nemůže, protože dovnitř záhlaví vyšších vrstev se protokol IP nedostane (je to pro něj prostě součástí dat, která je třeba poslat), navíc v některých záhlavích tato informace ani není. Ano, dozví se to z parametrů primitiv. Další informace najdete na <http://www.erg.abdn.ac.uk/users/gorry/course/intro-pages/service-prim.html>.



 *Socket* je kombinace síťové adresy (používané na vrstvě L3) a čísla portu (používaného na vrstvě L4). Pokud tento pár zapisujeme „ručně“ (třeba do adresního řádku webového prohlížeče), umístíme mezi oba údaje tečku. Místo síťové (tedy číselné) adresy může být použita doménová adresa.



Příklad

Socket na server s IP adresou 172.217.23.238 na portu 80 zapíšeme jako 172.217.23.238:80 tedy umístíme mezi adresu a port dvojtečku. Za socket se může považovat i údaj bez provedení DNS překladu, tj. například `www.google.com:80`.

S adresou IPv6 to je složitější, protože ta obsahuje dvojtečky sama o sobě, ale přitom zápis musí být jednoznačný. Proto například při zápisu do adresního řádku webového prohlížeče IPv6 adresu umístíme do hranatých závorek: `[2001:718:2201:208:5]:80`.



Sockety (sockets) najdeme na vrstvách L5–L7, víceméně fungují jako rozhraní mezi aplikačními protokoly a transportní vrstvou. Pro aplikace jsou dostupné jako Socket API (tj. rozhraní pro programování aplikací) ve formě knihoven obsahujících funkce pro práci se sockety (pro vytvoření socketu, akceptování na druhé straně spojení, naslouchání, čtení, zápis do socketu). Speciálním typem socketu je *stream socket*, který zaručuje dodání více bloků posílaných dat ve správném pořadí, a tedy na transportní vrstvě spolupracuje pouze s protokoly zajišťujícími spojovaný přenos (třeba TCP).

Se sockety se setkáme jak v UNIXových systémech (včetně Linuxu a MacOS), tak i ve Windows (WinSock API).

**Další informace:**

Na <https://www.binarytides.com/winsock-socket-programming-tutorial/> je tutoriál o programování socketů ve Windows v jazyce C/C++ s využitím WinSock API.



1.5.3 Protokolové zásobníky

 *Sada protokolů* (Protocol Suite) je definice (určení) skupiny protokolů, které vzájemně spolupracují. *Protokolový zásobník* (Protocol Stack) je implementace některé sady protokolů. Tyto dva pojmy se často používají jako synonyma.

Nemusí nutně jít o specifikaci protokolů pro naprosto všechny vrstvy modelu ISO/OSI, mohou být specifikovány pouze některé, přičemž se předpokládá spolupráce s jiným (doplňujícím) protokolovým zásobníkem.

 *TCP/IP Protocol Stack* (také se nazývá *Internet Protocol Suite*) je sada navzájem spolupracujících protokolů, kterou potřebujeme na zařízeních připojených k rozsáhlé síti (typicky Internetu). Kromě protokolů obsažených přímo v názvu (TCP, IP) zahrnuje ještě další, nejvíce na aplikační vrstvě. Přímou jsou určeny protokoly na vrstvách L3–L7, pro nižší vrstvy je pouze specifikováno rozhraní.

Tento protokolový zásobník je formalizován jako model TCP/IP, kterému se budeme podrobněji věnovat v následujícím textu (včetně protokolů).

 *IPX/SPX* je konkurenčním protokolovým zásobníkem k TCP/IP od společnosti Novell vytvořeným pro operační systém Novell Netware – IPX pracuje na L3 místo protokolu IP, SPX na vrstvě L4 místo TCP. Byl projektován spíše pro menší sítě, zatímco TCP/IP je určen i pro rozlehlé sítě. V současné době se už téměř nepoužívá.

 *Protokolové sady pro lokální sítě* jsou například Ethernet (IEEE 802.3), Wi-fi (IEEE 802.11) Token Ring (IEEE 802.5, už se nepoužívá) a další. Obvykle implementují pouze vrstvy L1 a L2, nad ně se nasouvá obvykle protokolový zásobník TCP/IP (ten pro změnu přímo nespecifikuje protokoly na L1 a L2).

 *Protokolové sady pro rozlehlé sítě* jsou například ATM, Frame Relay, MPLS a další. Také se obvykle napojují na TCP/IP, ale každá „trochu jinak“.

Například ATM se podsouvá pod síťovou vrstvu (L3), ale mezi ní a svou implementací vrstvy L2 vsouvá speciální přizpůsobovací vrstvu. Frame Relay implementuje vrstvu L2, na L1 předpokládá některé vhodné fyzické rozhraní, většinou dle standardů EIA/TIA (spoléhá na ISO/OSI).

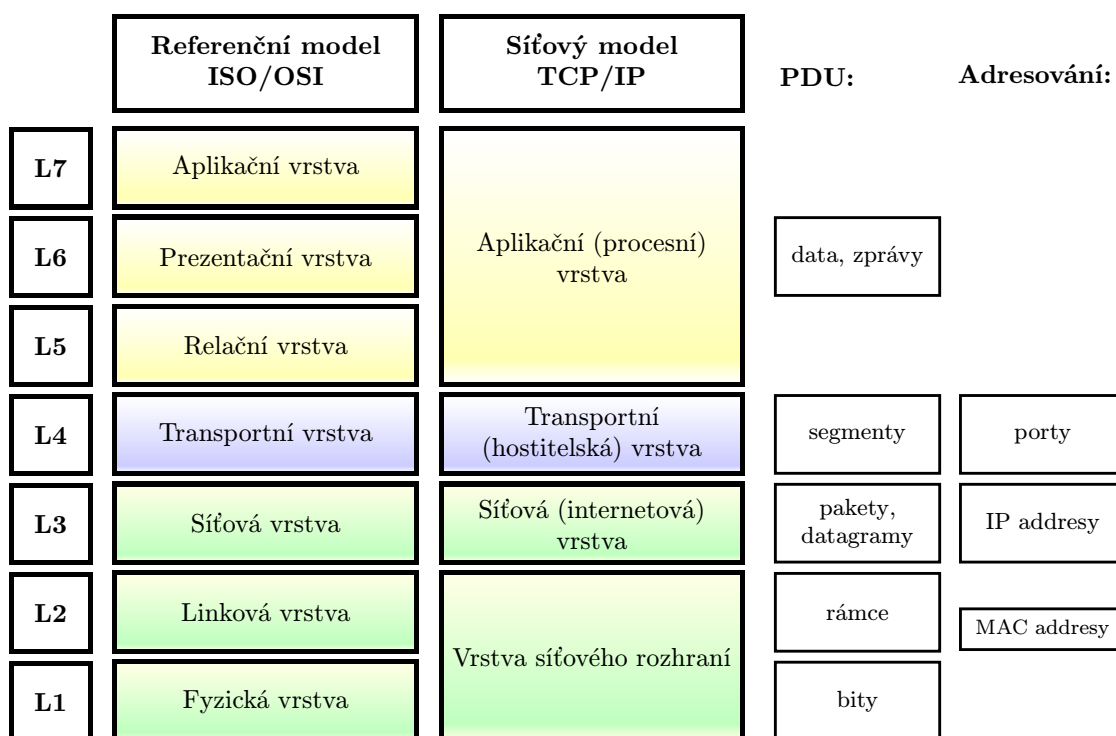
Oproti tomu MPLS se vsouvá mezi L2 a L3, tedy MPLS paket v sobě zabaluje paket z vrstvy L3 (většinou IP paket) a je zabalen do rámce vrstvy L2 (například do ethernetového rámce). Také může běžet nad ATM nebo Frame Relay, a tedy lze využít technologie ze starších zařízení. Také dokáže běžet nad PPP a dalšími protokoly pro přístupové sítě.

 *Protokolové sady pro mobilní sítě* jsou například sady pro LTE, GPRS, CDMA, UMTS a další. Implementují obvykle vrstvy L1 a L2 a předpokládají některé konkrétní protokoly i na vyšších vrstvách, ale ve skutečnosti záleží, o jaký typ zařízení jde (koncové zařízení bude potřebovat jinou sadu protokolů než základnová stanice nebo jiná specializovaná zařízení v mobilní síti).

1.6 Síťový model TCP/IP

Referenční model ISO/OSI je velmi komplexní, tudíž složitý, a příliš teoretický. Postupně bylo vytvořeno několik zjednodušených variant, z nichž je nejznámější právě síťový model TCP/IP, který je také nazýván DoD model (USA Department of Defense Model, tedy model Ministerstva obrany USA). Jeho součástí jsou také standardizovány organizací IETF a dostupné v RFC dokumentech.

Síťový model TCP/IP je vlastně formální popis síťového zásobníku TCP/IP, jeho napojení na spolupracující zásobníky a obecně možnost zapojení jiných protokolů. Skládá se ze čtyř vrstev, jejichž vztah k vrstvám RM OSI je naznačen na obrázku 1.5.



Obrázek 1.5: Srovnání modelů RM ISO/OSI a TCP/IP

Vztahy vyjádřené na obrázku platí i co se týče funkčnosti – aplikační (procesní) vrstva TCP/IP plní tutéž roli jako vrstvy L5–L7 v ISO/OSI. TCP/IP je navržen tak, aby

- byl co nejvíce decentralizovaný (žádná centrální správa),
- byl co nejodolnější vůči různým (i kritickým) podmínkám provozu a co nejodolnější vůči přenosovým chybám,
- nechával co nejvíce práce na koncových zařízeních (jádro sítě musí být rychlé a pružné), a aby byl spravovatelný distribuovaně (každá část si spravuje „to svoje“),
- dokázal propojit i sítě s hodně odlišnou síťovou architekturou a technologiemi (aby byl co nejuni-verzálnější při zachování předchozích vlastností).

 **Vrstva síťového rozhraní.** Tato vrstva v sobě sdružuje funkčnost vrstev L1 a L2 referenčního modelu. Přímo v TCP/IP pro ni nejsou stanoveny žádné protokoly, je jen určeno, jak mají komunikovat se síťovou vrstvou. Komunikuje s hardwarem (síťovým rozhraním), alespoň její část bývá hardwarově implementovaná.

Obvykle se do této vrstvy napojují protokolové sady lokálních a rozlehlých sítí (jinými slovy – přímo v TCP/IP pro tuto vrstvu nejsou žádné protokoly uvedeny), například

- IEEE 802.3 (Ethernet), resp. také protokol Ethernet II,
- IEEE 802.11 (Wi-fi),
- protokolové sady WAN sítí nebo jejich části, xDSL, mobilních sítí.

Kromě toho zde najdeme například „údržbové“ protokoly jako je STP, případně implementaci VLAN.

Na nižší části této vrstvy (obdoba L1) najdeme jednoduchá zařízení pracující pouze se signálem typu hub nebo repeater. Ve vyšší části této vrstvy (obdoba L2) pak o něco složitější zařízení pracující s rámci a spojující jednotlivé segmenty sítě, tedy switch a bridge.

 **Síťová (internetová) vrstva.** Na této vrstvě probíhá *internetworking* (propojování sítí), včetně *směrování* mezi sítěmi. Pojem „internet“ (s malým počátečním písmenem) značí obecně síť sítí, tedy síť propojující nikoliv jen jednotlivé uzly, ale jiné sítě. Typické L3 zařízení je router nebo L3 switch.

Co se protokolů týče, bude nás zajímat především protokol IP (Internet Protocol), a to jeho verze IPv4 a IPv6, které jsou dnes v praxi používány, a dále zde působí směrovací protokoly (OSPF, EIGRP, RIP, BGP a další) – pozor, neznamená to, že jsou zde plně implementovány. Z dalších protokolů to je například ICMP.

 **(Transportní vrstva).** Také ji nazýváme „hostitelská vrstva“, protože je implementována pouze na hostitelích v síti. *Hostitel* je název pro koncové zařízení (počítač, server, tablet, chytrá televize, atd.), které „hostí“ data, aplikace, služby; každý hostitel má svůj název (hostname).

Nejznámější protokoly transportní vrstvy jsou TCP a UDP.

 **Aplikační vrstva.** Také „procesní vrstva“, protože zdejší protokoly komunikují s procesy. Sdružuje v sobě vše, co je v RM ISO/OSI na nejvyšších třech vrstvách (L5–L7). Na této vrstvě najdeme velké množství protokolů, s většinou z nich komunikují aplikace potřebující přistupovat na síť. Příklady aplikačních protokolů: HTTP, HTTPS, SMTP, IMAP, POP3, DNS, DHCP, atd.

Poznámka:

Každý z modelů má svůj význam. Většinou se používá terminologie z ISO/OSI (označení vrstev, entity apod.), ale celkové rozvržení činností souvisejících se sítí a implementace postupů jsou obvykle podle TCP/IP. Ostatně, mnohé dnes používané protokoly pocházejí z protokolového zásobníku TCP/IP.



Definice (Internetworking)

Internetworking je mechanismus propojování sítí, tedy zajišťování komunikace mezi (inter) sítěmi (networks). Tento mechanismus probíhá na vrstvě L3, a to na zařízeních typu router (směrovač) nebo switch (přepínač) s funkcionalitou vrstvy L3.



1.7 Charakteristiky přenosu a přenosových cest

1.7.1 Přenos v základním a přeloženém pásmu

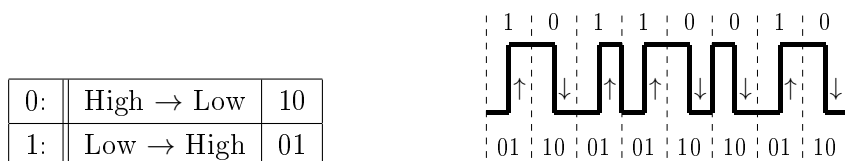
Přenos se provádí buď v základním pásmu nebo v přeloženém pásmu.

 **Přenos v základním pásmu (baseband).** Přenos probíhá tak, že sekvenci jedniček a nul, která má být přenesena, přímo zakódujeme do frekvenčního spektra (emitujeme signál) a takto vzniklý signál je přenesen komunikačním kanálem. Obvykle nám stačí nízké frekvence. Baseband se používá pro metalické lokální síť (Ethernet na měděném kabelu) a v základu pro optické síť.

Nevýhodou je omezený dosah (menší vzdálenost pro přenos) a problematická synchronizace (pokud bychom sekvenci nul a jedniček kódovali tak jak je, pak by dlouhé sekvence nul nebo dlouhé sekvence jedniček byly špatně dekódovatelné, nebylo by možné stanovit jejich délku). Odesílající a přijímající strana potřebují mít správně seřízené časovače, aby bylo možné synchronizovat intervaly mezi úseky představujícími jednotlivé bity, ale u dlouhých sekvencí bitů se stejnou hodnotou to nestačí.

Tento problém se obchází tak, že posloupnost bitů před kódováním na signál pozměníme podle určitého klíče tak, aby se v posloupnosti nevyskytovaly dlouhé sekvence stejných číslic, a samozřejmě aby bylo možné data po ukončení přenosu vrátit do původního stavu. Původní sekvence bitů se nahradí jinou sekvencí bitů tak, aby se dostatečně často „střídaly“ hodnoty 0 a 1. Dále si představíme několik obvyklých kódování pro přenos v základním pásmu.

 **Kódování Manchester** je velmi jednoduché. Spočívá v zakódování bitů do směru kmitu signálu – 0 je kódována jako přechod shora dolů, 1 je kódována jako přechod zdola nahoru (viz obrázek 1.6).



Obrázek 1.6: Kódování Manchester pro oktet $(10110010)_2$

Sekvence bitů 10111000 je při kódování Manchester zakódována na 01 10 01 01 01 10 10 10.

Jak vidíme, stejné symboly vedle sebe získáme pouze na těch místech, kde se v původní sekvenci měnily hodnoty bitů, a to nejvýše dva stejné symboly. To je výhodou kódování Manchester, nevýhodou je navýšení délky posílaných dat (délka se oproti původním zdvojnásobí). U vyšších rychlostí jsou další techniky jako data scrambling nebo samoopravitelný kód (Gigabit Ethernet).

Existují dvě základní varianty kódování Manchester:

- původní G.E. Thomas, kde jsou přechody opačně než na výše uvedeném obrázku (1: přechod dolů, 0: přechod nahoru),
- IEEE 802.3 pro 10b Ethernet s přechody přesně podle výše uvedeného obrázku.

Druhá varianta se tedy používá v 10Base-T, a také v jednom z typů spojení přes NFC.

Také se můžeme setkat s variantou differential Manchester, kde jsou přechody vždy jen při změně bitů (01 nebo 10). Tato varianta se používá například při zápisu na pevné disky.

 **Kódování 4B/5B:** z každé čtveřice bitů se vytvoří pět bitů tak, aby v celé sekvenci bylo co nejvíce jedniček (nejméně dvě na pěti), na začátku pěti nejvýše jedna nula, na konci nejvýše dvě nuly. Kódy jsou v tabulce 1.1.

Sekvenci 10111000, kterou jsme v kódování Manchester zakódovali do 16 znaků, zde zpracujeme na 1011110010 o délce 10 znaků.

4B/5B bylo původně určeno pro FDDI, v Ethernetu se s ním setkáme například v 100Base-X.

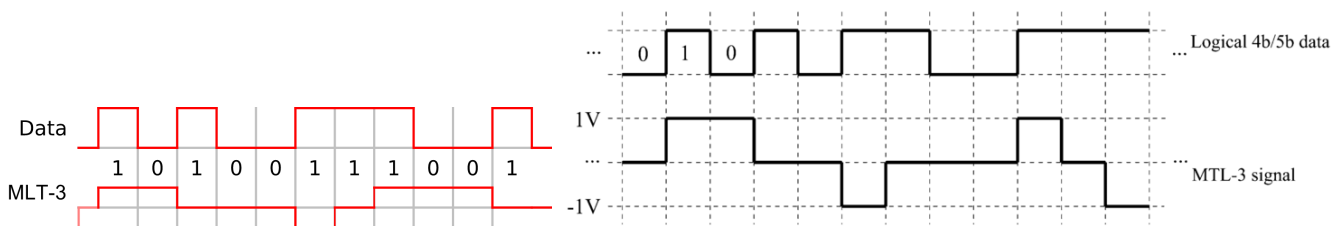
4B	5B	4B	5B	4B	5B	4B	5B
0000	11110	0100	01010	1000	10010	1100	11010
0001	01001	0101	01011	1001	10011	1101	11011
0010	10100	0110	01110	1010	10110	1110	11100
0011	10101	0111	01111	1011	10111	1111	11101

Tabulka 1.1: Tabulka kódů pro 4B/5B

Kódování 8B/6T znamená, že 8 bitů původní sekvence je zakódováno 6 změnami signálu (ternární symboly, tedy stavy -, 0, +). Používá se na kroucené dvojlince kategorie 3 se 4 páry, data se rozdělí do 3 párů, přenášejí se vždy jedním směrem. Čtvrtý pár je používán pro indikaci kolizí.

Další „lomítkové“ varianty také spočívají v tom, že určitý počet bitů ze zdroje se kóduje do určitého počtu bitů/stavů v cílovém signálu. Například *kódování 8B/10B* mapuje 8bitové sekvence do 10bitových sekvencí, používá se např. v 1000Base-X.

 *Kódování MLT-3* používá tři úrovně napětí (předchozí kódování používala jen dvě úrovně), a to -, základna, +. Změna napětí proběhne pouze na signál 1, a to na „sinusovce“. Toto kódování se obvykle kombinuje s některým jiným, například signál pro MLT-3 může být předzpracován kódováním 4B/5B.

Obrázek 1.7: Kódování MLT-3 (vlevo: samotné, vpravo: kombinace s 4B/5B)¹

Na obrázku vlevo vidíme kódování MLT-3 použité samostatně, vpravo v kombinaci s 4B/5B.

Oproti Manchesteru má signál jen čtvrtinovou frekvenci, proto méně vyzařuje, generuje mnohem méně rušení do okolí.

Například 100Base-TX kombinuje 4B/5B, NRZI (viz dále) a MLT-3.

 *Kódování PAM-5* používá čtyři úrovně napětí pro data a pátou pro řídicí bit. Data se kódují následovně:

- 00 → -1.0 V
- 01 → -0.5 V
- 10 → +0.5 V
- 11 → +1.0 V

PAM-5 se používá například v 1000Base-T.

 *Kódování NRZ (Non Return to Zero)* používá pro reprezentaci signálu dvě úrovně, z nich žádná není 0 (základna), například +1,-1. Existuje více variant – unipolární (1 je představována kladným napětím, 0 také kladným, ale menším), bipolární (0 kladné, 1 záporné), NRZI, atd.

¹Zdroj: podle https://cs.m.wikipedia.org/wiki/Linkov%C3%BD_k%C3%B3d,
https://www.researchgate.net/figure/Mapping-from-a-4b-5b-data-stream-to-MLT-3-signal-level_fig1_3056686

 **Kódování NRZI** (*Non Return to Zero – Inverted*) znamená, že 1 vyvolá změnu signálu, 0 ne. Je to obdoba MLT-3, ale jen na dvou úrovních (kladné a záporné).

Poznámka:

Pokud vám z předchozího studia nic neříkají označení 100Base-T, 1000Base-X apod., nalistujte si tyto zkratky v následující kapitole o Ethernetu.

 **Přenos v přeloženém pásmu (broadband).** Tento přenos probíhá tak, že sekvenci jedniček a nul sice zakódujeme také do frekvenčního spektra jako u basebandu, ale vzniklý signál přeložíme do takového pásma, které je pro tento konkrétní přenos určeno (nebo kódujeme data přímo do příslušného frekvenčního pásma). Tomuto překladu říkáme *modulace*.

Modulace tedy probíhá následovně:

- Určíme *nosnou*, což je vhodný signál na té frekvenci, na které mají být data přenesena, volí se signál s harmonickou frekvencí, obvykle sinusoida.
- Pozměníme parametry (amplitudu, frekvenci nebo fázi) tohoto signálu podle toho, jaká data mají být přenášena – modulujeme data na signál.
- Podle potřeby vše sloučíme a odešleme.

 Rovnice popisující průběh modulace je $s(t) = A \sin(\omega \cdot t + \varphi)$, kde t je čas (to je proměnná), A je amplituda signálu, ω je úhlový kmitočet a φ je fázový posun.

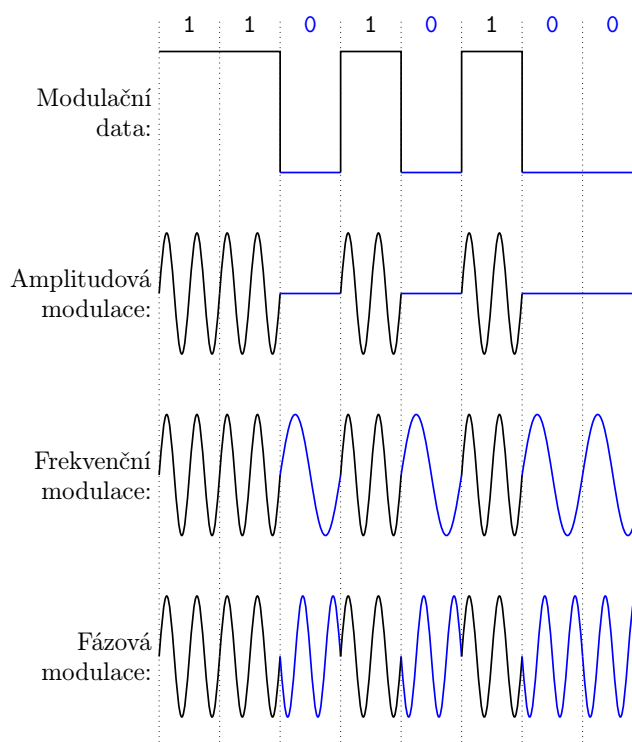
Na obrázku 1.8 na straně 24 jsou ukázky tří základních typů modulace – amplitudové (mění se amplituda), frekvenční (mění se frekvence signálu) a fázové (mění se fáze), vodorovná osa je proměnná t , svislá výsledek $s(t)$. Ve skutečnosti se často mění více než jen jeden z těchto tří parametrů, navíc různé modulace mohou tentýž parametr měnit s různou intenzitou. Podíváme se na několik nejznámějších modulací.

 **QAM** (*Kvadraturní amplitudová modulace*) používá kombinaci fázového (PSK) a amplitudového (ASK) posunu. Dokáže modulovat jak analogový, tak i digitální signál do příslušného frekvenčního rozsahu ve výsledném analogovém signálu. Existuje více variant: 16-QAM, 64-QAM, 256-QAM. V každé variantě je u určitého počtu nosných použita amplitudová modulace, u zbývajících fázová, kombinací získáme určitý počet stavů reprezentujících modulační data či signál. Například u nejjednodušší varianty 16-QAM se používají dvě nosné a existuje 36 kombinací, ale využíváno je pouze 16 (těch, které jsou od sebe nejsnáze rozlišitelné).

Modulace QAM v kombinaci s multiplexováním OFDM (viz dále) se dnes používá především v bezdrátových a mobilních sítích, například Wi-fi podle standardu IEEE 802.11n používá až 64-QAM, kdežto IEEE 802.11ac používá až 256-QAM (při špatném signálu „spadne“ na nižší variantu). V mobilních sítích čtvrté generace (LTE Advanced) se používá 256-QAM.

Další informace:

- <http://www.cs.vsb.cz/grygarek/PS/lect/PREZENTACE/fyzPrincipy.pdf>
- <http://www.earchiv.cz/l226/slide.php3?l=4>



Obrázek 1.8: Příklad amplitudové, frekvenční a fázové modulace digitálních dat

Baseband	Broadband
posíláme digitální signál	výsledkem je analogový signál
signál přímo emitujeme	modulační signál modulujeme na nosnou
na kratší vzdálenosti	na delší vzdálenosti
signál využívá celou šířku pásma	signál lze omezit na stanovenou šířku pásma

Tabulka 1.2: Rozdíl mezi přenosem v základním a přeloženém pásmu

1.7.2 Multiplex

Vezměme jeden přenosový kanál s určitými fyzikálními charakteristikami (frekvence apod.). Za normálních okolností bychom tímto přenosovým kanálem mohli přenášet jen jeden signál.

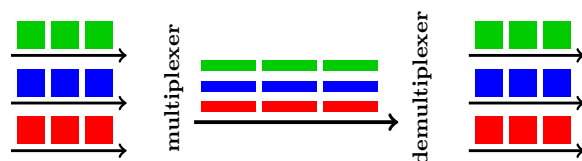
 **Multiplexing** je technika, která umožňuje rozdělit přenosový kanál s dostatečnou šířkou pásma na více logických subkanálů, zdánlivě samostatných, a v každém přenášet jiné bloky dat. V reálu to znamená, že se více signálů sloučí do jediného signálu. Multiplex vlastně znamená „multiple access“, tedy vícenásobný přístup (současný přístup více uživatelů k přenosovému kanálu).

Komponenta, která provádí multiplexing, se nazývá multiplexer (multiplexor, MUX). Opačnou operací k multiplexingu je demultiplexing (DEMUX, DMX): Na odesílaající straně přenosového kanálu se provádí multiplexing (rozdělení do subkanálů), na přijímající straně demultiplexing (odebrání ze subkanálů).

Existuje několik běžných typů multiplexování:

 **Frekvenční multiplex** (FDM, Frequency Division Multiplex): každému přenášenému signálu je přidělena část šířky pásma (mezi přidělenými pásmy musí být odstup, aby nedocházelo k rušení), a tento signál je namapován do této přidělené části. Každý subkanál má přidělen vlastní interval frekvencí.

FDM je použitelný pouze na analogový signál, a další jeho nevýhodou je, že je vhodný spíše pro



Obrázek 1.9: FDM – Frekvenční multiplex

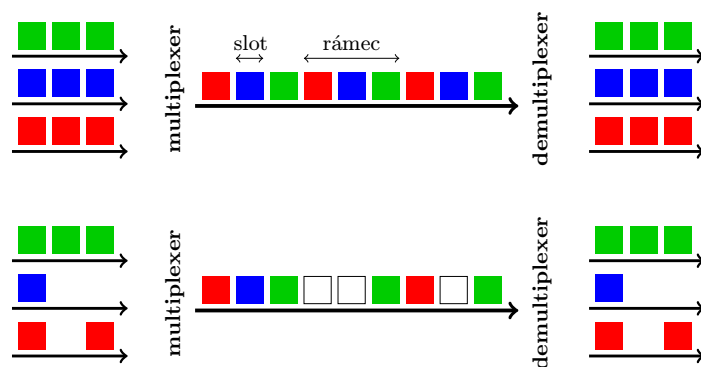
„stabilní počet“ uživatelů (resp. existuje strop pro množství uživatelů).

Frekvenční multiplex známe například z rozhlasu, kdy různé stanice mají přiděleny různé frekvence, dále například v satelitních přenosech a kabelových sítích. Také se používal na analogové telefonní síti – každý „telefonista“ měl vyhrazen subkanál o šířce 3.1 kHz.

Vlnový multiplex (vlnové dělení, WDM, Wave Division Multiplex): je to obdoba frekvenčního multiplexu, ale u optického signálu, kde místo frekvencí používáme dělení podle vlnových délek neboli barev světla (což je, jak víme, vlastně podobné). Každý signál dostane přidělen určitý rozsah vlnových délek, které jsou použity při jeho přenosu. Podle příkladu v předchozí sekci je zřejmé, že u světla se s vlnovými délkami pracuje lépe než s frekvencemi, jsou to menší čísla.

WDM se používá například při přenosu optickým kabelem, tedy především v optických WAN sítích a dále například v technologii FTTx (optické vlákno se vede co nejdříve zákazníkovi). Na signál generovaný laserem nebo LED diodou je po multiplexování modulován signál.

Časový multiplex (TDM, Time Division Multiplex): přenosový kanál je střídavě přidělován různým konkrétním přenosům. Princip je naznačen na obrázku 1.10. Přenosový kanál je rozškálován na tzv. *rámce* (pozor, to nejsou tytéž rámce jako na vrstvě L2, jen shoda názvů), v každém rámci je pro každého odesílajícího vyhrazen jeden *slot* (zásuvka), do kterého lze umístit paket (nebo ho nechat prázdný).



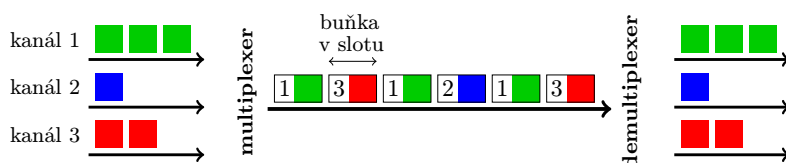
Obrázek 1.10: TDM – Časový multiplex (nahore plné vytížení, dole částečné vytížení)

Problém „čistého“ TDM je, že je taky vhodný spíše pro relativně konstantní počet uživatelů, navíc víceméně podobně komunikujících (co se týče počtu odesílaných paketů). Pokud některý uživatel odesílá výrazně méně paketů, jsou jeho sloty nevyužity, jak vidíme na obrázku 1.10 dole. Další nevýhodou časového multiplexu je nutnost neustálé synchronizace rámců – oběma komunikujícím stranám musí být jasné, kdy začíná a končí rámec a kdy začíná/končí který slot a komu patří. Nicméně na rozdíl od FDM je vhodný i pro digitální signál.

TDM se používá například v mobilních sítích při použití základního přenosu GSM.

Statistický multiplex je podobný časovému multiplexu v tom, že přenosový kanál je členěn na sloty. Ovšem u statistického nejsou sloty napevno přiděleny konkrétním subkanálům, ale jsou přidělovány

podle potřeby – vytíženější subkanál dostane víc slotů než méně vytížený. Aby bylo jasné, ke kterému subkanálu který slot patří, musí být k přenášeným blokům dat přidána informace o subkanálu (záhlaví).



Obrázek 1.11: Statistický multiplex

Výhodou statistického multiplexu je o něco nižší potřeba synchronizace (není nutné zabývat se rámci slotů), synchronizují se jen samotné sloty a celková komunikace zůstává asynchronní. Další výhodou je lepší vytěžování přenosového kanálu. Nevýhodou je nutnost opatřovat data záhlavími a s tím spojená režie.

Statistický multiplex se používá v některých WAN sítích, například v ATM (do slotů se skládají datové jednotky zvané buňky).

Kódový multiplex (CDM, Code Division Multiplex, také CDMA, Code Division Multiple Access, rozprostřené spektrum) je digitální metoda, jejíž vznik byl motivován především potřebou bezpečnosti (co nejvíc ztížit odposlech). Každý dílčí přenos je zakódován (kódování probíhá v koncových zařízeních, neprovádí je multiplexer), všechny paralelní (zakódované) přenosy jsou pak multiplexerem sloučeny a přenášeny sdíleným kanálem až ke koncovým zařízením. Cílové koncové zařízení pak s pomocí speciálního kódu dekóduje jen to, co mu ve skutečnosti patří, bez tohoto kódu se k obsahu subkanálu nelze dostat.

Problém CDM je potřeba složité synchronizace a horní limit pro počet subkanálů (příliš mnoho subkanálů by se navzájem rušilo a nebylo by možné je dekódovat). Také je nutné zajistit bezpečnou domluvu o kódu pro dekódování. Existuje víc různých variant CDM, většinou se používají v mobilních sítích třetí generace.

Ortogonalní multiplex (OFDM, Orthogonal Frequency-Division Multiplex) mapuje subkanály na různé frekvence podobně jako FDM, ale mnohem efektivněji. Zatímco FDM používá jedinou nosnou pro všechny subkanály, OFDM používá pro každý subkanál jinou nosnou, přičemž jednotlivé nosné jsou navzájem ortogonální, proto se mohou překrývat a přesto je lze na straně přijímače oddělit. Data přenášená přes subkanál jsou modulována na přidělenou nosnou některou vhodnou modulací, většinou se používá některá varianta modulace QAM.

OFDM (a taktéž jeho varianty) se dnes široce používá pro modulaci digitálních dat do analogového signálu v počítačových sítích (Wi-fi, WiMAX, LTE, xDSL apod.), ale také například pro přenos digitální televize.

Existují různé varianty OFDM přizpůsobené konkrétnímu způsobu využití v některé technologii. Například u mobilních sítí se často setkáváme s OFDMA (OFDM Access), kde jsou subkanály jednoho kanálu rozprostřeny po celém spektru a mohou být přidělovány různým klientům.



Další informace:

- <http://fyzika.jreichl.com/main.article/view/156-harmonicke-kmitani>
- <http://www.cs.vsb.cz/grygarek/PS/lect/PREZENTACE/SdileniMedia.pdf>
- http://www.wikiskripta.eu/index.php/Elektromagnetick%C3%A9_spektrum

-
- <https://khanovaskola.cz/video/8/27/1423-amplituda-perioda-frekvence-a-vlnova-delka-periodickeho-vlneni>
-



Souhrn kapitoly

V této kapitole jsme si zopakovali základní principy, na kterých fungují počítačové sítě, a pojmy, které budeme dále potřebovat. Soustředili jsme se na to, jak probíhá zpracování dat na různých vrstvách ISO/OSI, včetně toho, jak probíhá zpracování signálu na fyzické vrstvě.



Lokální síť – Ethernet

 *Rychlý náhled:* V této kapitole se budeme zabývat sítí Ethernet, resp. IEEE 802.3. Část obsahu je opět opakováním učiva z bakalářského stupně studia, část na tuto znalost navazuje. Zaměříme se na fungování Ethernetu na vrstvách L1 a L2 včetně práce se signálem, různých druhů rámců a průběhu přenosu, na konci kapitoly se budeme zabývat souvisejícími technologiemi včetně napájení zařízení přes Ethernetový kabel.

 *Klíčová slova:* Ethernet, IEEE 802.3, CSMA/CD, Backoff algoritmus, rámec, LAN, EtherType, LLC, MAC, LLC rámec, SNAP rámec, rámec Ethernet II, tabulka MAC adres, OUI, L/G bit, I/G bit, Collision Window, IFG, Burst mode, MAU, autonegociace, PoE, EFM, strukturovaná kabeláž.

 *Cíle studia:* Cílem této kapitoly je porozumět do hloubky problematice lokálních sítí se zaměřením na Ethernet, a také souvisejícím standardům.

2.1 Co je to Ethernet

 Pod pojmem lokální síť (LAN) budeme dále rozumět souhrn navzájem propojených zařízení (hostitelských/koncových zařízení, aktivních síťových prvků apod.), která patří do téže sítě (tj. jako aktivní síťové prvky jsou použity nejvýše switche/přepínače), obvykle v rámci jedné budovy či bytu, s typickou rozlohou jednotek až stovek metrů, výjimečně více.

V současné době jsou nejpoužívanějšími LAN technologiemi Ethernet (kabel) a Wi-fi (bezdrát). V této kapitole se zaměříme na Ethernet, Wi-fi budeme probírat v jedné z dalších kapitol.

Na původní specifikaci Ethernetu spolupracovaly společnosti Xerox, Digital a Intel, tato specifikace byla zveřejněna roku 1976 a označuje se jako DIX Ethernet (podle počátečních písmen společností).

 Později byl Ethernet standardizován jako IEEE 802.3, ale v tomto standardu se název „Ethernet“ vůbec nevyskytuje a s původním DIX Ethernetem je nekompatibilní. Postupně se objevovaly různé varianty – pro různá přenosová média (metalické kabely různých kategorií, optické kabely) a také se navyšovala rychlost.

**Definice (Ethernet, IEEE 802.3)**

Síť Ethernet, resp. IEEE 802.3, je standard popisující souhrn technologií pro lokální počítačovou síť používající kabely (metalické nebo optické), kde hostitelská zařízení sdílejí stejnou šířku komunikačního pásma a vzájemně o ni soupeří.

Standard IEEE 802.3 popisuje implementaci pro fyzickou a linkovou vrstvu (tj. celou vrstvu síťového rozhraní podle TCP/IP) včetně metody komunikace a řešení kolizí.



Typické vlastnosti sítě Ethernet (nebo IEEE 802.3):

- Všechna zařízení v síti jsou rovnocenná, žádné nemá prioritu.
- Jako kabeláž se používá buď kroucená dvojlinka (twisted pair – nestíněná nebo stíněná) nebo optika, v datových centrech se můžeme setkat také například s twin-ax kabelem.
- Typickým aktivním síťovým prvkem je switch.
- Všechna zařízení na segmentu sdílejí přenosové médium, ovšem při použití switche je logická topologie segmentu v podstatě point-to-point.
- Komunikuje se v polovičním nebo plném duplexu, dnes je typický plný duplex.

**Poznámka:**

V dalším textu budeme hovořit o Ethernetu, nicméně v standardu IEEE 802.3 ani jeho dodatcích se slovo „Ethernet“ vůbec nevyskytuje. Tento „rozpor“ bude dále vysvětlen.



2.2 Komunikace v Ethernetu



Přehled pojmů:

- *DTE* (Data Terminal Equipment, terminálové zařízení) – koncové zařízení v síti (počítač, server apod.),
- *DCE* (Data Circuit-terminating Equipment, komunikační zařízení, zařízení ukončující okruh) – zařízení v síti, které přeposílá provoz, obvykle není zdrojem ani cílem dat (obvykle switch),
- *kolizní doména* (segment) – souhrn zařízení v síti, která se navzájem „vidí“; pokud některé začne vysílat, všichni ostatní na segmentu jsou cílem a zároveň vysílat nemohou (jinak by došlo ke kolizi), vzhledem k aktivním síťovým zařízením:
 - hub neodděluje kolizní domény,
 - switch a router oddělují kolizní domény,
- *všesměrová (broadcastová) doména* – skupina zařízení v síti, kterým je doručeno broadcastové vysílání, jehož zdrojem je některé zařízení z této skupiny:
 - hub ani switch neoddělují broadcastové domény,
 - router odděluje broadcastové domény.



Přístupová (kolizní) metoda určuje, jakým způsobem se rozhoduje, zda zařízení může či nemůže začít vysílat. Pro Ethernet je specifikována přístupová metoda CSMA/CD.

**Definice (Přístupová metoda CSMA/CD)**

V Ethernetu se používá technologie vícenásobného přístupu k přenosovému médium s nasloucháním nosné a detekcí kolizních stavů – CSMA/CD:

- CS (Carrier Sense) – uzly v síti neustále naslouchají na nosné, zda nevysílá jiný uzel,
- MA (Multiple Access) – vícenásobný přístup, tedy kterýkoliv připojený uzel může začít vysílat, pokud nasloucháním zjistí, že nikdo jiný nevysílá,
- CD (Collision Detection) – pokud uzel v síti před vysíláním nestihne včas detekovat vysílání jiného uzlu (například tehdy, když oba začnou vysílat v přibližně stejné době), musí být schopen vzniklou kolizi signálů detekovat a patřičně na ni reagovat.



Použití přístupové metody má smysl tehdy, když se komunikuje v polovičním duplexu (nebo když používáme huby, ale to už je celkem historie).



Co se stane, když nastane kolize:

- Vysílající uzel buď chce vysílat a přitom zjistil, že vysílá jiný uzel, nebo detekoval kolizi (zjistil, že kromě něho vysílá i někdo jiný).
- Pokud už vysílal, přestane vysílat (ne okamžitě, musí umožnit i druhému vysílajícímu uzlu, aby kolizi detekoval). Vyšle *Jam signál*, což je signál oznamující kolizi na médium.
- Podle speciálního algoritmu (Backoff algoritmus) určí dobu čekání a po uplynutí této doby se pokusí znovu vysílat.
- Jestliže i další pokus selže (buď ještě před vysíláním zjistí, že linka není volná, nebo opět zjistí kolizi), vrací se k předchozímu bodu (doba čekání bez vysílání a nový pokus).

Backoff algoritmus určuje, jak dlouho má uzel čekat s vysíláním, když zjistí, že vysílá jiný uzel. Účelem je zajistit, aby v případě více zájemců o vysílání každý z nich počkal jinou dobu (určenou náhodně generovaným číslem), čímž by se snížila pravděpodobnost další kolize.

**Postup (Backoff algoritmus)**

Existuje víc různých (různě složitých) variant tohoto algoritmu, základní (Exponential Backoff) je následující:

- Při kolizi vyšle „Jam“ signál, čímž informuje o kolizi a zrušení právě odesílaného rámce.
- Čeká po dobu $0 \dots 51,2 \mu s$ (náhodně vygenerované číslo z tohoto intervalu), pak se znovu pokusí o přenos.
- Jestliže přenos znovu selže, čeká po dobu $2 \times 0 \dots 51,2 \mu s$ (opět se generuje náhodné číslo) a pak se znovu pokusí o přenos.
- Pokud dojde k dalším selháním přenosu, čeká vždy po dobu $K \times 0 \dots 51,2 \mu s$, kde K je náhodně vygenerované číslo z intervalu $0 \dots 2^c - 1$, kde c je dosavadní počet kolizí (neúspěšných pokusů o odeslání).

Po 10 pokusech se již c nezvyšuje, pak je vždy K z intervalu $0 \dots 2^{10} - 1$. Po 16 pokusech postup končí, rámec je považován za nedoručitelný.



Jak vidíme, po prvních dvou selháních přenosu se doba čekání určuje vygenerováním jednoho náhodného čísla, ale od třetího pokusu dále vlastně násobíme dvě náhodná čísla, přičemž pro první z nich se horní mez každým krokem exponenciálně zvyšuje. Tím se snižuje pravděpodobnost, že dva uzly opakují pokus o přenos ve stejnou dobu.



Poznámka:

Uvědomme si, že zatímco fyzická topologie se projevuje na vrstvě L1 (fyzické), logická topologie na vrstvě L2 (linkové). Původní Ethernet používající koaxiální kabel byl sběrníkový na fyzické i logické úrovni (na fyzické díky napojení zařízení na sdílené přenosové médium, na logické díky přístupové metodě CSMA/CD), novější je sběrníkový jen na logické úrovni, a to jen tehdy, když se používá CSMA/CD.



2.3 Ethernet na linkové vrstvě

2.3.1 Adresace na vrstvě L2

Na linkové vrstvě se používají MAC adresy (fyzické, hardwarové adresy). V každém rámci je MAC adresa cílového zařízení (adresáta) a MAC adresa zdrojového zařízení (odesílatele).

 Zařízení pracující na vrstvě L2 by mělo mít přehled o *fyzické topologii sítě*, tedy o adresách zařízení, s nimiž může komunikovat, a o tom, přes které síťové rozhraní (kterou cestou) je dotyčné zařízení dosažitelné. Tyto informace si každé zařízení vede ve své tabulce MAC adres (může se nazývat MAC tabulka, CAM tabulka apod., podle toho, jak ji nazve výrobce).

 *Hardwarová (MAC, fyzická, EUI-48) adresa* je vlastně nízkoúrovňovou adresou síťového rozhraní. Tuto adresu má každé síťové rozhraní, a to i tehdy, když zrovna není připojeno k síti.

Máme dva druhy MAC adres:

- adresa přidělená výrobcem (BIA – Burned-In-Address, „vypálená“), která je uložena „napevno“ v ROM nebo flash paměti síťového rozhraní,
- lokální (dočasná) MAC adresa, kterou jsme si nastavili sami.

V současných sítích se používají 48bitové MAC adresy (tj. 6 oktetů) a zapisují se většinou hexadecimálními číslicemi – MAC adresu zapíšeme pomocí 12 hexadecimálních číslic. Jednoznačnost je zajištěna takto:

- První polovinu adresy dostane výrobce přidělenou od sdružení IEEE (velcí výrobci mají přiděleno několik, menším stačí jedna), tedy první polovina je charakteristická pro výrobce a žádní dva výrobci nemají stejnou. Toto číslo se označuje OUI (Organizationally Unique Identifier).
- Druhou polovinu určuje již samotný výrobce, který si hlídá, aby byla tato čísla unikátní v rámci jemu přidělené první poloviny.

V tabulce 2.1 je výše popsaná struktura (dvě části) přehledněji naznačena.

Fyzickou adresu obvykle zapisujeme tak, že dvojice hexadecimálních číslic oddělíme dvojtečkou nebo pomlčkou, nebo dvojice oktetů oddělíme tečkou, případně je nijak neoddělujeme. Například:

50:E5:49:A2:80:61

50-E5-49-A2-80-61

50E5.49A2.8061

50E549A28061

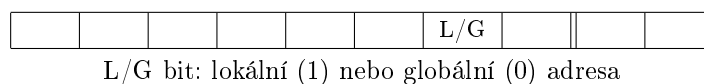
Změna MAC adresy není až tak běžnou věcí, ale někdy se hodí – například tehdy, když pro účely využití určité technologie či aplikace potřebujeme nutně MAC adresu v určitém konkrétním tvaru a není

24 bitů	+	24 bitů	=	48 bitů
OUI = identifikace výrobce přiděluje IEEE	+	identifikace konkrétního výrobku určuje výrobce	=	celá adresa globálně jednoznačná

Tabulka 2.1: Struktura MAC adresy síťového rozhraní

čas či možnost změnit konfiguraci, nebo ve virtualizovaném prostředí (třebaže i tam mohou být používány unikátní adresy). Hackeři používají pozměněné MAC adresy při pokusu o průnik do sítě chráněné blacklistem nebo whitelistem (seznamem zakázaných/povolených adres).

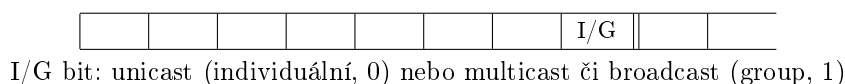
 Každá lokálně platná adresa (tj. ne BIA od výrobce) by správně měla mít nastavený *L/G bit*. Tento bit je v prvním oktetu adresy druhý zprava (v běžných síťových implementacích, například v Ethernetu), jak je vidět na obrázku 2.1.



Obrázek 2.1: Umístění L/G (local/global) bitu v MAC adrese

 Pokud jsou všechny hexadecimální číslice nastaveny na *F*, znamená to, že tato adresa je nejen lokálně platná (není globálně jednoznačná), ale navíc je *všeobecná* (broadcastová), tedy neoznačuje jedno konkrétní zařízení. Vypadá takto: **FF-FF-FF-FF-FF-FF**. Všeobecné MAC adresy se používají například tehdy, když je odeslán rámec určený pro všechna zařízení v lokální síti.

 Pokud adresa má označovat skupinu zařízení (ale ne nutně všechna), nazývá se *skupinovou* (multicast) adresou. V Ethernetu a některých dalších sítích používajících MAC adresy poznáme skupinovou adresu podle bitu v prvním oktetu nejvíc vpravo (nejméně významný bit prvního oktetu). U skupinové adresy je nastaven na 1, což znamená, že první oktet bude liché číslo.



Obrázek 2.2: Umístění I/G (individual/group) bitu v MAC adrese

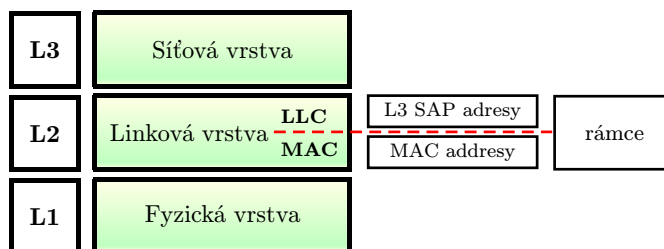
2.3.2 Podvrstvy L2

 Na vrstvě L2 (linkové) pracují tzv. *spojové protokoly* (protokoly vrstvy L2). Takový protokol buď sám určuje fungování celé vrstvy L2, nebo se používají dva protokoly, které totéž zvládnou společně. Vrstva L2 se totiž ve skutečnosti dělí na dvě podvrstvy:

- LLC (Logical Link Control) – horní podvrstva, zajišťuje spolupráci s vrstvou L3 (síťovou), přidává do rámce informaci o SAP pro spolupráci s L3 a další řídicí informace, zde se také provádí multiplexování.
- MAC (Media Access Control) – spodní podvrstva, zajišťuje spolupráci s vrstvou L1 (fyzickou), přidává do rámce fyzickou adresu příjemce a odesílatele, přizpůsobuje strukturu rámce pro stanovenou přenosovou rychlost a na začátek vloží synchronizační posloupnost bitů (aby bylo zřejmé, kde začíná rámec), zde je implementován mechanismus CSMA/CD.

Z funkce těchto podvrstev vyplývá, že na aktivních síťových prvcích, které implementují nejvýše vrstvu L2 (switche, mosty), vlastně ani není nutné mít celou podvrstvu LLC.

Na obrázku vpravo je naznačen vztah podvrstev linkové vrstvy k okolním vrstvám. Jak bylo výše řečeno, v protokolové sadě může buď obě podvrstvy „naplnit“ jeden protokol, nebo jsou použity dva protokoly – jeden pro podvrstvu LLC a druhý pro podvrstvu MAC.



Ethernet může být na vrstvě L2 implementován následovně (co se týče protokolů na L2):

1. Pro podvrstvu LLC je použit protokol IEEE 802.2 (přímo se nazývá LLC), data z nadřízené vrstvy se zapouzdří do LLC rámce. Ten je následně zapouzdřen do MAC rámce podle protokolu IEEE 802.3, čímž vznikne kompletní rámec vrstvy L2.
2. Podobně jako první bod, jen místo protokolu LLC se použije jeho rozšíření SNAP.
3. Data z nadřízené vrstvy (L3) se zapouzdří podle protokolu Ethernet II, který „zvládá“ obě podvrstvy – LLC i MAC.

První dvě možnosti mohou být využívány, ale momentálně se s nimi u Ethernetu téměř nesetkáváme. Protokol IEEE 802.2 se pro implementaci podvrstvy LLC používá spíše v bezdrátových sítích a také ve starších sítích typu Token Ring nebo FDDI. SNAP je vylepšením původního protokolu LLC, umožňuje napojovat do komunikace širší množinu protokolů. Nejčastěji se setkáváme právě s třetí možností, takové rámce se označují *Ethernet II*.

2.3.3 EtherType

Než se zaměříme na strukturu rámce, řekneme si něco o identifikačních kódech určujících, se kterým protokolem vlastně komunikujeme, tedy co konkrétně má být do rámce zapouzdřeno.

 *EtherType* je identifikátor určující typ dat, která jsou zapouzdřována do rámce, většinou určuje protokol, který na síťové vrstvě vytvořil odesílaný blok dat. Zapisuje se čtyřmi hexadecimálními číslicemi, takže maximální hodnota je $(FFFF)_{16}$ (zapisujeme $0 \times FFFF$), to je v dekadické soustavě 65 535. Ve skutečnosti existuje i spodní limit – pro EtherType se používají čísla o minimální hodnotě 0×0600 , to je 1536. To znamená, že EtherType je vždy z rozsahu $0 \times 0600 \dots 0 \times FFFF$.



Poznámka:

V záhlaví rámce vytvářeného na vrstvě L2 jsou dva speciální oktety (odtud maximum $0 \times FFFF$), do kterých se obvykle EtherType ukládá. Jenže totéž pole může být použito i pro jiný typ údaje – délku rámce. Ovšem v technice musí být vždy vše naprosto jednoznačné, proto musí být možné rozlišit, zda v daném poli je EtherType nebo délka rámce.

Délka SDU zapouzdřeného do ethernetového rámce je (téměř) vždy menší než $0 \times 05DC$ (dekadicky 1500), přičteme rezervu (zaokrouhlíme hexadecimálně nahoru) a získáme spodní mez pro hodnotu EtherType 0×0600 .

Pokud je nutné přenést velmi velký rámec (nad stanovený limit), pak se jedná o *jumbo rámec* (jumbo frame) a pro ten je v daném poli speciální hodnota.



Hodnot EtherType je velmi mnoho, podíváme se jen na několik:

- 0×8870: jumbo frames
- 0×8100: VLAN rámce podle 802.1Q
- 0×0800: IPv4
- 0×86DD: IPv6
- 0×0806: ARP
- 0×0835: RARP
- 0×08137, 0×08138: IPX
- 0×8847: MPLS unicast
- 0×8914: FCoE Initialization Protocol
- 0×814C: SNMP



Poznámka:

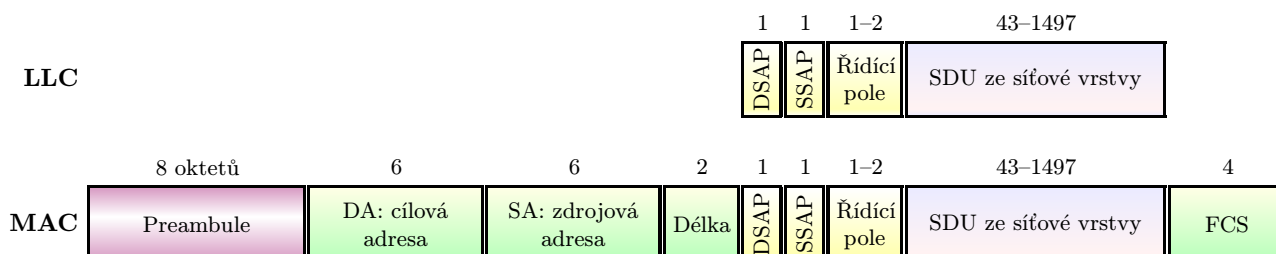
Z názvu „EtherType“ by se mohlo zdát, že tento identifikátor je používán jen v Ethernetu. Sice jako první byl v Ethernetu použit, ale ve skutečnosti se s ním setkáváme i v jiných typech sítí.



2.3.4 Formát ethernetového rámce

Nyní se podíváme, jak vlastně vypadá rámec posílaný přes síť Ethernet. Postupně si projdeme všechny tři možnosti. Formát rámce (a později formát různých dalších PDU) budeme reprezentovat formou tabulky, kde v záhlaví je naznačena délka jednotlivých polí.

LLC rámec. Struktura LLC rámce (podle protokolu IEEE 802.2) je jednoduchá – k SDU přijaté z nadřazené vrstvy se přidá záhlaví se třemi poli.



Obrázek 2.3: Rámec podle IEEE 802.3/802.2: LLC rámec zapouzdřený v MAC rámci podle 802.3

Na obrázku 2.3 je naznačen formát rámce v případě, že byly zkombinovány protokol IEEE 802.2 na podvrstvě LLC (vznikl LLC rámec, jeho záhlaví je žluté) a protokol IEEE 802.3 na podvrstvě MAC (LLC rámec jsme zapouzdřili do MAC rámce podle 802.3, záhlaví a zápatí MAC rámce je zelené a fialové). Jednotlivá pole v LLC rámci (žlutém) mají tento význam:

- *SSAP*, *DSAP* (po 1 oktetu) – přístupové body SAP na cílovém (*DSAP*) a zdrojovém (*SSAP*) zařízení, které zabírají 7 bitů z každého oktetu; zbývající (nejméně významný, vpravo) bit značí:
 - u cílového údaje I/G bit (individual/group) určující, zda jde o skupinovou adresu SAP,
 - u zdrojového údaje C/R bit (command/response) pro typ paketu (příkaz nebo odpověď),
- *řídicí pole* (většinou 1 oktet) – určuje typ rámce z hlediska služby, pořadové číslo apod.,
- data z vyšší vrstvy (*SDU*), která jsou v tomto rámci zapouzdřena.

Vraťme se k polím *DSAP* a *SSAP*. Jejich hodnoty jsou standardizovány, jmenujme si některé z nich:

- 0×06: DoD IP
- 0×42: IEEE 802.1 Bridge STP
- 0×98: Arpanet ARP
- 0×E0: Novell Netware (IPX)

Některé další hodnoty jsou používány například pro management LLC (tyto rámce jsou služební, s nadřazenou vrstvou nemají nic společného), další jsou vyhrazeny některým již nepoužívaným protokolům.



Příklad

Pokud LLC rámec obsahuje tyto údaje:

0×42	0×42	0×03	data
------	------	------	------

Pak to znamená, že na zdrojovém i cílovém zařízení vrstva L2 (její podvrstva LLC) komunikuje s protokolem STP (Spanning Tree Protocol). Číslo 0×42 je binárně 1000010, tedy I/G bit a C/R bit mají hodnotu 0 (individuální SAP, příkaz). Uvnitř (jako „data“) je PDU protokolu STP.



MAC rámec podle IEEE 802.3 přidává tato pole (na obrázku 2.3 zeleně a fialově):

- *Preamble* (8 oktetů) je identifikace začátku rámce, synchronizační informace. Prvních 7 oktetů preamble obsahuje střídající se jedničky a nuly, osmý oktet taky, až na poslední bit – ten je místo na nulu nastaven na jedničku. Celá preamble tedy vypadá takto:

10101010 10101010 10101010 10101010 10101010 10101010 10101010 10101011

V literatuře se setkáme taky s jinou charakteristikou – preamble na 7 oktetech se střídajícími se jedničkami a nulami, a pak jeden oktet, který toto pravidlo ve svém nejméně významném bitu porušuje. Tento oktet se nazývá SOF nebo SD (Start-of-frame Delimiter).

- *Destination Address* (DA, adresa cíle) a *Source Address* (SA, adresa odesílatele) jsou MAC adresy komunikujících uzlů v síti. Zdrojová adresa musí být vždy unicast, cílová může být unicast, multicast nebo broadcast.
- *Délka* vnořeného LLC rámce (takže délka dat z vyšší vrstvy plus LLC záhlaví).
- *FCS* (Frame Check Sequence) je kontrolní součet, který se vypočítává z polí adres (DA, SA), délka, záhlaví LLC a data z vyšší vrstvy (tj. ze všech kromě preamble, která je vždy stejná, a kromě FCS), plus obdoba preamble pro detekci konce rámce.



Kontrolní součet v rámci. Pro FCS se v Ethernetu používá algoritmus CRC-32, který se dá celkem jednoduše implementovat hardwarově, v obvodech. Výpočet je založen na operacích v Galoisově poli $GF(2^n)$, kde se jako operace využívá dělení polynomů nad binárními čísly modulo 2 (tj. zbytek po dělení dvěma).

Algebraické pole modulo dvěma funguje tak, že když jako výsledek operace vyjde číslo 2 nebo vyšší, nahradí se zbytkem po dělení dvěma. Například $1+1$ není dvě, ale nula. A navíc místo čísel pracujeme s polynomy.

Například polynom stupně 2 nad binárními čísly má tvar $a_2 * x^2 + a_1 * x^1 + a_0 * x^0$, kde koeficienty a_2 , a_1 , a_0 jsou binární číslice (0 nebo 1), $x^0 = 1$, tj. $a_2 * x^2 + a_1 * x + a_0$, například $1 * x^2 + 0 * x + 1$, respektive $x^2 + 1$ je polynom nad binárními čísly stupně 2. My potřebujeme polynom stupně 31, tj. s mocninami 31 až 0 (to znamená 32 členů polynomu, potřebujeme 32 binárních číslic pro koeficienty).

Polynomy se dají dělit stejně jako čísla (ovšem zde zásadně tak, abychom pořád zůstali u binárních koeficientů), a pokud jsou nesoudělné, získáme zbytek po dělení (tak jako u běžných čísel při dělení čísla 15 dvojkou dostaneme zbytek 1). Zbytek po dělení je opět binární polynom, a ten můžeme převést na posloupnost bitů tak, že u členů polynomu s určitou mocninou vezmeme koeficienty (0 nebo 1) a zřetězíme

je do 32bitového binárního čísla. Takže máme smyčku: binární číslo převedeme na polynom, vydělíme polynomem a zbytek převedeme zpět na binární číslo.

Postup výpočtu kontrolního součtu:

1. Posílané bity od pole pro cílovou adresu po data (tj. všechno mezi poli preamble a FCS) se rozdělí na posloupnost 32bitových slov (čísel), od toho CRC-32.
2. Tato 32bitová slova se dále chápou jako polynomy, kdy jednotlivé bity (0 nebo 1) určují, jestli určitý člen polynomu má nebo nemá být použit (0 – ne, 1 – ano). Například z posloupnosti 32 binárních čísel 1000111010000....0001 uděláme polynom $x^{31} + x^{27} + x^{26} + x^{25} + x^{23} + x + 1$.
3. Polynom z prvních 32 bitů vydělíme speciálním polynomem stupně vyššího než 31, kterému říkáme „ireducibilní“, v Galoisově poli binárních polynomů to je obdoba prvočísel, která jsou větší než maximální hodnota pro dělence (než to číslo, které dělíme). Na rozdíl od aritmetických operací s běžnými čísly zde dostaneme jako zbytek po dělení binární polynom stupně nejvýše 31.
4. Vezmeme polynom z následujícího 32bitového slova a přičteme polynom, který je zbytkem po dělení předchozího 32bitového slova. Součet opět vydělíme ireducibilním polynomem. Vezmeme polynom z dalšího 32bitového slova, přičteme zbytek předchozího dělení, vydělíme... atd. postupně přes všechna 32bitová slova.
5. Když dojdeme na konec dat a dostaneme poslední zbytek po dělení, tento zbytek (polynom) převedeme na 32bitové binární slovo přesně opačným postupem, než jak jsme vytvářeli polynomy – jednotlivé koeficienty 0,1 u členů polynomu (v pořadí od stupně 31 do stupně 0) vezmeme, zřetězíme a dostaneme 32bitové číslo. To uložíme do pole FCS.

Galoisovo pole $GF(2^n)$ má oproti běžným číslům jednu zvláštnost: když vezmeme dva polynomy a sečteme je, a tytéž dva polynomy od sebe odečteme, v obou případech získáme stejný výsledek. Taže sčítání a odčítání jsou zde stejná operace. Toho se využívá při kontrole kontrolního součtu. Takže když rámec prochází switchem, který opravdu pracuje s kontrolním součtem (tj. přepíná rámce metodou store-and-forward) nebo už přichází do koncového zařízení a je třeba zjistit, jestli rámec nebyl cestou poškozen, postupuje se takto:

1. Začneme úplně stejně jako při výpočtu CRC32, tj. bereme bity od adres včetně dál přes data po 32bitových slovech, postupně vytváříme binární polynomy, dělíme ireducibilním polynomem, zbytek přičteme k polynomu z dalších 32 bitů, vydělíme...
2. Místo abychom se zastavili na konci dat, sečteme výsledek posledního dělení s obsahem pole FCS (už nedělíme). Protože v předchozím algoritmu jsme na konci dat dostali binární číslo, které jsme uložili do FCS, teď bychom vlastně měli sčítat dvě stejné hodnoty. Protože sčítání a odčítání jsou zde shodné operace, je to totéž, jako když od sebe dvě stejná čísla odečítáme, tj. měli bychom dostat nulu.
3. Vyhodnocení: pokud je rámec ok, pak by výsledkem z bodu 2 měla být nula, ale pokud je rámec poškozený, vyjde něco jiného.

Výhodou je, že při kontrole, zda je rámec v pořádku, nemusíme nejdřív načíst rámec a pak se vracet na začátek, abychom spustili algoritmus, můžeme počítat souběžně s přijímáním rámce z rozhraní. Přijímáme, ukládáme do bufferu a speciální čip průběžně počítá. Po přijetí posledního 32bitového slova (což je FCS) stačí zkontrolovat výsledek. Když čip došel k výsledku 0, je rámec ok a přijmeme ho. Když došel k jinému výsledku, rámec zahodíme.

**Další informace:**

Další informace se dají najít na internetu, například

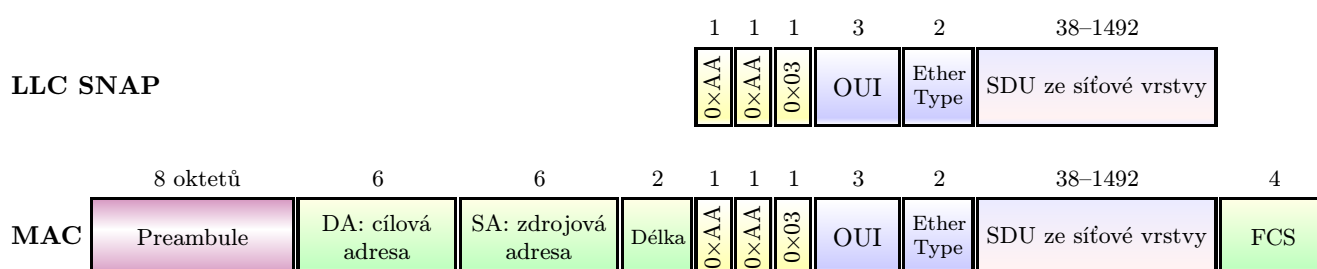
<https://www.autohotkey.com/boards/viewtopic.php?f=7&t=35671>

Na různých variantách Galoisova pole jsou založeny současné šifrovací algoritmy včetně známého AES, takže příslušné postupy najdete i v knihách o kryptografii, například:

OULEHLA, Milan a Roman JAŠEK. *Moderní kryptografie*. Praha: IFP Publishing, 2017. ISBN 978-80-87383-67-4.



SNAP rámec. SNAP (SubNetwork Access Protocol) vznikl rozšířením protokolu IEEE 802.2 LLC, pozměněním a prodloužením záhlaví podvrstvy LLC. SNAP rámec vypadá takto:



Obrázek 2.4: Rámec IEEE 802.3/SNAP: SNAP rámec zapouzdřený v MAC rámci podle 802.3

Na obrázku 2.4 jsou pole přidána rozšířením SNAP vybarvena modře. Jak vidíme, SNAP dosazuje do záhlaví protokolu LLC určité konstantní hodnoty – jako DSAP a SSAP se použije kód 0xAA (nebo je přípustný kód 0xAB) a do řídicího pole se uloží 0x03. Informaci, kterou by nám jinak tato pole dávala (tedy konkrétní protokoly na nadřazené vrstvě), najdeme právě v přidávaných modrých polích.

Rozšíření SNAP přidává k LLC tato pole (na obrázku 2.4 modře):

- *OUI* (Organizationally Unique ID, organizace) je většinou nastaveno na 0. Pokud ne, pak se jedná o identifikátor organizace, v rámci jejíž zařízení je tento rámec posílán. Například Cisco má $OUI = 0x00\ 00\ 0C$.
- *EtherType* (nebo obecněji Typ) určuje protokol nadřazené vrstvy, pokud je $OUI = 0$. Pokud není, pak je toto pole čistě v režii organizace, jejíž OUI je v předchozím poli, například pro přenos paketů podle proprietárních síťových protokolů.

O dvě stránky výše je seznam několika běžných hodnot pro pole SSAP a DSAP v LLC rámci. K těmto hodnotám si můžeme přidat ještě 0xAA jako identifikátor rozšíření SNAP.

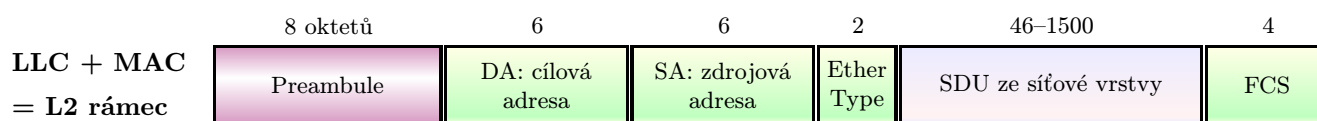
**Poznámka:**

Jak poznáme, jestli se jedná o rámec IEEE 802.3/LLC nebo IEEE 802.3/SNAP?

Několik polí (záhlaví MAC rámce podle 802.3) je stejných, odlišnost začíná na začátku vnořeného LLC/SNAP rámce. Pokud v „žluté části“, resp. třech oktetech za MAC záhlavím, najdeme hodnotu 0xAAA03, pak je jasné, že jde o SNAP rámec. Jinak jde o LLC rámec.



Rámec Ethernet II. V současné době se v Ethernetu používá téměř výhradně tento typ rámce, a není divu – je nejjednodušší a nejefektivnější.



Obrázek 2.5: Rámec Ethernet II

Na obrázku 2.5 vidíme, že struktura rámce je vlastně velmi podobná tomu, co v předchozích případech na MAC podvrstvě dodával protokol IEEE 802.3. Pole pro preamble, obě MAC adresy a kontrolní součet mají stejný význam. Rozdíl je jen v poli, ve kterém byla v předchozích případech délka rámce – v rámci Ethernet II tam najdeme hodnotu EtherType. Proč?

- Evidovat délku rámce vlastně ani nepotřebujeme. Konec rámce jednoduše poznáme tak, že na přenosovém médiu nebude žádný signál (přenáší se v základním pásmu, kdy je signál na médiu generován, nikoliv modulován existující). Navíc je součástí zápatí obdoba synchronizační informace z preamble, třebaže kratší.
- Naopak nutně potřebujeme vědět, co je zapouzdřeno uvnitř rámce, aby bylo jasné, kterému protokolu vyšší vrstvy má být přenášéný paket v cíli předán.

Proto místo pole s délkou máme pole s typem protokolu, tedy EtherType, o kterém byla řeč dříve.



Poznámka:

Podle čeho poznáme, že se jedná o rámec Ethernet II a ne žádný z předchozích?

V předchozí sekci je psáno o hodnotách a významu pole EtherType. Kromě jiného se tam v jedné poznámce dočteme, že EtherType a délka rámce jsou ukládány do téhož pole, přičemž číslo větší než 00600 znamená EtherType, menší znamená délku rámce. Takže přepínač během načítání prvních $8 + 6 + 6 = 20$ oktetů ještě typy rámců nedokáže rozlišit, ale hned podle 21. a 22. oktetu už dokáže rozlišit, zda se jedná o Ethernet II.



Kdy se který typ rámce používá:

- rámce Ethernet II se používají pro běžné datové přenosy (uvnitř tedy většinou bývají pakety IPv4 nebo IPv6),
- rámce LLC jsou obvykle určeny pro jiné než datové přenosy (například STP nebo jiné standardizované protokoly),
- rámce SNAP jsou výhodné zejména pro provoz proprietárních protokolů, kdy neexistuje „standardní“ číslo zapouzdřeného protokolu.



Poznámka:

Pokud se někde uvádí skutečná, minimální či maximální délka rámce, preamble do ní obvykle nepočítáme. Když se vrátíme k nákresům struktury jednotlivých typů rámců a sečteme čísla uvedená nad jednotlivými poli (délky těchto polí) kromě preamble, zjistíme, že ve všech případech je minimální délka rámce $46 + 18 = 64$ oktetů, maximální délka je $1500 + 18 = 1518$ oktetů.



2.3.5 Tabulka MAC adres na switchi

Přepínací tabulky s MAC adresami zařízení mohou různí výrobci říkat různě, také se setkáváme s názvem CAM (Content Addressable Memory) tabulka.

V tabulce najdeme k jednotlivým zařízením především tyto údaje:

- MAC adresu zařízení,
- port, na kterém je zařízení dostupné,
- další informace (VLAN, časové razítko apod.).

Nás zatím budou zajímat první dva údaje. Shrňme si, jak switch zpracovává příchozí rámce.

- Pokud je adresátem rámce zařízení evidované v tabulce MAC adres, switch najde řádek tabulky s MAC adresou příjemce, na tomto řádku zjistí port, ke kterému je adresát připojen, a na ten port rámec pošle.
- Pokud je adresa příjemce neznámá (není v tabulce), je rámec odeslán na všechny porty kromě toho, ze kterého přišel. V tom případě se switch chová jako hub.
- Pokud je cílová adresa broadcast, je rámec poslán na všechny porty kromě toho, ze kterého přišel.

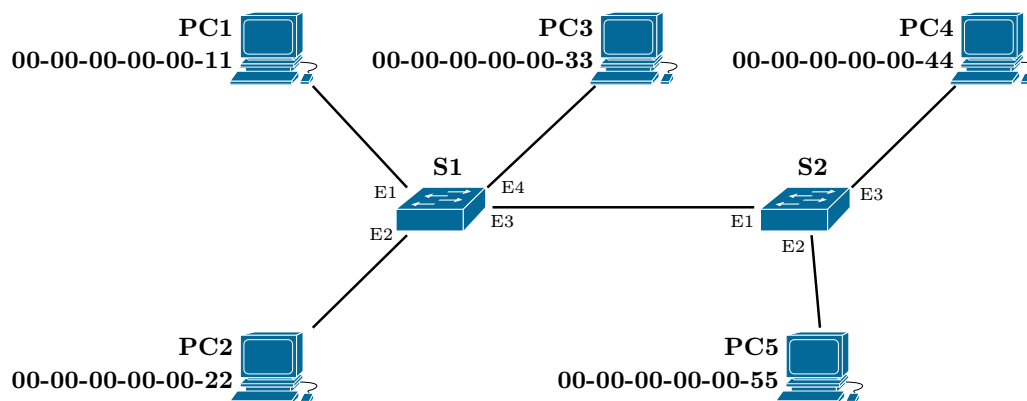
Ale jak se vlastně dotyčný adresát do té tabulky dostane?

 Předpokládejme, že máme nový (nebo restartovaný) switch a poprvé ho zapojíme, tabulka je prázdná. V takovém případě jsou dvě možnosti, jak tabulku naplnit. Buď jsou do ní příslušné údaje „natvrdo“ přidány pomocí ruční konfigurace (či pomocí skriptu), nebo to jednoduše necháme na dotyčném switchi, ať se tedy v síti sám rozkouká a „seznámí se se sousedy“.

Ethernetové switche používají tento druhý způsob. Kdykoliv je doručen rámec, switch se zajímá nejen o cílovou adresu (ta je důležitá, aby věděl, kam rámec poslat), ale taky o zdrojovou adresu. U zdrojové adresy je víceméně jasné, na kterém portu je dotyčné (zdrojové) zařízení dostupné – na tom portu, ze kterého právě přišel rámec. Tedy pokud adresu zdroje v tabulce nemá, přidá ji s informací o portu, ze kterého rámec přišel.

Příklad

Podívejme se na následující obrázek. Jsou na něm dva switche, z nichž první má aktivní čtyři porty, druhý tři. Předpokládejme, že switch S1 má svou tabulku zatím prázdnou.



Následuje tento provoz (zkratka DA je Destination Address, tedy cílová adresa, SA je Source Address, tedy zdrojová adresa), vše na switchi S1:

- Na portu E2 je přijat rámeček, v němž jsou tyto adresy:

- DA = 00-00-00-00-00-33
- SA = 00-00-00-00-00-22

⇒ switch usoudí, že zařízení s MAC adresou 00-00-00-00-00-22 je na portu E2, tabulka:

MAC adresa	Port	...
00-00-00-00-00-22	E2	...

adresáta (adresu 00-00-00-00-00-33) nezná, takže rámeček přepošle na porty E1, E3 a E4.

- Na portu E3 je přijat rámeček, v němž jsou tyto adresy:

- DA = 00-00-00-00-00-11
- SA = 00-00-00-00-00-55

⇒ zařízení s MAC adresou 00-00-00-00-00-55 je na portu E3, tabulka:

MAC adresa	Port	...
00-00-00-00-00-22	E2	...
00-00-00-00-00-55	E3	...

adresáta (adresu 00-00-00-00-00-11) nezná, takže rámeček přepošle na porty E1, E2 a E4.

- Na portu E1 je přijat rámeček, v němž jsou tyto adresy:

- DA = 00-00-00-00-00-22
- SA = 00-00-00-00-00-11

⇒ zařízení s MAC adresou 00-00-00-00-00-11 je na portu E1, tabulka:

MAC adresa	Port	...
00-00-00-00-00-22	E2	...
00-00-00-00-00-55	E3	...
00-00-00-00-00-11	E1	...

adresáta (adresu 00-00-00-00-00-22) už v tabulce máme (hned na prvním řádku), takže rámeček přepošleme pouze na takto zjištěný port E2.

- Na portu E1 je přijat rámeček, v němž jsou tyto adresy:

- DA = 00-00-00-00-00-44
- SA = 00-00-00-00-00-11

⇒ zařízení s MAC adresou 00-00-00-00-00-11 už v tabulce máme, takže teď do tabulky nebudeme nic přidávat,

adresáta (adresu 00-00-00-00-00-44) nezná, takže rámeček přepošle na porty E2, E3 a E4.

Tak bychom pokračovali i dál, MAC tabulka by byla kompletní až tehdy, když by se postupně „prozradila“ všechna zařízení v síti.



2.4 Ethernet na fyzické vrstvě

Na fyzické vrstvě v Ethernetu jsou zajištěny tyto operace:

- spolupráce s vrstvou L2, přebírání SDU z této vrstvy (rámeček), v opačném směru předávání dat vrstvě L2,

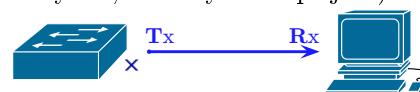
- kódování, multiplexing,
- vysílání a přijímání signálu, synchronizace se zařízením „na druhé straně“,
- auto-negociace (vyjednávání) – síťové rozhraní si potřebuje dojednat se síťovým rozhraním druhého zařízení veškeré přenosové parametry (rychlost, poloviční nebo plný duplex apod.).

2.4.1 Křížení a krimpování

Zařízení na svém síťovém rozhraní vlastně provádí dvě základní operace:

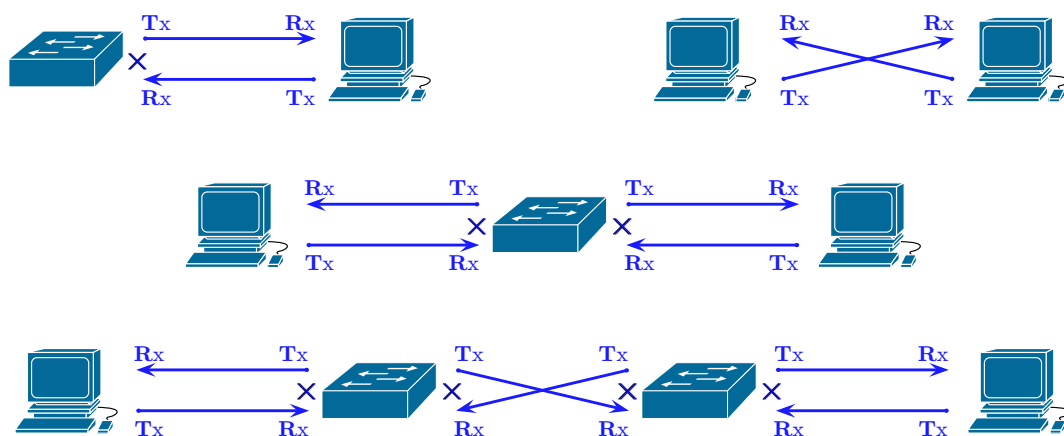
- Tx (Transmit) – odesílá,
- Rx (Receive) – přijímá (tj. naslouchá a když zjistí, že druhá strana vysílá, toto vysílání přijme).

V reálu to znamená, že například u kroucené dvojlinky (twisted pair) jsou páry vodičů vyhrazené pro operaci Tx a páry vodičů vyhrazené pro operaci Rx, případně je toto přiřazení určováno dynamicky (ale



vždy platí, že v jednom okamžiku je toto přiřazení jednoznačné) – daný pár vodičů nemůže v jednom okamžiku provádět totéž. Jinými slovy – síťové rozhraní se skládá ze dvou částí – Tx a Rx, a je důležité, na kterou z těchto částí je který pár v kabelu napojen. U dvojice přímo komunikujících zařízení (tj. propojených tímto kabelem) musí platit, že na konkrétním páru jedno zařízení přijímá a druhé vysílá – kdyby na tomtéž páru obě zařízení odesílala, došlo by ke kolizi, a kdyby na tomtéž páru obě zařízení naslouchala, k žádnému přenosu by naopak vůbec nedošlo. Vysílající zařízení odešle signál do části Tx svého síťového rozhraní, a cílové zařízení tento signál přijme na části Rx svého síťového rozhraní, a tyto dvě části musí být fyzicky propojeny tímto párem vodičů.

Z toho vyplývá, že někde musí být přenosová cesta *překřížena*. Pokud jsou dvě zařízení propojena kroucenou dvojlinkou, pak (minimálně jeden) pár vodičů vedoucí z Tx jednoho zařízení musí být napojen za Rx druhého zařízení a naopak.



Obrázek 2.6: Princip křížení na kroucené dvojce

Na obrázku 2.6 vidíme několik řešení – v „první řadě“ křížení zajistí jedno ze zařízení (switch, bridge) nebo je přímo v kabelu – křížení zajištěné zařízením je naznačeno křížkem. Níže jsou případy, kdy máme jedno nebo dvě mezilehlá zařízení (která zajišťují křížení).

 Z obrázku vyplývá, že existují tři typy křížení: v kabelu, v rozhraní, a pak přepínací prvky (switche) kříží datový tok na logické úrovni (L2), tj. to, co přijmou na RX, dále pošlou na TX. Na celé cestě mezi koncovými zařízeními je vždy lichý počet křížení.

 Rozlišujeme dva typy kabelů typu kroucená dvojlinka:

- *přímý kabel* – žádné křížení neobsahuje, používá se obvykle k propojení DTE a DCE (tedy například počítače ke switchi),
- *křížený kabel* – obsahuje křížení, používá se k propojení zařízení stejného typu (tedy buď DTE a DTE, nebo DCE a DCE); pokud chceme propojit dva počítače přímo bez mezilehlého DCE, měli bychom použít křížený kabel.

Ve skutečnosti se v obou případech používá tentýž kabel; to, zda je přímý nebo křížený, se určuje nasazením koncovek. Přímý kabel má na obou stranách koncovky nasazené stejně (vodiče jsou na obou koncích ve stejném pořadí), kdežto křížený kabel má na jednom konci vodiče jinak seřazené než na druhém konci.



Poznámka:

V současné době máme situaci jednodušší – novější síťová rozhraní většinou dokážou automaticky rozpoznat, jestli mají nebo nemají zajistit vnitřní křížení, takže prostě použijeme přímý kabel.



 *Krimpování (konektorování)* je postup nasazování konektoru na kabel. V případě kroucené dvojlinky se většinou používají konektory a zásuvky typu *8p8c* (to znamená 8 pozic a 8 kontaktů), kterým se poněkud nepřesně říká RJ-45.



V konektoru (i zásuvce) jsou vodiče umístěny v řadě za sebou a jejich pořadí je určeno standardem TIA/EIA-568-B z roku 2001, nověji TIA/EIA-568-C z roku 2014. Standard ve skutečnosti určuje dvě pořadí – T568A a T568B (pozor, neplést si písmenka s verzí standardu), přičemž u přímého kabelu se na obou koncích použije pořadí podle T568B, kdežto u kříženého na každém konci jedno z těchto pořadí. Pozor, nejde jenom o obrácení pořadí!

V případě optických vláken je křížení zajištěno vždy v kabelu (zde jde o pořadí optických vláken).

2.4.2 Parametry

 Naprostá většina druhů Ethernetu používá přenos v základním pásmu (baseband). Dalšími důležitými parametry jsou rychlost a typ přenosového média (kabelu), obojí má vliv na použité kódování a průběh multiplexování. Na fyzické vrstvě se v označení objevuje:

- rychlost: 10, 100, 1000 apod. – v Mb/s, u vyšších se píše s písmenem G, např. 10G,
- přenos v základním (base) nebo přeloženém (broad) pásmu,
- typ kabelu – většinou „T“ znamená kroucenou dvojlinku (twisted), „F“ optický.

Takže pojďme od historie, ale přeskočíme nejstarší standardy. Nebudeme se zabývat všemi, u každé rychlosti vybereme jen některé typické zástupce.

 **10Mb Ethernet.** Všichni zástupci dovolují rychlost 10 Mb/s, konkrétně:

- 10Base-T: baseband, kabel nestíněná kroucená dvojlinka (UTP),
- 10Base-F: baseband, kabel optický,
- 10Broad-36: broadband (přeložené pásmo), kabel koaxiál o impedanci 75 Ω; nerozšířil se, ale stal se základem pro pozdější Ethernet provozovaný v sítích kabelových televizí (dnes se místo toho v kabelových televizích používá pro přenos dat standard DOCSIS).

 **Fast Ethernet.** Rychlost 100 Mb/s a baseband, přibyla možnost auto-negociace (síťová rozhraní se dokážou automaticky dohodnout na parametrech přenosu). Zástupci:

- 100Base-TX: kabel UTP alespoň Cat5, využívá pouze dva páry (jeden pro vysílání a druhý pro příjem); v podstatě pokračovatel 10Base-T, v síti bylo možné kombinovat síťové karty těchto typů,
- 100Base-FX: používá dvě optická vlákna, jedno pro každý směr. Není kompatibilní s 10Base-F, protože používá signál o jiné vlnové délce.

 **Gigabit Ethernet** umožňuje rychlost 1 Gb/s, baseband. Zástupci:

- 1000Base-T: pro kabel UTP minimálně Cat.5 (doporučeno minimálně 5e), používají se všechny čtyři páry (takže žádný light se dvěma páry nelze použít), směr pro jednotlivé páry se určuje adaptivně – žádný není vyhrazen pro konkrétní směr,
- 1000Base-TX: jedná se o nepříliš úspěšný pokus o úpravu standardu 1000Base-T tak, že pro každý směr jsou vyhrazeny dva páry a je vyžadován kabel Cat.6; zatímco 1000Base-T je standard od organizace IEEE (IEEE 802.3ab), 1000Base-TX pochází od TIA (TIA/EIA-854); komerčně neúspěšný – nelze použít na starších rozvodech Cat.5 nebo 5e,
- 1000Base-X: souhrn standardů pro (většinou) optická vlákna
 - 1000Base-SX („S“ jako short) optická vlákna mnohavidová, typický dosah je ve stovkách metrů,
 - 1000Base-LX („L“ jako long) optická vlákna buď jednovidová (dosah v jednotkách kilometrů) nebo mnohavidová (stovky metrů),
 - 1000Base-CX je výjimka z pravidla – nepoužívá optiku, ale stíněnou kroucenou dvojlinku, konektory jsou podobné jako u UTP, ale jiné přiřazení pinů (jinak se krimpuje); typický dosah je do 25 metrů, používá se v datových centrech k připojení serverů ke switchům (například u některých produktů společnosti IBM).



Poznámka:

Velmi často se setkáváme se záměnou názvu – místo 1000Base-T je udáno 1000Base-TX. Obvykle to je proto, že u „dřívější generace“ (tedy Fast Ethernetu) byla přípona TX. U síťových rozhraní podporujících různé rychlosti bývá označení 100/1000Base-TX, třebaže u druhé uvedené rychlosti jde ve skutečnosti o „Téčkový“ standard.



 **10Gigabit Ethernet** (10GE, 10GigE) umožňuje rychlost 10 Gb/s, přenos baseband. Přístupová metoda CSMA/CD je zcela odbourána, komunikuje se pouze ve full-duplexu.

- 10GBase-R je skupina standardů pro optická vlákna:
 - 10GBase-SR (short range) pro mnohavidové vlákno, dosah do 62 m,
 - 10GBase-LR (long range) pro jednovidové vlákno, dosah cca 10 km,
 - 10GBase-LRM (long reach multimode) pro mnohavidové vlákno, dosah do 220 m,
 - 10GBase-ER (extended range) pro jednovidové vlákno, dosah 40 km,

liši se nejen kabely, ale také vlnovými délkami, na kterých je vyslán signál,

- 10GBase-CX4 používá kabel twin-ax (podobně jako koaxiál, jen místo jednoho středového vodiče jsou dva), dosah je jen 15 metrů, ale na druhou stranu má velice dobré hodnoty odezvy (latence) a cena je vcelku příznivá; používá se v datových centrech pro připojení serverů k DCE,

- 10GBase-T je pokračovatelem 1000Base-T, používá kroucenou dvojlinku minimálně Cat.6; doporučuje se však použít vyšší kategorii, protože na Cat.6 má dosah jen max. 55 m,
- 10GBase-W už míří mimo LAN sítě – na vrstvu L1 přidává novou podvrstvu WIS, která umožňuje komunikovat s WAN sítí (SONET/SDH, viz dále); funguje podobně jako 10GBase-R, jen navíc ethernetový rámec zapouzdří do WIS rámce, ve kterém data procházejí napojenou WAN sítí.

**Poznámka:**

Zvědavý čtenář se určitě ptá: a kde jsou tedy ty standardy? Tak například 100Base-TX a 100base-FX (tedy kroucená dvojlinka a optika pro Fast Ethernet) jsou standardizovány v IEEE 802.3u, kdežto 1000Base-T (kroucená dvojlinka pro Gigabit Ethernet) je v standardu IEEE 802.3ab a Gigabit Ethernet na optice najdeme v IEEE 802.3z.

Standard IEEE 801.3 a jeho dodatky (písmenka na konci) definují vrstvu L1 a spodní část vrstvy L2 (MAC podvrstvu), přičemž se v praxi využívá spíše to, co platí pro vrstvu L1 (už dřív jsme si vysvětlili, že na L2 se setkáváme spíše s rámcem podle Ethernet II).



40Gigabit Ethernet a 100Gigabit Ethernet (40GigE, 100GigE) předepisuje optické kabely s laserem jako zdrojem světla (na vzdálenost cca 100 m pro mnohavidová vlákna, resp. kilometry a desítky kilometrů pro jednovidová s různými parametry), na vzdálenost několika metrů kabel twinax (v datových centrech pro připojení výkonných serverů do sítě), a na vzdálenost do 30 m lze pro rychlost 40GigE použít kroucenou dvojlinku Cat.8.



2.5 a 5Gigabit Ethernet byly jako standardy zveřejněny až v roce 2016, tj. po „rychlejších“ sourozencích. Jedná se o komunikaci po kroucené dvojlince, formálně 2.5GBase-T a 5GBase-T, standard IEEE 802.3bz. Fyzická vrstva byla převzata z 10GBase-T a přizpůsobena „horší“ kabeláži. 2.5G-Base-T může využívat nestíněnou kroucenou dvojlinku kategorie Cat.5e, 5GBase-T používá kategorii Cat.5e (menší vzdálenost než 100 m) nebo Cat.6.

Zatímco 10GBase-T nepodporuje PoE, 2.5GBase-T a 5GBase-T by měly napájení přes Ethernet zvládat.

**Poznámka:**

Jak je možné řádově zvyšovat rychlost, když se pořád používají v podstatě tytéž kabely (i když generaci od generace lepší)? Zvýšení se dosahovalo kombinací více způsobů, například:

- vytížením všech čtyř párů kroucené dvojlinky místo pouhých dvou,
- použitím kvalitnější kabeláže s lepším stíněním, kroucením s vyšší hustotou a větší šířkou přenosového pásma,
- použitím lepšího kódování dat na signál (v reálu se přenáší menší objem dat).



2.4.3 Implementace fyzické vrstvy

Fyzická vrstva je implementována typicky v obvodech, a to buď ve formě samostatného čipu (na základní desce nebo na síťové kartě NIC – Network Interface Card) nebo integrovaná v čipu s jinými komponentami. Může se skládat ze dvou částí.

 *MAU* (Medium Attachment Unit) neboli *transceiver* (zkráceno ze slov transmitter/receiver) je prvek na NIC, který zajišťuje rozpoznání přítomnosti signálu, kolizi (signál jam) a vysílání/příjem signálu. Může být

- externí – u starého Ethernetu na koaxiálu a naopak u nových optických modulů (XENPAK, XFP, SFP, SFP+ apod.),
- integrovaný v čipu s ostatními částmi síťového rozhraní.

Na síťových prvcích (switche, routery) se můžeme setkat s pozicemi pro SFP+ moduly, které mohou být optické nebo metalické, s různou specifikací.



Obrázek 2.7: SFP+ transceivery pro 10GBase-LR/LW¹

 Fyzická vrstva se také dělí na podvrstvy, a to závislé/nezávislé na médiu s vlastními komunikačními rozhraními (interface).

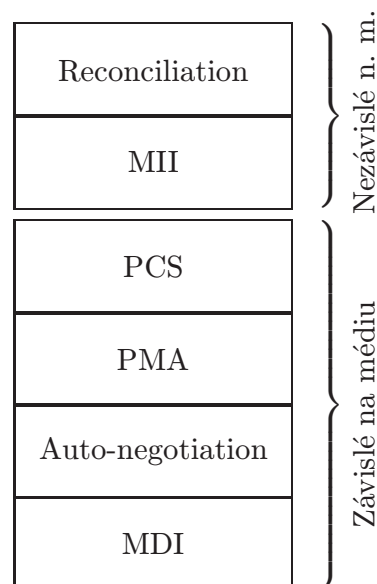
Podvrstvy nezávislé na médiu:

- vyrovnávací podvrstva (reconciliation),
- MII (10Mb a 100Mb), GMII (Gb) nebo XGMII (vyšší) – zvlášť odesílací a přijímací datové cesty o šířce 1 bit pro 10Mb (bit-serial), 4 bity pro 100Mb (nibble-serial) nebo širší.

Zajišťují logické spojení mezi MAC a různými sadami podvrstev závislých na médiu.

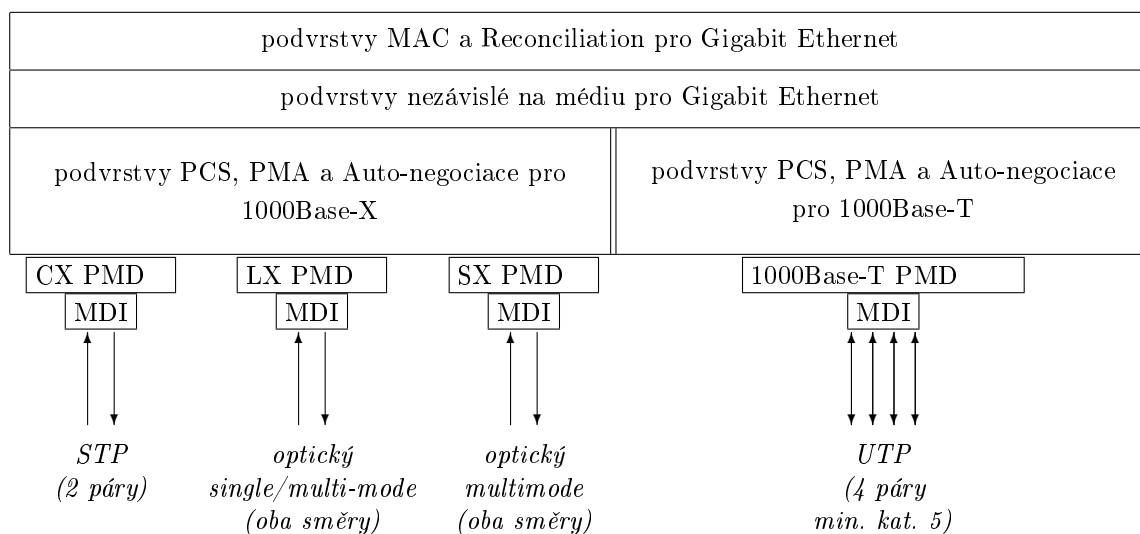
Podvrstvy závislé na médiu:

- PCS (Physical Coding Sublayer) – kódování, multiplexing, synchronizace odchozích a opačně u příchozích dat,
- PMA (Physical Medium Attachment) – vysílání a přijímání signálu, mechanismus zotavení časovače (zajištění sesynchronizování časovače na začátku proudu dat),
- Auto-negotiation (vyjednávací podvrstva) – tyto podvrstvy na začátku přenosu dohodnou konkrétní vlastnosti spojení vyhovující oběma NIC,
- MDI (Medium-Dependent Interface) – ke konektoru kabelu.



Příklad schématu vrstev (zde pro 1Gb Ethernet) vidíme na obrázku 2.8 na straně 46. Fyzická vrstva je zcela předefinována ve standardech od 10G Ethernetu.

¹Zdroje: <https://www.alternetivo.cz>, <https://www.atcmarket.cz>



Obrázek 2.8: Souhrnné schéma pro Gigabit Ethernet

2.5 Průběh přenosu

V předchozí kapitole jsme se dozvěděli, že obousměrný přenos mezi dvěma zařízeními může probíhat jedním z těchto způsobů:

- v polovičním duplexu (half duplex), kdy vysílat mohou obě zařízení, ale střídavě (v jeden okamžik může vysílat jen jedno zařízení), tedy sdílejí tentýž komunikační kanál,
- v plném duplexu (full duplex), kdy mohou vysílat třeba i obě zařízení najednou, protože pro každý směr je vyhrazen (minimálně) jeden komunikační kanál.

2.5.1 Přenos v polovičním duplexu

Režim polovičního duplexu je možné používat maximálně do rychlosti 1 Gb/s. Protože zařízení na segmentu sdílejí komunikační kanál, používá se kolizní metoda CSMA/CD, která byla popsána výše, včetně Backoff algoritmu pro řešení kolizí.



Definice (Kolizní okno)

Kolizní okno (Collision Window, Slot Time) je doba, po kterou musí zařízení naslouchat na nosné, aby zachytilo vysílání jiného zařízení a dokázalo detekovat nebezpečí kolize.



Co se kolizního okna týče, platí tyto vztahy:

- největší možná vzdálenost mezi zařízeními v téže kolizní doméně (čím větší vzdálenost, tím déle musí zařízení naslouchat, protože signálu to „déle trvá“) – čím větší je maximální povolená vzdálenost, tím větší musí být kolizní okno,
- minimální délka rámce (zařízení se musí dozvědět ještě během vysílání rámce, že došlo ke kolizi, aby mohlo vysílání zrušit, počkat dle Backoff algoritmu a vysílat znovu) – čím menší je minimální délka rámce, tím menší musí být kolizní okno a menší kolizní doména,
- pokud je malé kolizní okno, musí být malá i kolizní doména.

Z toho vyplývá, že existuje velmi úzký vztah mezi kolizním oknem, velikostí kolizní domény (max. vzdáleností mezi zařízeními v síti) a minimální délkou rámce.



Poznámka:

 Představme si tuto situaci: v kolizní doméně jsou dvě zařízení – Z_a a Z_b . Signálu trvá dobu t , než se dostane od jednoho zařízení k druhému. Zařízení Z_a začne vysílat v době t_a rámec, jehož délka je taková, že vysílání rámce trvá po dobu d_a .

Pokud zařízení Z_b začne vysílat ještě před tím, než k němu dojde signál od zařízení Z_a , dojde ke kolizi. Teď je nutné, aby se zařízení Z_a včas dozvědělo, že ke kolizi došlo, aby mohlo vysílání zrušit a po čekání vysílat znovu. Musí se to tedy dozvědět ještě před tím, než skončí vysílání, tedy nejpozději v okamžiku $t_a + d_a$. Zařízení Z_b zjistí kolizi nejdříve v čase $t_a + t$, a pokud okamžitě zareaguje a pošle jam signál, tento signál se dostane k zařízení Z_a nejdříve v době $t_a + 2 \cdot t$. Takže tu máme nerovnici:

$$\begin{aligned} t_a + d_a &> t_a + 2 \cdot t \\ d_a &> 2 \cdot t \end{aligned}$$

To znamená, že minimální doba, po kterou má trvat vysílání jednoho rámce, musí být větší než doba, kterou potřebuje signál pro cestu tam a zpět mezi dvěma nejvzdálenějšími zařízeními na segmentu (v kolizní doméně). Proto pokud povolíme příliš velkou kolizní doménu (= chceme příliš dlouhé kabely), musíme přikázat vyšší minimální délku rámce.



 Takže tady balancujeme mezi dvěma požadavky:

- Chceme sesítovat co největší prostor, tedy jsou žádoucí dlouhé kabely a velká kolizní doména.
- Nechceme zbytečně zahlcovat linku v případě, že potřebujeme posílat i malé rámce (pokud je třeba poslat množství dat menší než je limit pro minimum, musí se přidat „vycpávka“ – data bez významu), takže je efektivnější mít spíše nízký limit pro minimální délku.

K tomu se přidává ještě třetí faktor – rychlost přenosu. Pokud pouze zvýšíme rychlost přenosu, musíme zvýšit limit pro minimální délku rámce nebo zmenšit doménu, protože zvýšení rychlosti ovlivní dobu přenosu.

Zaměříme se na různé rychlosti – 10, 100, 1000 Mb/s (u rychlejších se už nepoužívá CSMA/CD), ke každé si vysvětlíme, jak se tento problém řeší.

 Pro 10 Mb/s je stanovena minimální délka rámce bez preamble na 512 bitů = 64 oktetů (podívejte se na strukturu rámce Ethernet II na straně 37), tedy odvysílání 512 bitů trvá $51,2 \mu\text{s}$ (to je doba d_a). Doba, po kterou může maximálně trvat cesta mezi dvěma nejvzdálenějšími zařízeními v kolizní doméně, je polovina této hodnoty. Vzdálenost v metrech pak záleží na tom, jaké přenosové médium je zvoleno (metalickými a optickými kabely se signál šíří odlišnou rychlostí).

Pokud je minimální délka MAC rámce bez preamble určena na 64 oktetů, pak by to, co je uvnitř zapouzdřeno, mělo být dlouhé minimálně $64 - 18 = 46$ oktetů (odečteme délku polí v záhlaví a západí MAC rámce).

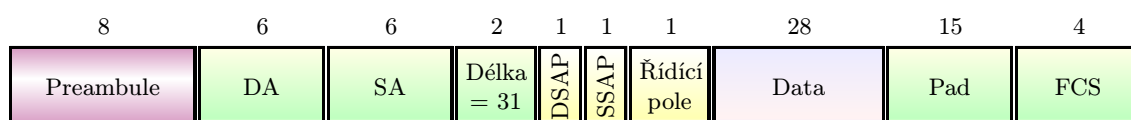
Teoreticky může nastat situace, že by uvnitř rámce mělo být přeneseno méně než 46 oktetů dat. Standard IEEE 802.3 ve spolupráci s LLC a SNAP to řeší „vycpávku“ (pad) přidanou za data, přičemž délka vycpávky se dá odvodit z hodnoty v poli *Délka*. Řešení si ukážeme na příkladu.

**Příklad**

Pokud je v poli *Délka* LLC/SNAP rámce číslo < 46 (zde konkrétně 31), pak platí:

$$\begin{aligned} \text{DSAP} + \text{SSAP} + \text{Řídicí pole} + \text{Data} + \text{Pad} &= 46 \\ \text{Délka} + \text{Pad} &= 46 \\ \text{Pad} &= 46 - \text{Délka} \end{aligned}$$

Pro případ naznačený na obrázku 2.9 na straně 48 platí $\text{Pad} = 46 - (28 + 1 + 1 + 1) = 15$. Tento údaj je důležitý, aby bylo zřejmé, na kterém oktetu rámce začíná pole s kontrolním součtem. Pozor, pole *Délka* v sobě započítává i délku LLC záhlaví (na obrázku žlutě).



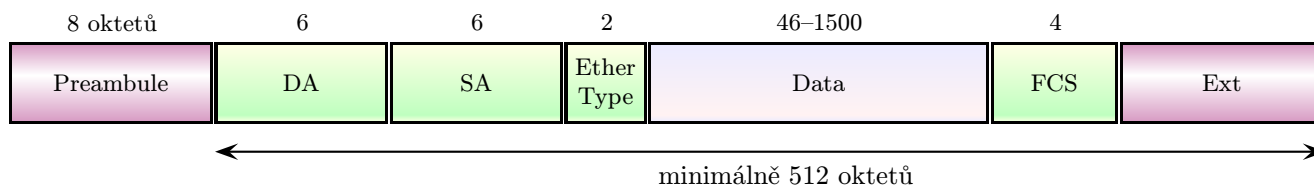
Obrázek 2.9: Velmi krátký LLC rámec v MAC rámci podle IEEE 802.3

Rámce podle Ethernet II toto neřeší, předpokládá se, že zapouzdřená data zachovávají určitou minimální délku. Protože je „uvnitř“ většinou IPv4 paket, jehož záhlaví je minimálně 20 oktetů dlouhé, na to, co je zapouzdřeno uvnitř IP paketu, zbývá minimálně 26 oktetů, což není tak moc. Připomeňme, že rámce Ethernet II se používají zejména pro přenos dat, nikoliv pro správu.

✂ Při přechodu na vyšší rychlost, tedy na 100 Mb/s, bylo rozhodnuto zachovat minimální délku rámce (tj. pro formát rámce platí totéž, co bylo napsáno v předchozím odstavci) i velikost kolizního okna, tedy bylo třeba zmenšit velikost kolizní domény mezi DCE (přibližně $10\times$).

Na metalickém kabelu to znamenalo zmenšení kolizní domény z 2000 m na 200 m, což znamená, že mezi DTE a DCE je vzdálenost max. 100 m (to zůstalo) a mezi dvěma DCE max. 200 m.

✂ Při přechodu na Gigabit Ethernet (1 Gb/s) toto řešení nebylo možné (20 m jako nejvyšší možná vzdálenost by bylo opravdu málo), tedy se pro změnu zvýšila minimální délka rámce bez započítání preamble na 512 oktetů (4096 bitů). Protože se často přenáší menší množství dat, je u rámců s délkou pod limitem nutné přidat další „vypávku“ (Carrier Extension, nedatovou příponu), a to za kontrolní součet. Umístění je naznačeno na obrázku 2.10 (pole zcela vpravo).



Obrázek 2.10: Nedatová přípona u rámce pro Gigabit Ethernet

Nedatová přípona se při odesílání přidává až po vypočtení kontrolního součtu na MAC podvrstvě, při přijetí rámce je na MAC podvrstvě cílového systému odstraněna ještě před kontrolou podle kontrolního součtu. Skladba symbolů v nedatové příponě je patentovaná – musí být odlišitelná od pole s kontrolním součtem.

**Další informace:**

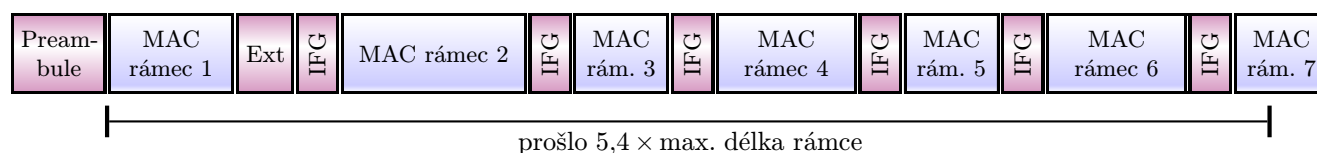
Údaje o patentu jsou dostupné na <http://www.google.com/patents/US5940401>.



Pokud by se posílalo velmi často větší množství rámců s malým množstvím zapouzdřených dat, linka by samozřejmě byla zbytečně vytěžována víc než je nutné. Takový případ se řeší jinak:

 **Burst mode (shlukový režim)** je právě určen pro posílání sekvencí krátkých úseků dat. Je používán pouze u přenosů v rychlostech od 1 Gb/s výše, tj. od Gigabit Ethernetu. „Krátké“ rámce se posílají ve shluku, který má následující strukturu:

- první rámec ve shluku má výše popsanou strukturu, včetně případné nedatové přípony, pokud je menší než 512 oktetů,
- následuje *IFG* (InterFrame Gap) – speciální sekvence bitů vyplňující prostor mezi rámci, slouží k rozpoznatelnému oddělení rámců ve shluku,
- další rámce v posloupnosti navzájem oddělené IFG sekvencemi; pokud jsou kratší než 512 oktetů, není nutné je doplňovat o nedatovou příponu.



Obrázek 2.11: Přenos v burst (shlukovém) módu

Celková délka shluku by neměla překročit přibližně 5,4násobek maximální možné délky rámce (přesněji 8 192 oktetů), ale pokud je tato hranice překročena až během odesílání posledního rámce ve shluku, může být jeho odesílání ještě dokončeno a shluk je ukončen až poté.

**Poznámka:**

To znamená, že během přenosu v burst modu je třeba sledovat, zda se neblížíme k hranici 5,4násobku maximální délky rámce. Pokud jí dosáhneme, pak rámec, který právě odesíláme, je posledním rámcem shluku. Je to důležité, protože během odesílání shluku rámců blokuje linku a znemožňujeme odesílat ostatním zařízením v segmentu (téže kolizní doméně), a blokování nemůže trvat příliš dlouho.



2.5.2 Přenos v plném duplexu

V plném duplexu je situace mnohem jednodušší. Na logické úrovni se v podstatě setkáváme jen se spoji typu point-to-point a veškeré spoje jsou vlastně vyhrazené, zařízení může vysílat kdykoliv. Nepoužívá se přístupová metoda CSMA/CD a parametry spojení nejsou omezeny vlastnostmi kolizní domény a kolizního okna, pouze fyzikálními vlastnostmi přenosové cesty (například útlumem).

Vysílání může prakticky pořád probíhat v burst modu tak, jak bylo popsáno u polovičního duplexu. Vysílání a příjem probíhají po jiné cestě, ale může jít o jeden společný kabel (například v mnoha standardech pro fyzickou vrstvu je pro UTP se čtyřmi páry kroucené dvojlinky vždy jedna dvojice vyhrazena pro

příjem a druhá dvojice pro vysílání, nebo se směr určuje adaptivně). U optických kabelů je buď pro každý směr jeden kabel, nebo pravděpodobněji v kabelu máme jedno či více vláken zvlášť pro různé směry.

 Jediný problém (který může nastat i u polovičního duplexu, ale pro plný duplex je typický) nastává tehdy, když vysílající zařízení neustále vyvíjí aktivitu, ale přijímající zařízení nestíhá rámce přijmout. Přijetí rámce totiž není až tak triviální, kromě zpracování na fyzické vrstvě je třeba také postupně rámec rozbalit a patřičně ho zpracovat, a procesor (který se na tom taky podílí) může mít i jinou práci než zpracovávat rámce od dotyčného zařízení. Tento problém se řeší kombinací těchto způsobů:

- zařízení má dostatečně velký *buffer*, tam si ukládá to, co zatím nemohlo být zpracováno,
- odesílatele je třeba zpomalit, sdělit mu, že má určitou dobu počkat s vysíláním.

První možnost je celkem jasná, ale nemusí být dostačující. Druhá možnost znamená použití tzv. *pause frame*, tedy speciálního rámce, který je poslán „pilnému“ odesílateli; tento rámec obsahuje především informaci o tom, na jak dlouho je třeba přerušit vysílání.



Poznámka:

Plný duplex se dá použít pouze na spojích typu point-to-point (v kolizní doméně jsou tedy pouze právě dvě fyzicky propojená zařízení), což například naprosto vylučuje použití hubu (rozbočovače). Navíc obě propojená zařízení „musí umět“ plný duplex, aby bylo možné ho využít.

Od 10G Ethernetu (ale také 2.5G a 5G) se používá pouze plný duplex, poloviční není podporován. Ke kolizím z principu nemůže dojít, tedy nemá smysl využívat jakoukoliv přístupovou metodu (CSMA/CD se v tom případě nepoužívá).



2.6 Technologie související s Ethernetem

2.6.1 Autonegociace

Aby vůbec bylo možné zprovoznit přenos přes ethernetovou síť, je třeba, aby si síťová rozhraní komunikujících zařízení „rozuměla“, tedy aby se shodla na následujících parametrech:

- rychlost přenosu,
- pro danou rychlost konkrétní standard, pokud je víc možností pro použité přenosové médium (například aby bylo jasné, jak konkrétně budou jednotlivé páry UTP využívány),
- poloviční nebo plný duplex.

Zároveň je potřeba pravidelně provádět kontrolu funkčnosti spojení.

 Pokud propojíme dvě zařízení kabelem a zprovozníme příslušná síťová rozhraní, pak se tato síťová rozhraní nejdříve dohodnou na výše uvedených parametrech – provede se *autonegociace* (tedy automatická dohoda o parametrech) a následně se pravidelně ozývají v době, kdy se nepřenášejí žádné rámce (aby bylo jasné, že dané zařízení je pořád připojeno a linka funguje). Autonegociace se provádí na fyzické vrstvě, je popsána přímo v standardu IEEE 802.3 a dodatecích.

První pokusy o autonegociaci (ale ještě ne zcela v plném smyslu toho slova) byly u Ethernetu 10Base-T, kdy síťové rozhraní vysílalo v pravidelných intervalech impulsy, které sloužily ke zjišťování, zda je linka funkční. Tento signál se označoval *NLP* (Normal Link Pulse). Pokud od zařízení po určitou dobu nepřišel žádný rámec ani NLP signál, linka byla považována za nefunkční.

Dnes, v době vyšších rychlostí, se odesílá signál *FLP* (Fast Link Pulse), kde se série původních NLP signálů prokládá dalšími informacemi (protože se postupně přidávaly další parametry, které bylo možné oznámit), a zatímco původně šlo o nepovinnou funkčnost, od Gigabit Ethernetu je pro metalické kabely autonegociace povinná.

Celý mechanismus je zpětně kompatibilní, a v reálu se pro linku použijí nejvyšší možné parametry z těch, které splňují obě komunikující strany. Například pokud jedno zařízení zvládá rychlost 100 Mb/s v polovičním duplexu a druhé rychlost 1 Gb/s s plným duplexem, použije se rychlost 100 Mb/s a poloviční duplex.

2.6.2 Napájení přes Ethernet

 *PoE* (Power over Ethernet) je standard popisující možnost napájení menších zařízení přes síťový (metalický) kabel. První PoE standard IEEE 802.3af je z roku 2003. Brzy přišla nová verze IEEE 802.3at, v současné době se setkáváme s další verzí IEEE 802.3bt.

Výhodou je, že k dotyčnému zařízení nemusíme přivádět napájecí kabel. Takto napájené zařízení si samozřejmě musí vystačit s nižším odběrem, typicky jde například o IP telefony, IP kamery a některé Wi-fi access pointy připojené kroucenou dvojlinkou k zařízení, které může fungovat jako zdroj PoE (což je většinou switch).

 PoE rozlišuje obě strany: jedna strana je cílem energie (PD, Powered Device) a druhá strana poskytuje napájení (PSE, Power Sourcing Equipment), což je tedy obvykle switch.

Switch sloužící jako zdroj PoE musí mít samozřejmě dostatečně dimenzovaný napájecí zdroj (což je ta komponenta, která stejně jako u běžných počítačů provádí transformaci střídavého napětí z napájecí sítě na stejnosměrné napětí), a i přesto obvykle nedokáže napájet zařízení na všech svých portech, zbývající připojená zařízení musí mít vlastní napájení.

Podle starších standardů se v kabelu typu kroucená dvojlinka (se 4 páry vodičů) dva páry používaly pro data a zbývající dva páry pro napájení přes PoE. Ale jak víme, od rychlosti 1 Gb/s jsou pro data používány všechny čtyři páry. Proto podle nového standardu IEEE 802.3bt se tyto čtyři páry v případě PoE sdílejí, tedy přenášejí se po nich data i energie.

 Rozlišujeme *výkonové třídy* 0–4 (starší standardy), resp. 0–8 (nový standard) podle toho, jaký proud a maximální příkon napájené zařízení potřebuje, přičemž třída 0 je pro zařízení nepodporující PoE. Čím vyšší příkon, tím vyšší třída.

Podobně jako při stanovování parametrů připojení se v rámci autonegociace dojednává rychlost a duplex, u PoE se při připojení určuje třída. Nejdřív proběhne detekce, při níž se k připojenému zařízení pouští proud o nízkém napětí, které by v případě, že zařízení „neumí“ PoE, nemělo ublížit (odpovídá třídě 0). Jestliže zařízení zareaguje, pokračuje detekce určením výkonové třídy.



Další informace:

- <https://vyvoj.hw.cz/produkty/ethernet/princip-cinnosti-power-over-ethernet.html>
- https://www.arubanetworks.com/techdocs/AOS-CX/10.11/HTML/monitoring_6200/Content/Chp_PoE/pow-ove-eth-oveview.htm





Souhrn kapitoly

V této kapitole jsme se zaměřili na lokální síť Ethernet. Zabývali jsme se průběhem komunikace na vrstvách L2 a L1, přístupovou metodou včetně fungování Backoff algoritmu, formátem rámce, způsobem výpočtu kontrolního součtu rámce, průběhem přenosu v polovičním a plném duplexu, a na konci jsme se zaměřili na způsob vyjednávání spojení a fungování napájení přes Ethernet (PoE).



Další témata k lokálním sítím

 **Rychlý náhled:** V této kapitole se budeme zabývat technologiemi, které souvisejí s lokálními sítěmi: nejdřív si ujasníme, které MAC a IP adresy se používají na jednotlivých úsecích cesty rámce/paketu, pak po přehledu standardů rodiny IEEE 802 se zaměříme na srovnání různých přístupových metod v lokálních sítích, následně se budeme zabývat virtuálními lokálními sítěmi (VLAN), dále protokolem STP pro odstraňování smyček v síti switchů, podíváme se na technologii slučování linek EtherChannel a na konci kapitoly si projdeme nejznámější SAN síť.

 **Klíčová slova:** Adresa, IEEE 802, přístupová metoda, CSMA/CD, token, VLAN, přístupový port, trunk, STP, spanning-tree, konvergence, bridge ID, BPDU, EtherChannel, LACP, PaGP, SAN.

 **Cíle studia:** Cílem v této kapitole je zopakovat si vztah mezi MAC a IP adresou na různých úsecích cesty, zorientovat se v rodině protokolů IEEE 802, umět porovnat různé přístupové metody, pochopit princip virtuálních lokálních sítí a metody jejich konfigurace, ujasnit si roli STP v síti switchů s redundantními cestami, ujasnit si význam a princip fungování slučování linek pomocí technologie EtherChannel a získat přehled v používaných SAN sítích v datových centrech.

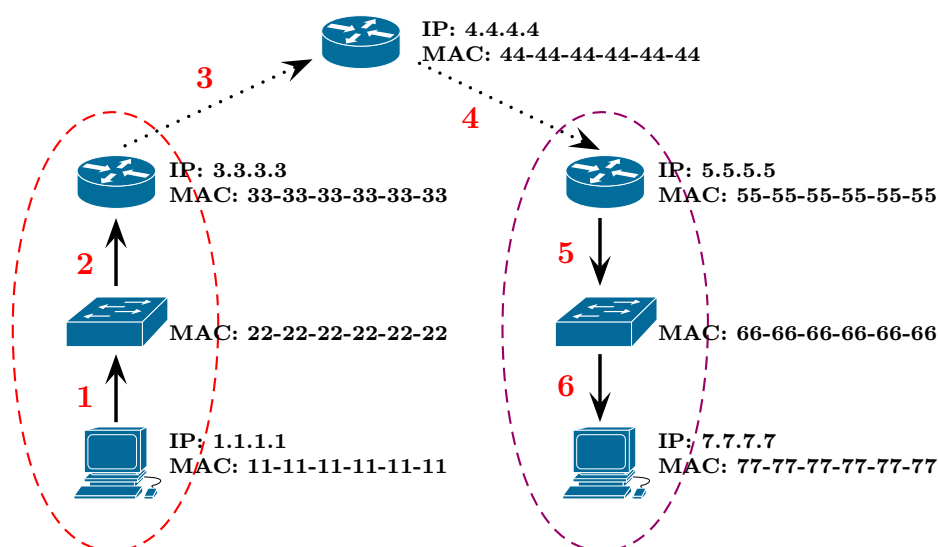
3.1 Adresy v rámcích a paketech

Když odesíláme IP paket, potřebujeme samozřejmě IP adresu zdroje (naši) a příjemce, ale musíme si uvědomit, že IP paket se má zapouzdřit do rámce, a tam mají být MAC adresy zdroje a cíle.

Do rámce, který odesíláme, jako zdrojovou dáme vlastní MAC adresu, ale jakou dát tu cílovou? Určitě nemůžeme použít MAC adresu příjemce, protože tu neznáme. O příjemci (třeba serveru) máme totiž jen omezenou informaci (jenom IP nebo jmennou adresu), navíc pod toutéž jmennou či IP adresou může v reálu existovat více než jedno fyzické zařízení, například z důvodu vyrovnávání zátěže nebo kvůli údržbě. Jediná MAC adresa, kterou známe, je MAC adresa brány.

Příklad

Na obrázku 3.1 je naznačena struktura sítě, přes kterou máme poslat paket. Zdrojovým zařízením je počítač s IP adresou 1.1.1.1, cílovým zařízením počítač s IP adresou 7.7.7.7.



Obrázek 3.1: Vztah mezi L2 a L3 adresováním

Jak to proběhne:

- Na prvním úseku cesty směřuje PDU k switchi, který vidí jen do rámce; rámec a paket zapouzdřený v rámci obsahují adresy:

Zdroj:	IP 1.1.1.1	MAC 11-11-11-11-11-11	adresy odesílatele
Cíl:	IP 7.7.7.7	MAC 33-33-33-33-33-33	IP adresa příjemce, MAC adresa brány

Všimněte si, že IP a MAC adresa cíle k sobě nepatří, ale to vůbec nevadí.

- Na druhém úseku se PDU dostává k bráně a adresy se zatím nezměnily:

Zdroj:	IP 1.1.1.1	MAC 11-11-11-11-11-11	adresy odesílatele
Cíl:	IP 7.7.7.7	MAC 33-33-33-33-33-33	IP adresa příjemce, MAC adresa brány

- Třetí úsek cesty již vede mimo síť odesílatele. Router, který plní roli brány, přijal a rozbalil rámec, rozbalil vnořený IP paket a znovu ho zabalil, tentokrát s jinými MAC adresami.

Zdroj:	IP 1.1.1.1	MAC 33-33-33-33-33-33	IP adresa odesílatele, MAC adresa začátku úseku
Cíl:	IP 7.7.7.7	MAC 44-44-44-44-44-44	IP adresa příjemce, MAC adresa konce úseku

- Podobně to bude na čtvrtém úseku:

Zdroj:	IP 1.1.1.1	MAC 44-44-44-44-44-44	IP adresa odesílatele, MAC adresa začátku úseku
Cíl:	IP 7.7.7.7	MAC 55-55-55-55-55-55	IP adresa příjemce, MAC adresa konce úseku

- Pak již jsme v síti příjemce. Brána v síti příjemce „zná“ MAC adresy uzlů ve své síti.

Zdroj:	IP 1.1.1.1	MAC 55-55-55-55-55-55	IP adresa odesílatele, MAC adresa cizí brány
Cíl:	IP 7.7.7.7	MAC 77-77-77-77-77-77	adresy příjemce

- Na šestém úseku už jen zůstáváme na vrstvě L2 a nemá smysl adresy měnit.

Zdroj:	IP 1.1.1.1	MAC 55-55-55-55-55-55	IP adresa odesílatele, MAC adresa cizí brány
Cíl:	IP 7.7.7.7	MAC 77-77-77-77-77-77	adresy příjemce

Tento příklad je trochu zjednodušený, v reálu bude mít každé síťové zařízení na každém rozhraní jinou MAC adresu.

Z toho plyne, že opravdu nemusíme znát MAC adresu cíle, stačí nám IP adresa, a cíl nepotřebuje naši MAC adresu. Cíl dostane rámec obsahující MAC adresu odesílatele pouze v případě, že oba jsou v téže síti. Taky z toho plyne další informace – switch, ke kterému jsme přímo připojeni, je pro nás vlastně „neviditelný“, k jeho MAC adrese se v rámci nedostaneme (k jeho IP adrese taky ne).

3.2 Přehled rodiny standardů IEEE 802

Víme, že Ethernet je standardizován jako IEEE 802.3, a že existují standardy IEEE 802.2 (pro podvrstvu LLC linkové vrstvy) a standard IEEE 802.1 (s jeho některými částmi jsme se setkali v této kapitole). Ve skutečnosti je rodina standardů IEEE 802 (Standardy pro místní sítě) mnohem širší. Obsahuje standardy definující vrstvu síťového rozhraní podle TCP/IP, resp. vrstvy L1 a L2 podle ISO/OSI, nebo část této vrstvy (vrstev). Obvykle se počítá s „podsunutím“ pod síťový zásobník TCP/IP.

Standard	Význam
802.1	Technologie mostů a správa sítí na vrstvě L2
802.2	Logické řízení spoje na vrstvě L2 – podvrstva LLC vrstvy L2 (pracovní skupina neaktivní)
802.3	Lokální síť Ethernet, přístupová metoda CSMA/CD (vrstva L1 a podvrstva MAC vrstvy L2)
802.4	Lokální síť Token Bus – sběrnice LAN (pracovní skupina rozpuštěna)
802.5	Lokální síť Token Ring – kruhová LAN (pracovní skupina neaktivní)
802.6	Metropolitní síť DQDB (pracovní skupina rozpuštěna)
802.7	Širokopásmové (broadband) sítě LAN (pracovní skupina rozpuštěna)
802.8	Sítě LAN/MAN na optických vláknech (pracovní skupina rozpuštěna)
802.9	Integrované služby – izochronní síť (pracovní skupina rozpuštěna)
802.10	Bezpečnost v sítích LAN/MAN (pracovní skupina rozpuštěna)
802.11	WLAN Working Group: Bezdrátová lokální síť – Wi-fi
802.12	Vysokorychlostní síť – 100VG-AnyLAN (pracovní skupina rozpuštěna)
802.14	Širokopásmová síť na kabelech pro kabelové televize (pracovní skupina rozpuštěna)
802.15	WPAN Working Group: bezdrátové osobní sítě (například Bluetooth, ZigBee, UWB), včetně koexistence s IEEE 802.11
802.16	Širokopásmové bezdrátové MAN sítě (WiMAX) – neaktivní
802.17	RPR – Resilient Packet Ring, používá se v sítích SONET/SDH, neaktivní
802.18	Regulace rádiových vln
802.19	Koexistence nelicencovaných bezdrátových sítí
802.20	Mobilní širokopásmové bezdrátové LAN/MAN sítě včetně dopravní mobility
802.21	Handover Services Working Group – předávání mezi různými typy mobilních a bezdrátových sítí (GSM, GPRS, Wi-fi, Bluetooth, WiMAX, atd.)
802.22	Bezdrátové regionální sítě využívající nepoužívané frekvence televizního vysílání
802.23	Emergency Services Working Group (neaktivní)
802.24	Smart Grid Technical Advisory Group: vertikální aplikace – Internet of Things (IoT), komunikace Machine-to-Machine, Smart Grid, propojení dopravních prostředků...

Tabulka 3.1: Rodina standardů IEEE 802

V tabulce 3.1 je přehled momentálně existujících standardů z rodiny IEEE 802. Každý standard je vypracováván konkrétní pracovní skupinou (working group), například standard IEEE 802.3 má na starosti *Ethernet Working Group* a standard IEEE 802.11 má „na svědomí“ *Wireless LAN Working Group*.

Nejdůležitější standardy (tedy z našeho pohledu) jsou zvýrazněny tučným písmem. Některé z nich už známe, s jinými se seznámíme později.



Poznámka:

Všimněte si, že chybí řádek IEEE 802.13. Oficiálně je číslo 13 vyhrazeno pro další vývoj fast Ethernetu, ve skutečnosti je důvod podobný důvodu toho, proč v mnoha hotelech chybí třinácté patro.



Standard IEEE 802.1 je asi nejrozmanitější. Zatím jsme se setkali s těmito jeho podřízenými standardy (velké písmeno v označení) a dodatky (malé písmeno):

- IEEE 802.1Q – VLAN (Virtuální lokální síť),
- IEEE 802.1p – stanovení priority (hodnota využívající 3 bity, tedy v rozsahu 0–7) používané například standardem IEEE 802.1Q, ve skutečnosti to není standard, pouze jednoduchý předpis,
- IEEE 802.1D – protokol STP (technologie mostů a síť bez smyček),
- IEEE 802.1w – původně protokol RSTP (opět není chápáno jako samostatný standard), později je specifikace přidána do revize IEEE 802.1D-2004,
- IEEE 802.1s – původně protokol MSTP, později přidán do revize IEEE 802.1Q-2005.

Z těch, kterým jsme se dosud nevěnovali, je zajímavý například IEEE 802.1X – autentizace zařízení přistupujícího do LAN. Autentizace podle tohoto standardu se často používá ve firemních sítích pro ověřování přístupu uživatele/zařízení do sítě a určování konkrétních oprávnění pro daného uživatele.



Poznámka:

Mnohé z pracovních skupin IEEE 802 jsou buď neaktivní nebo dokonce rozpuštěné. Neaktivní znamená, že dlouho nebyla vydána žádná aktualizace či dodatek (což není problém, když dodatky nejsou nutné), rozpuštěná znamená, že se ani v budoucnu neplánují žádné aktivity (typicky u technologií, které se neprosadily). Například síť Token Ring, Token Bus a 100VG-AnyLAN byly postupně vytlačeny Ethernetem, Izochronní síť nabízející možnost přidělovat prioritu některým typům dat (třeba při přenosu hlasu) také (Ethernet s využitím IEEE 802.1p/Q to umí taky a je rychlejší). Jinými slovy, právě Ethernet stojí za upozaděním většiny standardů z rodiny IEEE 802.



3.3 Přístupové metody



Pojem *přístupová metoda* už známe – víme, že se jedná o metodu pro stanovení vysílacího přístupu připojeného zařízení na sdílené přenosové médium, přičemž se může jednat i o kolizní metodu (určující buď způsob vyhýbání se kolizím nebo reakci na kolizi). Zatím jsme se seznámili s přístupovou/kolizní metodou CSMA/CD, která se používá v Ethernetu provozovaném v polovičním duplexu a jejím účelem je, aby zařízení dokázalo zjistit kolizi a správně na ni reagovat. V sítích IEEE 802.11 (Wi-fi) se pro změnu používá přístupová metoda CSMA/CA, kde se klade důraz na to, aby ke kolizi pokud možno vůbec nedošlo.

V této podkapitole se podíváme na další lokální sítě a jejich přístupové metody. Vzhledem k tomu, že tyto sítě byly prakticky vytlačeny Ethernetem, nebudeme zacházet do podrobností.

 Přístupová metoda je *deterministická*, pokud zaručuje každému připojenému zařízení možnost vysílání, tedy že (při určitém čekání) dostane příležitost vyslat rámeček, přístup všech zařízení k přenosovému médiumu je spravedlivý a rovnocenný. Přístupová metoda, která toto nezaručuje, je *nedeterministická*.

Přístupové metody CSMA/CD a CSMA/CA jsou nedeterministické. Není stanoveno žádné pořadí, ve kterém zařízení mohou vysílat, a při větším zarušení sítě se může stát, že rámeček nebude možné vyslat vůbec (vzpomeňte si na Backoff algoritmus – po určitém počtu čekacích cyklů to odesílatel vzdá).

3.3.1 Token Ring

Token Ring je síť původně od společnosti IBM, později byla standardizována jako IEEE 802.5.

 *Fyzická topologie* je hvězda (nebo více propojených hvězd), přičemž ve středu hvězdy máme speciální zařízení nazývané MSAU (Multistation Attachment Unit) – obdobu switchu z Ethernetu. *Logická topologie* je kruhová (proto je v názvu „Ring“), tedy připojená zařízení typu DTE jsou seřazena do logického kruhu, což souvisí právě s přístupovou metodou.

 Používá se *přístupová metoda Token Passing* (předávání tokenu = peška, jako v dětské hře „chodí pešek okolo“). Síť prochází v logickém kruhu speciální rámeček zvaný *token*, vysílat může pouze to zařízení, které tento token právě „vlastní“. Když chce zařízení vyslat rámeček, postupuje takto:

- počká, až obdrží token, čímž získá právo vysílat,
- vyšle rámeček s daty,
- počká na potvrzení o doručení (tedy potvrzující rámeček dojde po kruhu zpět k odesílateli),
- vyšle do sítě původní token, čímž dovolí vysílat někomu dalšímu.

Pokud zařízení obdrží token a nechce vysílat, pak token přepoše sousedovi na kruhu. Naopak když chce vysílat, „pozdrzí“ token u sebe po celou dobu od vyslání rámce do přijetí potvrzení.

Přístupová metoda Token Passing je deterministická – neobsahuje žádný náhodný prvek a zaručuje každému zařízení, že dostane možnost vyslat rámeček.

Protokol IEEE 802.5 definuje pouze vrstvu L1 a podvrstvu MAC vrstvy L2 (podobně jako IEEE 802.3). Záhlaví a zápatí rámce je složitější než u Ethernetu, má o několik polí více. MAC rámeček podle IEEE 802.5 obsahuje v záhlaví bit indikace tokenu. Pokud je tento bit nastaven na 1, jde o krátký tokenový rámeček a uvnitř není nic zapouzdřeno, kdežto pokud je tento bit nastaven na 0, máme uvnitř LLC nebo SNAP rámeček s daty.

	Ethernet	Token Ring
Přístupová metoda	nedeterministická CSMA/CD	deterministická Token Passing
Kolize	jen při half duplexu	žádné
Fyzická topologie	hvězda/strom	hvězda/strom
Logická topologie	sběrnice/point-to-point	kruh
Složitost	jednoduché rámce	složitější rámce
Cena	relativně levný hardware	mnohem dražší hardware

Tabulka 3.2: Srovnání sítí Ethernet a Token Ring

Z údajů v tabulce 3.2 plyne, proč nad sítěmi Token Ring zvítězil Ethernet – v prvních dvou řádcích je Ethernet v nevýhodě, ale nedeterminismus přístupové metody je obcházen přechodem na plný duplex. V praxi je jedním z nejdůležitějších kritérií cena, která také má svůj podíl na momentálním stavu.

3.3.2 100VG-AnyLAN

Síť 100VG-AnyLAN (standard IEEE 802.12) vznikla jako alternativa během vývoje Fast Ethernetu. Název 100VG AnyLAN znamená:

- rychlost 100 Mb/s (jako Fast Ethernet),
- VG je zkratka z „Voice Grade“ – hlasová kvalita (protože se počítalo s kabeláží UTP Cat3 používané také pro telefonní linky),
- AnyLAN je odkaz na to, že jde o LAN s tím, že se pro data používají všechny v té době nejpoužívanější typy rámců (Ethernet II, IEEE 802.3 a IEEE 802.5).

Z toho je zřejmé, že v IEEE 802.12 nejsou přímo specifikovány rámce pro data, ale používají se rámce z jiných typů sítí.

 *Fyzická topologie* je hvězda nebo strom, jako DCE jsou použity huby (rozbočovače), případně mosty. Síť je přísně hierarchická – jeden hub je kořenem stromu, každý uzel má přehled o svých podřízených uzlech z nižší vrstvy a o svém nadřízeném DCE. *Logická topologie* je také hvězda či strom.

 Používá se přístupová metoda *Demand Priority* (DPAM, Demand Priority Access Method), která je centralizovaná a vysílat může pouze takové zařízení, které k tomu dostane povolení. Pokud některé zařízení v síti chce vysílat, musí poslat nadřízenému uzlu speciální rámec s žádostí o vysílání s určením priority vysílaných dat (vysoká priorita pro multimediální data). Až po obdržení povolení vysílá rámec.

Kořenový DCE si vede *prioritní frontu žádostí o vysílání* (to znamená, že žádosti o vysílání dat s vyšší prioritou předbíhají nižší prioritu). Pokud žádost s nižší prioritou čeká na vyřízení déle než 300 ms, dostane vyšší prioritu. DCE, který není kořenový, pouze přeposílá obdržené žádosti směrem nahoru a obdržená povolení směrem dolů. Jak vidíme, rozhodování je plně centralizované.

Přístupová metoda DPAM je *deterministická* – řízení pomocí front neobsahuje žádnou náhodnost a zaručuje, že každé zařízení dříve či později dostane povolení vysílat.

	Ethernet	100VG-AnyLAN
Přístupová metoda	nedeterministická CSMA/CD	deterministická Demand Priority
Kolize	jen při half duplexu	žádné
Fyzická topologie	hvězda/strom	hvězda/strom
Logická topologie	sběrnice/point-to-point	strom, centralizovaná
Max. propustnost	50 % při použití hubů	i přes 90 %

Tabulka 3.3: Srovnání sítí Ethernet a 100VG-AnyLAN



Poznámka:

Principy použité při návrhu 100VG-AnyLAN byly původně zamýšleny pro Fast Ethernet – na vývoji Fast Ethernetu původně pracovala skupina lidí, která přišla právě s touto myšlenkou. Vedením však byli odmítnuti („to přece není Ethernet – vytvořte si vlastní pracovní skupinu“), a tak vznikla samostatná pracovní skupina pro IEEE 802.12 a nový typ sítě.



Proč zvítězil Ethernet? Princip 100VG-AnyLAN byl původně životaschopný a mohl Fast Ethernetu konkurovat (deterministická přístupová metoda, lepší chování při vysokém vytížení sítě). Jenže došlo ke zlevnění switchů, a při použití switchu jako DCE se dá implementovat plný duplex při mnohem lepším využití spoje a nevýhody Ethernetu odpadají. Navíc může být na závadu centralizovanost řízení – decentralizovaný systém lépe odolává výpadkům.

3.4 VLAN

VLAN (Virtuální LAN, Virtual LAN) je logické seskupení určitých zdrojů (například počítačů) nacházejících se v jedné nebo několika běžných lokálních sítích. Členění na virtuální síť se sice konfiguruje na vrstvě L2, ale „táhne se“ na vrstvu L3, protože každá VLAN odpovídá některé podsíti naší sítě. Zařízení patřící do určité VLAN má přiřazenu IP adresu spadající do příslušné podsítě.

Kdyby nešlo o VLANy, ale o „obyčejné“ podsítě, měli bychom router (nebo L3 switch) a za každým jeho L3 portem by byla jiná podsít', tedy všechna zařízení nacházející se za tím jedním portem by musela patřit do stejné podsítě. VLANy nám dávají možnost oprostit příslušnost k podsíti od fyzického umístění (od toho, ke kterému switchi/portu routeru je zařízení připojeno: k témuž switchi mohou být připojena zařízení patřící do různých VLAN.

V rámci téže VLAN (a v důsledku podsítě) je možné komunikovat jednoduše bez jakéhokoliv filtrování a přeposílání, přes switchu. Pokud mají komunikovat zařízení patřící do různých VLAN, pak tyto pakety musí jít přes L3 zařízení, které má jedno rozhraní v jedné VLAN/podsíti a druhé rozhraní v druhé VLAN/podsíti. L3 zařízení pak může například filtrovat, určovat, jestli ta komunikace má či nemá proběhnout.

 Proč něco takového dělat? Například z těchto důvodů:

- Zabezpečení – máme možnost omezit přístup ke konkrétnímu prostředku (třeba serveru s citlivými firemními informacemi) jen na někoho (členy jedné konkrétní VLANy).
- Omezení komunikace mezi VLAN sítěmi platí i pro broadcastové vysílání.
- Máme možnost zprehlednit si síť rozdělením do částí bez toho, abychom toto rozdělení museli řešit na úrovni kabelů a umístění switchů.

Rozdělení VLAN má smysl nejen kvůli bezpečnosti, ale například taky pro nastavení kvality služby (QoS), tedy priorit, zároveň s dalšími možnostmi řízení provozu. Například je typicky používána voice VLAN (pro IP telefony v rámci firmy).

 Členství ve skupině (VLAN) se dá stanovit podle různých kritérií, záleží na tom, na které vrstvě ISO/OSI chceme VLAN implementovat, tedy na konkrétní technologii.

1. *VLAN podle portů* – na přepínačích v síti jsou konkrétní porty (a tedy stanice k nim připojené) roztrženy do jednotlivých VLAN. Tato metoda je poměrně jednoduchá. Při změně struktury sítě (například stanice bude přenesena a připojena k jinému portu) je však nutné provést rekonfiguraci, a další nevýhodou je nemožnost zařazení stanice (portu) do více než jedné VLAN.
2. *VLAN podle MAC adres* (na L2) – tato metoda je náročnější na počáteční konfiguraci, protože MAC adresy musejí být do skupin přiřazeny ručně. Odpadá však ruční rekonfigurace při změně umístění stanice a jedna stanice může být i ve více VLAN. Nevýhodou je větší časová náročnost přepínání a s tím související snížená propustnost sítě (zvláště pokud jedna MAC adresa přísluší do více VLAN). Je třeba dát pozor na možnost pozměnění MAC adresy.

3. *VLAN na síťové vrstvě* – seskupujeme stanice podle IP adresy podsítě.
4. *VLAN podle politiky* – příslušnost do VLAN je určena kombinací více kritérií (tj. politikou, „zásadou“) a může být dynamicky měněna podle chování stanice v síti (umístění, typu a množství zasílaných rámců, atd.). Chování stanic a provoz na síti jsou automaticky monitorovány a politiky mohou být přizpůsobovány.
5. *VLAN s autentizací* – můžeme použít IEEE 802.1X pro autentizaci uživatele a autorizace pak obsahuje zařazení uživatele do příslušné VLAN. Je také možné nastavit, že uživatel, který nebyl autentizován (host) bude zařazen do speciální VLAN pro hosty.

V současné době se často kombinuje VLAN podle portů na L2 pro komunikaci v rámci téže VLANy a VLAN na síťové vrstvě pro zajištění (řízené) komunikace mezi VLANy, přičemž můžeme mít v síti IEEE 802.1X server řadící stanice do VLAN v rámci autorizace, případně Active Directory či jinou implementaci protokolu LDAP s politikami.

V praxi se pro VLAN výhradně používá implementace protokolu IEEE 802.1Q (určuje, jak se má informace o příslušnosti zdroje dat k určité VLAN distribuovat mezi switchi). V dalším textu budeme popisovat základ daný zmíněným protokolem, ten platí vždy, ať použijeme jakékoliv další technologie. Každé koncové zařízení by mělo do některé VLAN patřit, což se obvykle určuje nastavením příslušné VLAN na portu switche, přes který je zařízení přímo dostupné. VLANy se rozlišují především svým číslem, ale každá VLAN může mít i slovní označení, aby se nich lidé lépe orientovali.

 Používáme tyto termíny:

- *datová VLAN* – VLAN pro běžný provoz,
- *default* – pokud na portu není nastaveno nic jiného, je tam defaultní VLAN (u většiny výrobců je to VLAN číslo 1),
- *native* – pokud se mezi dvěma switchi přenáší něco, co nemá žádnou VLAN nastavenou, spadne to do „kanálu“ nativní VLAN (přes nativní VLAN bývají přenášeny například také STP rámce); jestliže jsme v konfiguraci žádnou nativní nenastavili ručně, VLAN 1 se použije jako nativní,
- *management* – je určená pro management rámce (SSH, Syslog, SNMP, ...), výchozí nastavení je VLAN 1, ovšem rozhodně je dobré ji přenastavit.

 Konkrétní port na switchi je

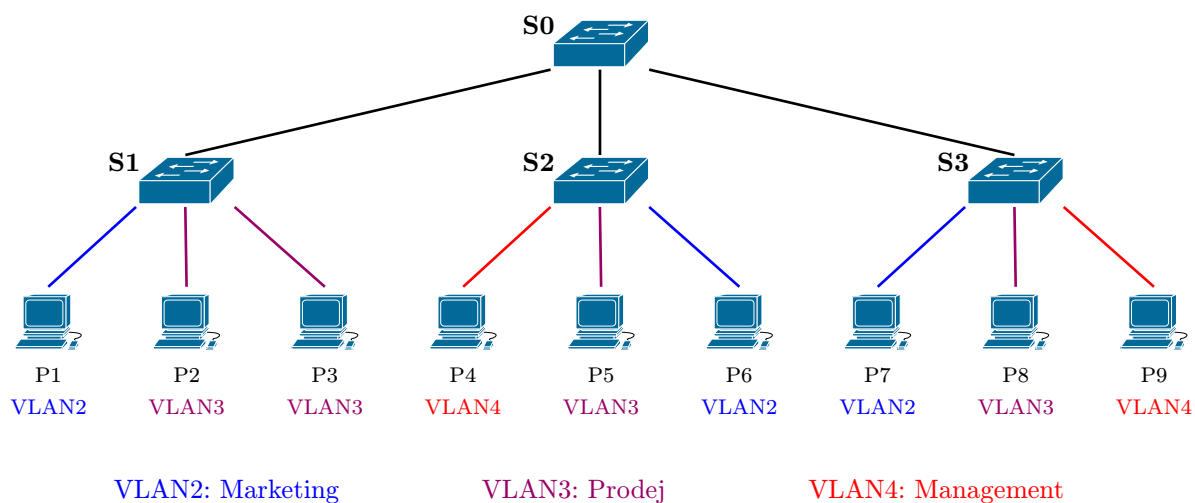
- *přístupový* – za ním je přímo jedno zařízení, má nastavenou příslušnost do jediné VLAN,
- *trunkový* – za ním je další switch (tedy potenciálně celá řada různých zařízení patřících do různých VLAN), sousední switch má příslušný port taky nastavený jako trunkový.

Přes trunkový port lze přenášet rámce patřící do více různých VLAN (nastavuje se seznam „povolených“ VLAN, seznam musí být stejný na obou koncích trunku).

3.4.1 VLAN na jednom switchi

Princip VLAN je jednoduchý – rozdělíme zařízení v síti do skupin, každé skupině přiřadíme identifikátor (číslo VLAN) a zajistíme oddělení komunikace mezi těmito skupinami. Na obrázku 3.2 je naznačeno, jak by to mohlo vypadat (poněkud zjednodušeně).

Na switchi je pro každý port určeno, do které VLAN patří zařízení přes tento port přístupné. Tedy v našem případě je na switchi S1 první port přiřazen do VLAN2 a další dva porty do VLAN3. Port vedoucí směrem nahoru ke switchi S0 bude přenášet rámce z různých VLAN, proto jde o trunk.



Obrázek 3.2: Ukázka využití VLAN

Zatímco spoje ke koncovým zařízením jsou znázorněny různými barvami podle příslušnosti k VLAN, trunk je černě, ale můžeme si ho představit jako celý svazek různobarevných virtuálních linek (pro rámce patříící do různých VLAN) a k tomu „dodatkovou“ linku pro nativní VLAN, po které jdou rámce nepatřící do žádné VLANy.

 Na každém switchi potřebujeme tabulku MAC adres. Protože přepínání má fungovat pouze mezi zařízeními patřícími do stejné VLAN (kromě trunku), musí být ke každé MAC adrese a portu ještě přidána informace o číslu VLAN.

Přepínání pak bude fungovat takto:

- Switch přijme na daném portu rámec, přičemž tento port patří do konkrétní VLAN.
- V tabulce MAC adres se zaměří pouze na záznamy patřící k téže VLAN a trunky, jiné nekontroluje.
- Mezi těmito „profiltrovanými“ záznamy najde ten, který odpovídá adrese cílového zařízení.
- Pokud vyjde přístupový port (s tímž číslem VLAN), pak na něj rámec jednoduše přepošle. Jestliže vyjde trunk (tj. k cíli vede cesta přes jiný switch), upraví rámec (později si ujasníme jak) a pošle do daného trunku.

Zjednodušeně si to můžeme představit tak, že na switchi je pro každou VLAN samostatná tabulka MAC adres (ve skutečnosti tomu tak není, protože bychom pak neměli kam zařadit ta zařízení, která do běžných „komunikačních“ VLAN nepatří). Pokud komunikace zůstává v rámci jednoho switchu, je přepínání rámců celkem jednoduché. Pokud například počítač P2 posílá rámec počítači P3, přijme switch S1 tento rámec na „fialovém“ portu (VLAN3) a tedy se řídí podle řádků tabulky pro VLAN3, kde je přímo uvedeno, na kterém portu je cíl (P3) dostupný.



Poznámka:

V předchozí kapitole jsme si ukázali, že pokud cíl není v tabulce MAC adres, switch použije flooding (záplavu) – pošle rámec na všechny porty kromě toho, ze kterého přišel. Ale když se používají VLANy, tak to asi bude jinak. Kam myslíte, že bude takový rámec poslán? A jak se řeší broadcasty?

Zmínka o broadcastech je nápoředou – již dříve tu bylo zmíněno, že různé VLANy odpovídají různým broadcastovým doménám, tedy broadcast se pošle jen na porty, které přísluší do stejné VLAN jako

odesílatel, a do trunků s touto VLAN nakonfigurovanou (jinými slovy, broadcast se nedostane do jiné VLANy). No, a rámec, jehož cílovou adresu nemáme v tabulce, dopadne stejně – pošle se na porty patřící do téže VLAN jako odesílatel (a na trunky s touto VLAN).



Poznámka:

Uvědomte si, že koncová zařízení naprosto netuší nic o nějaké virtualizaci. Když odesílají rámec, vědí jen to, na kterou MAC adresu má dojít, ale číslo VLAN neznají, toto číslo (zatím) není ani součástí odesílaného rámce. VLAN dokáže určit až switch, protože ten si ke každému portu eviduje informaci o VLAN, do které port (a k němu připojené zařízení) patří.



3.4.2 VLAN rámce na cestě přes trunk

Dále potřebujeme mechanismus, který nám umožní správně přepínat rámce, které mají jít od zdroje k cíli přes více než jeden switch. Dokud jsme zůstali v rámci jediného switchu, bylo to jednoduché – switch určil VLAN podle portu, ze kterého rámec přišel, a tím byly stanoveny i konkrétní řádky tabulky, na kterých má hledat adresu a port cílového zařízení.

Pokud však je cíl připojen k jinému switchi, pak je třeba rámec nejdřív odeslat na tento switch, který pak již zajistí předání cíli. Jenže jak „tomu druhému“ switchi předat informaci o tom, do které VLAN patřil zdroj, aby následující switch věděl, na které řádky se podívat?

Zaměříme se opět na obrázek 3.2. Předpokládejme, že počítač P1 posílá rámec počítači P6. Nějak to jít musí, protože oba počítače patří do VLAN2.

Switch S1 přijme rámec od počítače P1, podle tabulky zjistí, že ho má předat switchi S0, ale musí *přidat informaci o číslu VLAN*. Switch S0 přijme rámec včetně této dodatečné informace a obojí předá dál switchi S2. Až switch S2 zjistí, že tento rámec může předat přímo na port s přiřazeným číslem VLAN, a tedy dodatečnou informaci o číslu VLAN sice použije (aby určil tabulku adres), ale na cílový port odešle pouze rámec bez dodatečné informace.

 Spoj, kterým posíláme rámce patřící do různých VLAN, se nazývá *trunk*. Vede obvykle mezi dvěma switchi nebo mezi switchem a routerem, výjimečně mezi switchem a serverem. Z toho vyplývá, že switche S1–S3 mají oba druhy portů – přístupové i trunkové.

Při komunikaci mezi dvěma počítači patřícími do téže VLAN je vždy prvním a posledním spojem na cestě spoj u přístupového portu, mezilehlé spoje jsou trunkové. Pouze přes trunkové spoje se přenáší dodatečná informace o VLAN.

Nyní k té dodatečné informaci o VLAN. Už jsme si řekli, že při komunikaci v rámci jednoho switchu (tj. pouze přes přístupové porty) něco takového není zapotřebí (dokonce platí, že pokud přijde přes přístupový port rámec s VLAN informací, bude zahozen, protože pravděpodobně se někdo pokouší „vpašovat“ do VLAN, do které nepatří), ale při komunikaci přes trunkový port je tato informace nutná, protože další switch na cestě nemůže vědět, ze kterého přístupového portu dotýčný rámec původně přišel.

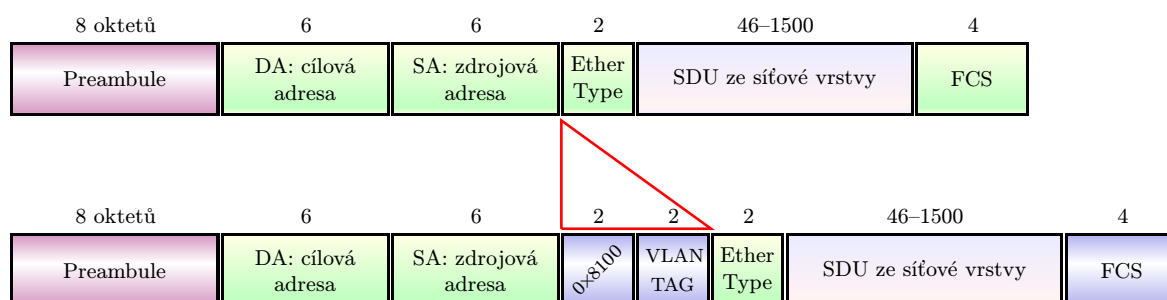
Existují dva přístupy: první mění záhlaví rámce (přidá informaci o VLAN), kdežto podle druhého přístupu rámec neměníme a místo toho jej zabalíme do speciálního rámce s informací o VLAN. První přístup je popsán protokolem IEEE 802.1Q, druhý je používán některými proprietárními protokoly.

 **VLAN podle IEEE 802.1Q.** Podle tohoto standardu se informace o VLAN vkládá přímo do

záhlaví přenášeného rámce, a to tak,

- aby bylo jasné, že to je VLAN rámec, a zároveň
- aby se ze záhlaví neztratily žádné původní položky.

Takže všechna původní pole necháme, ale dvě pole s informací o VLAN přidáme. Na obrázku 3.3 vidíme, kam jsou tato dvě pole vsunuta – hned za adresy.



Obrázek 3.3: VLAN rámec podle IEEE 802.1Q

Obsah nových polí je následující:

- *EtherType hodnota pro VLAN* (podle toho switch pozná, že doručený rámec obsahuje VLAN informaci). Pro VLAN to je hodnota 0×8100 .
- *VLAN tag*, který v sobě zahrnuje
 - prioritu podle IEEE 802.1p (3 bity), obvykle je tady 0 (to znamená Best Effort – bez priorit),
 - indikace kanonického tvaru adresy (1 bit), obvykle 0,
 -  číslo VLAN (12 bitů).

Protože se tímto změnil formát rámce, je třeba znovu vypočíst kontrolní součet (FCS).

Pokud switch (například switch S2 podle obrázku 3.2) přijme na trunkovém portu rámec, zkontroluje pole EtherType. Pokud tam najde číslo 0×8100 , je jasné, že se jedná o VLAN rámec, a v následujícím poli hledá identifikátor té VLAN, do které patří odesílatel.

 Switch má přepínat rámce mezi porty (z jednoho rámec přijme, na druhý ho pošle) podle tabulky MAC adres, a totéž je třeba provést i zde.

- Switch přijme rámec (bez VLAN ID) na přístupovém portu, podle portu určí VLAN. Pokud zjistí podle tabulky (vyfiltrováním řádků podle VLAN), že cílem nebude jiný přístupový port, přidá do záhlaví obě pole související s VLAN, vypočte kontrolní součet a odešle hotový VLAN rámec na trunkový port.
- Pokud switch přijme rámec na trunkovém portu a zjistí, že cíl je dostupný na přístupovém portu (určil VLAN podle záhlaví rámce a podle toho vyfiltroval řádky tabulky), odstraní ze záhlaví obě VLAN pole, vypočte kontrolní součet a odešle na dotyčný přístupový port.
- Pokud máme switche v stromové hierarchii, pak nejvyšší úrovně obvykle pracují pouze v režimu trunku, tedy nekontrolují VLAN ID – aby páteřní síť zbytečně nezdržovala provoz. Takže takový switch přijme rámec, přečte adresu cíle a jednoduše přepne na příslušný port.



Poznámka:

Pojmy *přístupový port* a *trunkový port* patří do terminologie Cisco. U jiných výrobců se můžeme setkat

s jinou terminologií (ale jinak to funguje úplně stejně, podle IEEE 802.1Q), například na zařízeních HP mluvíme o *untagged member* a *tagged member*.

K přístupovému portu (resp. untagged member) je připojeno zařízení, které „nerozumí“ VLAN (běžný počítač, server apod.), kdežto na trunkovém portu (resp. tagged member) je připojeno zařízení, které „rozumí“ VLAN (typicky switch) a umí pracovat s VLAN rámci.



 Pokud používáme VLAN podle portů (což je celkem běžné) ve velkých sítích, lze použít podpůrné protokoly, které umožní parametry VLAN spravovat automaticky, případně centrálně. Tyto protokoly bývají obvykle proprietární. V případě zařízení Cisco lze používat protokoly *VTP* (VLAN Trunking Protocol) a *DTP* (Dynamic Trunking Protocol).

VTP usnadňuje centralizovanou správu VLAN na switchích. V tzv. VLAN doméně je stanoven VTP server, ostatní jsou VTP klienti. Na VTP serveru vytvoříme všechny potřebné VLANy a ty se automaticky přenesou na klienty. Ovšem přiřadit porty do jednotlivých VLAN musíme ručně, účelem tohoto protokolu je pouze oprostit administrátory od ručního vytváření VLAN na switchích. Pokud by totiž byl některý port přiřazen do VLAN, která na daném zařízení není vytvořena, nefungoval by.

VTP je poněkud problematický protokol – jestliže uděláme v konfiguraci chybu (například na VTP serveru omylem odstraníme některou VLAN), pak se chyba distribuuje do celé VLAN domény a všechny porty, které jsou do této VLAN přiřazeny, přestanou fungovat. Jestliže se do konfigurace zařízení dá dostat pouze přes takový port, čeká nás osobní cesta do datového centra s postiženými switchi.

DTP umožňuje nechat trunk, aby se vytvořil automaticky v případě, kdy obě zařízení tuto možnost dovolují. Na Cisco switchi může být port v těchto stavech:

- access (ručně nastaveno na přístupový port)
- trunk (ručně nastaveno na trunk)
- dynamic auto (DTP je funkční, druhá strana může spustit dojednání trunku)
- dynamic desirable (DTP je funkční, tento port může spustit dojednání trunku)

Defaultní hodnota je dynamic auto, což znamená, že pokud se k portu připojí nepovolené zařízení, které se prohlásí za switch a spustí vyjednávání trunku, opravdu se vytvoří trunk se všemi souvisejícími bezpečnostními riziky. Toto nastavení je tedy vhodné změnit (pokud k portu nebude připojen „náš“ switch, pak ho přepneme do access módu).

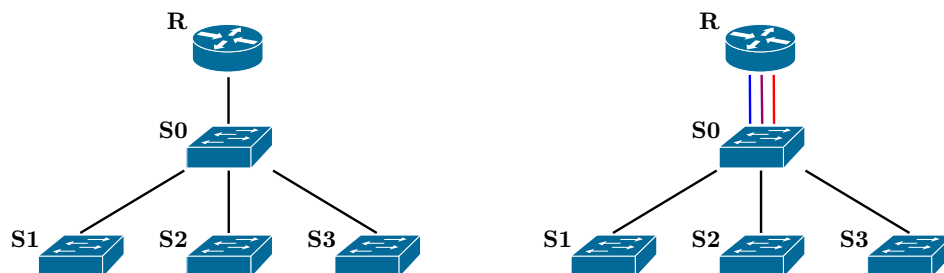
 **Další protokoly.** Druhý přístup implementace VLAN rámců spočívá v tom, že se do původního rámce nezasahuje, ale je zapouzdřen do speciálního rámce podle určitého protokolu – přidá se nové záhlaví obsahující kromě jiného identifikaci VLAN. Na (některých) zařízeních Cisco je k tomuto účelu používán *protokol ISL*.

Kromě ISL existují i další protokoly používané k tomuto účelu (proprietární od určitého výrobce, navzájem nekompatibilní). Vznikly ještě v době, kdy IEEE 802.1Q neexistoval, a proto tehdy pro jejich používání existoval důvod.

V současné době se pro identifikaci VLAN na trunkových portech téměř výhradně používá IEEE 802.1Q, tedy se jinými protokoly (včetně ISL) ani nebudeme zabývat. Navíc pokud konkrétní zařízení podporuje některý z proprietárních VLAN protokolů, pak obvykle podporuje taky IEEE 802.1Q.

3.4.3 Komunikace mezi různými VLAN sítěmi

Jenže my sice chceme tyto podsítě navzájem oddělit, ale ne úplně. Určitou možnost komunikace bychom chtěli zachovat i mezi různými VLANy, ale zároveň chceme tuto komunikaci určitým způsobem ovlivňovat (například skladníkovi nedovolíme navázat spojení se serverem s osobními údaji zaměstnanců), případně monitorovat, logovat. A právě k tomu nám může sloužit router, který je na obrázku připojený ke kořenovému switchi.



Obrázek 3.4: Možnosti pro zapojení L3 zařízení pro propojení VLAN sítí

Obecně – komunikaci mezi různými VLANy může zajistit pouze zařízení pracující na vrstvě L3: router nebo switch s funkcionalitou vrstvy L3. Dále budeme pro zjednodušení zmiňovat a zobrazovat většinou router.

Při komunikaci mezi zařízeními z různých VLAN je rámec posílán přes zařízení vrstvy L3. Jestliže v tomto případě zašle počítač P1 z „modré“ VLAN2 rámec počítači P5 z „fialové“ VLAN3, bude cesta v síti následující:

$$P1 - S1 - S0 - R - S0 - S2 - P4$$

přičemž na routeru proběhne určité testování a směrování a také se změní identifikátor VLAN v rámci (ve skutečnosti se na jakémkoliv L3 zařízení vždy odstraní původní záhlaví L2 a vytvoří nové). Proběhne směrování paketu vnořeného v rámci z podsítě určené pro VLAN zdroje do podsítě určené pro VLAN cíle, což lze jen na vrstvě L3.

Ovšem je třeba si uvědomit, že jakákoliv komunikace mezi zařízeními z různých VLAN opravdu půjde přes zařízení vrstvy L3, se všemi důsledky včetně prodloužení cesty rámce v síti a případně i snížení propustnosti spojů v síti (i kdyby zdroj a cíl byly připojeny k těmto switchi).



Poznámka:

Pro místo, které může brzdit provoz (včetně výše popsaného trunkového spoje) se vžil název *úzké hrdlo* (bottleneck). Pokud je síť špatně navržena, může L3 zařízení být úzkým hrdlem.



Různé spoje na L3 zařízení nutně musí být v různých podsítích (samozřejmě s nepřekrývajícím se adresním prostorem), tj. každý port by měl mít IP adresu ze „své“ vlastní podsítě. Ale není nutné, aby to platilo pro fyzické porty a fyzické spoje.

Jsou dvě možnosti, jak napojení VLAN na L3 zajistit. V prvním případě povedeme mezi L2 zařízeními a L3 zařízeními jediný fyzický kabel, který je na logické úrovni rozdělen na virtuální spoje – pro každou VLAN vlastní virtuální spoj, v druhém případě povedeme mezi těmito zařízeními tolik kabelů, kolik je VLAN (tj. každá VLAN bude mít vlastní fyzický kabel). V prvním případě hovoříme o virtuálních

subrozhraníh jednoho fyzického síťového rozhraní, a místo abychom nastavovali IP adresu na síťovém rozhraní, budeme ji nastavovat zvlášť na jednotlivých subrozhraníh. K oběma možnostem podrobněji:

 **Pokud router podporuje vytváření subrozhraní:** centrální switch a router propojíme z *pohledu switche* trunkovým spojem, *na straně routeru* vytvoříme na daném síťovém rozhraní tolik virtuálních subrozhraní, kolik VLAN bude moci projít přes trunk od switche. Každá VLAN bude mít své subrozhraní s vlastním označením a IP adresou. Situace je naznačena na obrázku 3.4 vlevo.

Konkrétně – na switchi tento port nakonfigurujeme jako trunkový, na routeru jde o síťové rozhraní s definovanými subrozhraními (například pro port `g0/1` to budou subrozhraní `g0/1.1`, `g0/1.2`, . . . , `g0/1.10`, . . . podle potřeby, za tečkou je číslo VLANy). IP adresu přiřazujeme jednotlivým subrozhráním, například subrozhraní `g0/1.10` bude mít IP adresu patřící do podsítě, kterou jsme určili pro VLAN 10.

 **Pokud máme router a L2 switch:** Starší způsob řešení spočíval v tom, že mezi routerem a switchem se vedlo tolik kabelů, kolik mělo existovat VLAN, pro každou VLAN jeden kabel (a také port na každém z těchto dvou zařízení). Tyto spoje pak na straně switche musí fungovat jako přístupové porty (tj. na každé je definována jedna konkrétní VLAN) a směrem k zařízení L3 (kde je síťové rozhraní s IP adresou patřící do podsítě pro příslušnou VLAN) půjdou již rámce bez VLAN polí.

Tento postup se týká jak routerů, které „neumí“ IEEE 802.1Q, tak i L3 switchů, které sice daný protokol zvládají, ale obvykle nepodporují vytvoření subrozhraní na jednom fyzickém rozhraní.

Takže zatímco za trunkovým portem bývají obvykle L2 switche, za přístupovým (access) portem bývají buď koncová zařízení, nebo L3 zařízení bez definovaných subrozhraní (což je z pohledu Ethernetu také koncové zařízení).

 **Máme multilayer switch:** Multilayer switche sice umí pracovat na vrstvě L3 jako routery, ale neumí vytvářet subrozhraní uvnitř jednoho rozhraní, ovšem ani to není třeba. Multilayer switch si můžeme představit jako L2 switch s integrovaným L3 modulem. Tento modul umí plnohodnotně směrovat, a taktéž může svou působnost přenést na některé porty. Multilayer switch tedy může mít dva druhy portů:

- switched port (přepínaný, L2 port) – funguje jako porty na obyčejných L2 switchích, včetně možnosti nastavení port security, může být v accessed nebo trunk módu,
- routed port (L3 port) – funguje (téměř) tak jako rozhraní na routeru (až na některé drobnosti, například nepodporuje subrozhraní), lze na něm nastavit IP adresy.

Zprovoznění přeposílání provozu mezi různými VLAN je právě na multilayer switchi nejjednodušší. Stačí zapnout L3 modul příkazem

```
ip routing
```

a máme hotovo. Jednotlivé porty zřejmě zůstanou přepínané, a pokud povedou k jinému switchi, pak zřejmě v trunk módu. Pokud mají pracovat v accessed módu, musíme na nich nastavit VLAN.



Poznámka:

Na předchozích stránkách bylo uvedeno, že koncová zařízení neobsahují implementaci protokolu IEEE 802.1Q (tedy „nerozumějí VLAN“). Existují však i výjimky. U serverů může být výhodné zařazení do více než jedné VLAN (abychom se vyhnuli neustálému přeposílání provozu přes zařízení vrstvy L3), což se právě dá udělat implementací IEEE 802.1Q přímo na tomto zařízení. Jak se to dá provést:

- Pokud se jedná o server s Linuxem, pak je to jednoduché, protože Linux samotný obsahuje implementaci IEEE 802.1Q.

- U serverů s Windows je situace složitější, protože samotné Windows tento protokol nepodporují. Jediná možnost, jak tam rozjet podporu VLAN, je pořídit speciální síťovou kartu, která IEEE 802.1Q implementuje.



 Řekli jsme si, že existuje *management VLAN* (VLAN určená pro správu, jsou po ní přenášeny například rámce s konfiguračními příkazy, např. SSH). Přístup na zařízení pro konfiguraci s pomocí Telnetu a SSH (vlastně i přes HTTP a HTTPS) funguje jen tehdy, když má zařízení nastavenou alespoň nějakou IP adresu (ta se uvádí jako parametr při navazování spojení), a měla by to být právě adresa z podsítě navázané na management VLAN.

Na L2 zařízení nenastavujeme IP adresu na konkrétním fyzickém rozhraní (to děláme jen na L3 zařízení), ale na virtuálním rozhraní (Switch Virtual Interface – SVI). Stejně jako koncová zařízení, i switch odpovídá na ARP/NDP dotazy typu „kdo má tuto IP adresu?“ ohlášením své IP a MAC adresy.

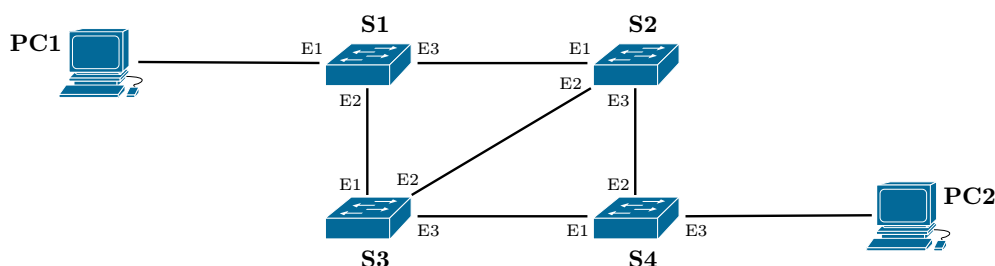
Nás spíše zajímá opačný problém – jak zajistit, aby „nepovolané osoby“ naopak to spojení navázat nemohly? Například právě tak, že určíme management VLAN a na L3 zařízení (třeba routeru) budeme určovat, kdo do podsítě navázané na management VLAN může a kdo ne. Ale pozor, tento postup není všemocný, dá se obejít, nejen využitím IP spoofingu (podvržení IP adresy).

3.5 STP – kostra v grafu

 V následujícím textu se zaměříme na postup, který je standardizován pro zařízení typu bridge, nicméně my budeme hovořit o switchi.

3.5.1 Protokol STP

Představme si větší síť, ve které máme několik switchů. Protože chceme, aby síť pracovala i v případě, že některá přenosová cesta bude přerušena nebo některý switch přestane být funkční, jsou switche fyzicky propojeny více než je nutné – některé cesty jsou *redundantní* (znásobené).



Obrázek 3.5: Síť se smyčkami

Redundantní cesty jsou praktické pro případ výpadku části sítě, ale mohou přinést určité problémy. Představme si několik možných scénářů:

1. Počítač PC1 odesílá rámec počítači PC2. První část cesty (switch S1) je jednoznačná, ale co potom? Má být rámec odeslán ze switche S1 přes port E2 nebo přes port E3?
2. Počítač PC2 odesílá broadcastový rámec. Switch S4 rámec přijme na portu E3 a odešle na porty E1 a E2. Co provedou ostatní switche?

- Switch S2 rámec přijme na portu E3 a odešle na porty E1 a E2.
- Switch S3 rámec přijme na portu E3 a odešle na porty E1 a E2. Tentýž rámec (o čemž neví) přijme i na portu E2 (od switchu S2) a odešle na porty E1 a E3.
- Switch S2 opět přijme rámec na portu E2 (od switchu S3) a odešle na porty E1 a E3.
- Switch S1 tentýž rámec přijme dvakrát z portu E2 a dvakrát z portu E3, odešle ho postupně celkem dvakrát na port E2, dvakrát na port E3 a čtyřikrát na port E1, čímž se rámec dostane k počítači PC1.
- Mezitím se broadcastový rámec několikrát dostal zpět ke switchi S4, počítači PC2 a následně opět k dalším switchům...

3. Předpokládejme, že se počítač PC1 připojuje do sítě, tedy zatím není v MAC tabulce žádného switchu. Počítač PC2 mu odešle rámec (MAC adresa počítače PC2 je cílová). Switch S4 tuto adresu nezná, tedy rámec odešle na porty E1 a E2, atd. – rámec putuje sítí stejným způsobem jako broadcastový. K počítači PC1 se sice dostane, ale ještě dlouho budou jeho kopie bloudit sítí.

Tato situace nemusí nastávat jen tehdy, když připojujeme nové zařízení do sítě. Uvědomte si, že v tabulce MAC adres na switchi jsou záznamy obvykle drženy jen po dobu jednotek minut (obvykle 5 minut, záleží na výrobci a případné změně konfigurace).



Poznámka:

Scénář popsáný v bodu 2 se nazývá *všesměrová bouře* (broadcast storm). Vyznačuje se tím, že sítí neustále bloudí totožné kopie téhož broadcastového rámce, čímž je síť defacto zahlcena a po určité době nepoužitelná.



Z toho vyplývá, že něco je špatně. Redundance cest je dobrá, ale zároveň bychom měli zabránit nekontrolovanému opakovanému bloudění rámců sítí. Tomu zabráníme, pokud ze sítě *odstraníme smyčky* (co se týče komunikace; ve fyzické topologii mohou smyčky kvůli redundanci zůstat).

 **Protokol STP** (Spanning Tree Protocol, také protokol kostry v grafu) je standardizován jako IEEE 802.1D jako mechanismus odstranění (logických) smyček v grafu sítě mostů (tedy v našem případě switchů). Tento protokol „vidí“ pouze mosty (resp. switchy), žádná jiná zařízení ho nezajímají (takže v naší síti bude počítače PC1 a PC2 ignorovat).

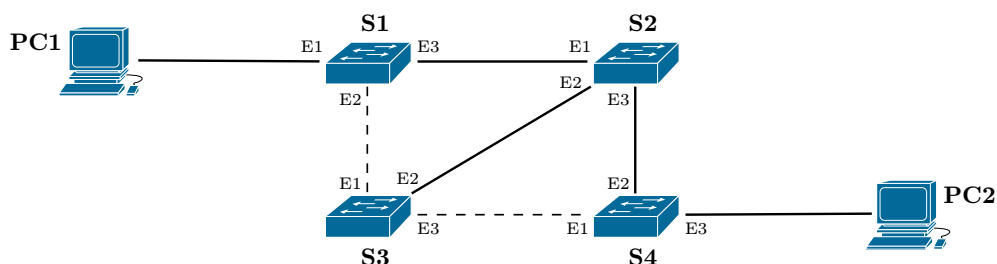
Účelem je překonfigurovat switchy v síti takovým způsobem, aby některé porty nebyly pro komunikaci používány (jsou buď deaktivovány nebo používány jen pro servisní účely), čímž ze sítě odstraníme smyčky. Protokol STP v podstatě řeší „problém hledání kostry v grafu“, jehož účelem je v grafu nechat pouze takové spoje, aby mezi kterýmikoliv dvěma uzly sítě (switchy) existovala právě jedna cesta (ne více, ne méně), ideálně ta nejrychlejší.

 Vlastně ze sítě switchů vytvoříme *strom* („spanning tree“ znamená pokrývající strom). Jeden ze switchů je prohlášen za kořen stromu (*kořenový switch*) a z každého dalšího switchu je hledána nejrychlejší cesta ke kořenovému switchi. Všechny ostatní cesty jsou deaktivovány. V našem případě by mohl být kořenovým switchem třeba switch S2 a v síti by zůstaly aktivní pouze spoje S1–S2, S3–S2 a S4–S2. Spoj S1–S3 a S3–S4 by byly deaktivovány.



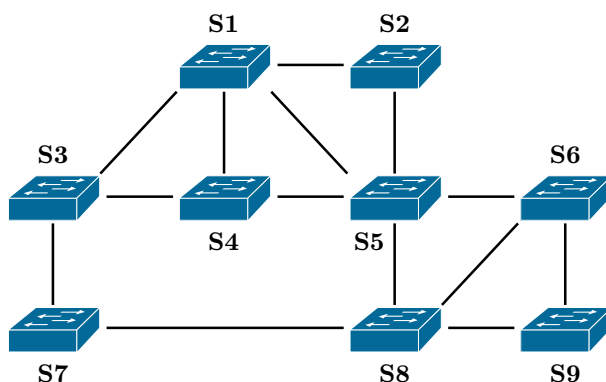
Příklad

V síti na následujícím obrázku vidíme celkem devět switchů, které jsou propojeny ve fyzické topologii

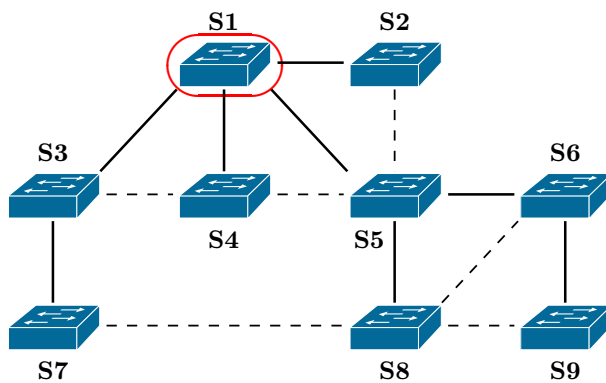


Obrázek 3.6: Síť po odstranění smyček

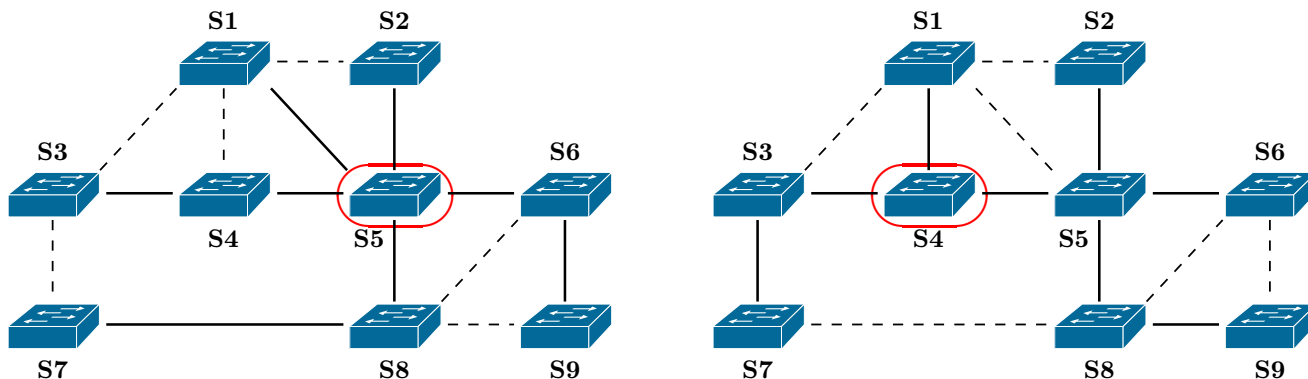
mesh. Máme tady smyčky, které je třeba vyřešit pomocí protokolu STP.



Prvním úkolem je zvolit kořenový switch. Máme na výběr, ovšem v reálné situaci bychom pro tuto volbu zřejmě měli určitá kritéria. Zde jednoduše zvolíme switch S1. Některé porty blokujeme, čímž pro běžnou komunikaci zpřístupníme celkem šest cest.



Jsou i další možnosti – jako kořenový switch můžeme zvolit třeba switch S5 nebo S4:



Také v těchto případech se některé porty deaktivovaly pro běžný provoz, přičemž od každého switchu existuje právě jedna (pokud možno neoptimálnější) cesta ke kořenovému switchi.

Všimněte si, že v každém z těchto tří řešení jsou cesty mezi některými dvojicemi switchů různě dlouhé. Také je třeba brát v úvahu, že tyto nákresy jsou pouze topologie, neberou v úvahu skutečné délky spojů a jejich vytížení, takže v reálné situaci by administrátor při volbě kořenového switchu promýšlel různá kritéria.



3.5.2 Konvergence sítě switchů

Podle čeho určuje protokol kořenový switch? Každý switch, který „rozumí“ protokolu STP (je na něm tento protokol implementován), má přiřazen speciální identifikátor (bridge ID, resp. switch ID) o délce 8 oktetů, který se skládá ze dvou částí:

- Priorita (2 oktety)
- MAC adresa (6 oktetů).

1	2	3	4	5	6	7	8
Priorita		MAC adresa					

Standardně je v poli pro prioritu číslo 0×8000 , což je desítkově 32 768. Přesto žádné dva switchy nemají tento identifikátor stejný (liší se alespoň v MAC adrese, každý switch má jinou).

 Protokol STP vybere jako kořenový strom ten switch, jehož switch ID je *nejmenší číslo*. Pokud bychom spoléhali jen na tento postup, mohl by být jako kořenový zvolen ten switch, u kterého se nám to moc nehodí – řešením je změnit prioritu námi vybraného switchu na *menší hodnotu*. Díky tomu, že v switch ID je nejdřív priorita a pak až MAC adresa, má tato změna vždy očekávaný účinek.

Zbývá stanovit, jak určíme, které spoje (resp. porty) mají být blokovány a které mají zůstat aktivní. Pokud ke kořenovému switchi vede cesta přes víc než jeden port, pak se pro každý port vypočte optimalita cesty (záleží na šířce pásma/rychlosti, vytížení, hovoříme o ceně cesty). Aktivní bude jen port pro neoptimálnější cestu. Pokud pro oba porty vyjde stejná cena cesty, použije se při rozhodování switch ID sousedů, ke kterým tyto porty vedou – soused s nižším switch ID vyhrává, spoj k němu vedoucí zůstane aktivní.

 Takže port buď je nebo není aktivní z pohledu přeposílání rámců s běžným provozem. Tuto vlastnost označujeme jako *stav*, rozlišujeme tyto stavy portu:

- *předávání* (forwarding) – funguje, přeposílá veškerý provoz,
- *blocking* – nepřeposílá běžný provoz, ale přijímá správní rámce protokolu STP (tzv. BPDU rámce), aby se v případě potřeby dokázal aktivovat,
- *naslouchání* (listening) a *zjišťování* (learning) – mezilehlé stavy pro případ, že port má přejít ze stavu blokování do stavu předávání; ve stavu listening switch zpracovává BPDU rámce, ve stavu learning přijímá (ale nepřeposílá, zahazuje) rámce s běžným provozem a začíná si tvořit tabulku MAC adres,
- *disabled* – zcela vypnutý port, nepřeposílá vůbec nic.

Porty se mezi těmito stavy přesouvají, třebaže většinu času stráví právě ve stavech předávání nebo blokování.

 Čas od času (především při změnách v síti) probíhá rekonfigurace – *proces konvergence*. Účelem je zajistit, aby byla síť v konvergentním stavu (tj. v našem případě aby v síti nebyly žádné smyčky a aby

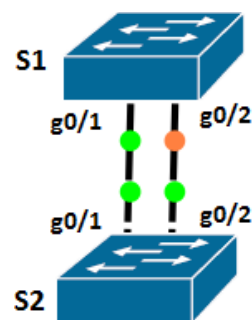
byly aktivní pouze porty pro neoptimálnější cestu). Může být nutné určit nový kořenový switch, a také stavy portů se mohou měnit (například pokud má port přejít ze stavu blokování do stavu předávání, stráví určitý čas ve stavech naslouchání a zjišťování, aby si dokázal nastavit správně všechny parametry a doplnit tabulku MAC adres).

 Každý port plní určitou roli, která je dána pozicí switche v hierarchii. Podle role se v dané chvíli určuje, v jakém stavu tento port bude. Jaké role tedy existují?

- *root port* – port, který směřuje k root switchi, tedy každý switch kromě kořenového má jeden root port mířící „nahoru“ (předpokládejme, že máme switche nakresleny tak, že root je nahoře), root port je vždy aktivní, tj. ve stavu forwarding,
- *designated port* – port, který je aktivní a v hierarchii míří dolů, tedy k podřízeným switchům (takže za designated portem máme podstrom switchů); kořenový switch má všechny své porty kromě blokových v roli designated, naopak switche v hierarchii zcela dole (takové, od kterých je ke kořenovému switchi nejdelší cesta) nemají žádné designated porty,
- *blocked port* – tento port nebyl vybrán, nepoužívá se pro běžný provoz (tedy obvykle je ve stavu blocking), vyšší verze STP místo tohoto stavu rozlišuje stavy
 - backup port – blokový port, který je záložním portem k root portu (tj. když selže cesta nahoru ke kořenovému switchi, tento port se může stát root portem),
 - alternate port – blokový port, který je záložním portem k některému designated portu (tj. když selže některá cesta vedoucí k podřízeným switchům, alternate port se stane designated portem a začne předávat provoz).

 Speciálním případem je situace, kdy právě mezi dvěma sousedními switchi máme natažené redundantní spoje (prostě mezi dvěma switchi vedou dva kabely), situace je naznačena na obrázku 3.7. Tuto situaci nazýváme *kravata (tie)*. To, který switch bude nadřízený (tj. z těch dvou potenciálně root), se určí tak, jak bylo výše naznačeno (tj. switch ID), ale co s těmi dvěma spoji, který z nich bude aktivní?

Pro oba tyto spoje platí logicky stejná cena cesty a nedá se rozhodnout ani podle switch ID souseda, protože za oběma porty vlastně mám téhož souseda. V tomto případě se použije další kritérium – i na portech totiž máme dvojici údajů: prioritu portu a číslo portu. Priorita portu obvykle bývá číslo 128, takže rozhodující je číslo portu. Prioritu samozřejmě můžeme změnit a tím ovlivnit určení aktivního portu.



Obrázek 3.7: STP tie

Příklad

Na obrázku 3.7 máme switche S1 a S2. Podle LED světél můžeme usoudit, že kořenovým switchem se stal S2 a spoj mezi porty g0/2 na obou switchích bude blokový. Výpis základního nastavení STP na switchi S1:

```
S1#show spanning-tree
VLAN0001
  Spanning tree enabled protocol ieee
  Root ID    Priority    32769
             Address     0007.EC99.A154
             Cost        4
```

```

Port 25(GigabitEthernet0/1)
Hello Time 2 sec Max Age 20 sec Forward Delay 15 sec

Bridge ID Priority 32769 (priority 32768 sys-id-ext 1)
Address 000D.BDC5.2941
Hello Time 2 sec Max Age 20 sec Forward Delay 15 sec
Aging Time 20

Interface Role Sts Cost Prio.Nbr Type
-----
Gi0/1 Root FWD 4 128.25 P2p
Gi0/2 Altn BLK 4 128.26 P2p

```

V části výpisu Bridge ID vidíme prioritu a MAC adresu switchu S1, nad tím v části výpisu Root ID jsou tyto údaje o kořenovém switchi (každý switch musí znát „svého roota“). V částech Root ID a Bridge ID jsou různé MAC adresy, z toho plyne, že S1 není kořenovým switchem. Dole je pak tabulka se všemi používanými porty, ke každému vidíme roli, stav, cenu spoje, prioritu portu s číslem portu a typ linky.

Port g0/1 je root port, tj. vede ke kořenovému switchi, a je ve stavu forwarding (takže předává veškerý provoz, běžně funguje). Port g0/2 má roli Altn, což je „alternated“, neboli je blokový (právě u něj je oranžové světlo). Číslo těchto portů nejsou 1 a 2, ale 25 a 26 (předposlední sloupec) – určují se globálně jednoznačně za celý switch, nižší čísla mají fast-ethernetové porty f0/1–f0/24.

Tentýž výpis na switchi S2:

```

S2#show spanning-tree
VLAN0001
Spanning tree enabled protocol ieee
Root ID Priority 32769
Address 0007.EC99.A154
This bridge is the root
Hello Time 2 sec Max Age 20 sec Forward Delay 15 sec

Bridge ID Priority 32769 (priority 32768 sys-id-ext 1)
Address 0007.EC99.A154
Hello Time 2 sec Max Age 20 sec Forward Delay 15 sec
Aging Time 20

Interface Role Sts Cost Prio.Nbr Type
-----
Gi0/1 Desg FWD 4 128.25 P2p
Gi0/2 Desg FWD 4 128.26 P2p

```

Oba switchy mají stejnou prioritu (32769), liší se adresou. MAC adresa switchu S2 je menší číslo, tedy S2 se stává rootem.

Porty g0/1 a g0/2 jsou oba v roli designated a oba předávají, což se může zdát zvláštní, když vlastně protějšek toho druhého je blokový. Jenže toto je běžný případ – blokována je jen jedna strana linky, aby v případě nutnosti bylo možné ji co nejrychleji aktivovat. Takže na straně nadřazeného switchu je port v roli designated ve stavu preposílání, na druhé straně je port v roli alternate ve stavu blokový. Běžný provoz neprojde (je zahozen na blokový port), takže smyčka se nevytvoří.



 Aby to vše mohlo fungovat, switchy se navzájem domlouvají pomocí speciálních rámců, které označujeme *BPDU* (Bridge PDU). Slouží k údržbě a případné aktualizaci informací o síti potřebných pro

STP, posílají se na multicast adresu označující všechna zařízení se spuštěným STP (01:80:c2:00:00:00). BPDU se odesílají pravidelně každé 2 sekundy, a pak při každé změně sítě switchů.

Součástí BPDU je kromě jiného informace o tom, že se jedná o rámec STP a jaké verze, který switch je kořenový (jeho bridge ID/switch ID), dále ID odesílajícího switche, údaje o době platnosti rámce, cena cesty od odesílatele ke kořenovému switchi, atd. Tyto údaje jsme viděli v příkladu na výpisech. BPDU se zapouzdřuje do LLC rámce podle IEEE 802.2 a následně do MAC rámce IEEE 802.3.

 <https://wiki.wireshark.org/STP>

 Pokud se do sítě switchů začleňuje nový switch, musí se také zapojit do STP stromu. Nový switch nejdřív sám sebe považuje za kořenový switch, ale své porty zatím nechává ve stavu naslouchání (přijímá BPDU, ale nic neodesílá). Pokud v přijatém BPDU najde ID kořenového switche, které je nižší než jeho ID, přestane se považovat za kořenový switch a postupně se zařadí do struktury (pro každý port si z přijatých BPDU vypočte cenu cesty k ostatním switchům a pro každý switch si vybere jen jeden port, ostatní blokuje).

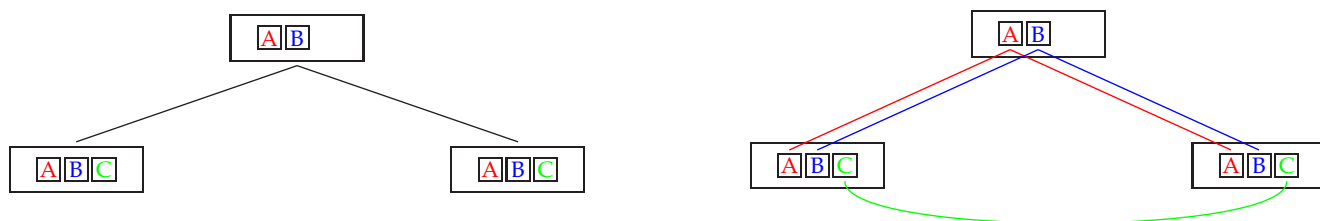
3.5.3 Varianty protokolu STP

Původní protokol STP byl standardizován jako IEEE 802.1D.

 **RSTP.** Protože u rychlejších standardů přestala původní specifikace stačit, byl vytvořen nový dodatek standardu nazvaný RSTP (Rapid STP) jako IEEE 802.1w, nicméně nyní je zahrnut do revize IEEE 802.1D-2004 (takže původní STP a novější RSTP splynuly v jednom standardu, třebaže někteří výrobci v konfiguraci rozlišují původní STP a novější RSTP).

Zatímco u původního STP trval proces konvergence 50 sekund (po tuto dobu byla celá síť nefunkční, což je strašně dlouho), u novějšího RSTP to trvá kolem 1 sekundy, i méně. Mezi původním STP a RSTP jsou sice i jiné rozdíly, ale tento rozdíl je nejvýraznější.

 **MSTP.** Varianta MSTP (Multiple STP) byla standardizována jako dodatek IEEE 802.1s, také se označuje jako MST (Multiple Spanning Tree). Účelem je zefektivnit činnost protokolu STP pro případ, že jsou v síti definovány virtuální sítě (VLAN). Taktéž byla dodatečně přidána do standardu, ale tentokrát IEEE 802.1Q, a to jako revize IEEE 802.1Q-2005 jako rozšíření RSTP pro VLAN sítě.



Obrázek 3.8: Přerušená cesta pro jednu VLAN při použití STP

V principu jde o to, aby pro každou VLAN nebo skupinu VLAN mohla (nemusela) existovat samostatná instance protokolu STP, tedy pro každou VLAN (či několik) by mohla být trochu jiná struktura aktivních linek a taky jiný kořenový switch. Účelem je jak zefektivnění komunikace v rámci jednotlivých VLAN (aby například zbytečně nebyla vedena komunikace přes takové switchy, které s dotyčnou VLAN nemají co do činění), ale také možnost vyrovnaní zátěže v síti (určitý port bude pro jednu VLAN blokován a pro jinou aktivní, u jiného portu to může být naopak, každá bude využívat jiné cesty).

Pokud bychom v prostředí s VLAN sítěmi použili pouze STP, je možné, že bychom některé VLANy na určitých spojkách prakticky vyřadili z činnosti nebo bychom některé z linek zbytečně zatížili příliš. Na obrázku 3.8 vidíme podobný problém: VLAN C (zelená) je aktivní jen na spodních dvou switchích. Pokud bychom použili klasický STP bez podpory VLAN, zbytečně bychom rámcům z VLAN C prodloužili cestu. Jestliže ale použijeme MSTP, můžeme pro zelenou VLAN určit jiný kořenový switch než pro ostatní, třeba jeden ze spodních switchů, a spoj mezi nimi necháme aktivní (jen pro tuto VLAN).



Další informace:

<http://www.samuraj-cz.com/clanek/cisco-ios-10-rapid-spanning-tree-protocol/>



Kromě standardizovaných variant existují i proprietární varianty vytvořené různými výrobci. Se standardy to je totiž tak, že nepřicházejí tak rychle, jak by si komerční prostředí představovalo, a tak se často stává, že nejdřív výrobce používá svůj vlastní proprietární protokol (pouze pro zařízení jeho značky), a až se objeví standard, doplní ho do specifikace, takže zařízení „umí“ více různých vzájemně alternativních protokolů.

Například u Cisca se můžeme setkat s protokolem PVST+, který do původního IEEE 802.1D přidává podporu VLAN (pro každou VLAN vlastní STP instance), a Rapid PVST+, který zároveň s podporou VLAN urychluje konvergenci jako v RSTP. Ale pozor – Rapid PVST+ a MSTP nemají úplně stejné vlastnosti: zatímco Rapid PVST+ opravdu vytváří pro naprosto každou VLAN samostatnou STP instanci v vlastním kořenovém switchem apod., MSTP dovoluje vytvořit více než jednu instanci, ale nevynucuje si vytvoření tolik instancí, kolik je VLAN (může být v jedné instanci i více než jedna VLAN), což v případě velkého množství VLAN šetří výpočetní a přenosové kapacity sítě switchů.

3.6 EtherChannel

3.6.1 Vlastnosti



EtherChannel je technologie umožňující spojit (agregovat) až 8 fyzických ethernetových linek do jedné. Dva uzly (typicky servery, přepínače nebo směrovače podporující tuto technologii) jsou propojeny více aktivními fyzickými ethernetovými spoji, PDU procházejí některým z těchto spojů. Dále je podporováno navíc až 8 záložních fyzických linek (failover) pro případ, že aktivní linky budou nefunkční.



Poznámka:

V případě této technologie se používá dvojí terminologie – zatímco u firmy Cisco se mluví o EtherChannel, v systému Solaris se pro totéž používá pojem *trunk*, což ale u Cisco znamená něco úplně jiného. Proto když slyšíme pojem „trunk“, měli bychom mít jasno v tom, čí terminologie je vlastně právě používána. U serverů se také používá pojem *NIC Teaming*, což však není úplně totéž (je to širší termín, sdružování síťových karet do jedné „virtuální“, logické).



Spoje, které jsou součástí EtherChannelu, jsou typu Fast Ethernet nebo rychlejší, a všechny musejí být stejného typu (tj. například nelze kombinovat spoje Fast Ethernet a 10Gb Ethernet – stejná rychlost, dále stejný duplex apod.). Pokud používáme VLAN (virtuální síť), měly by spoje být v rámci jedné VLAN nebo v trunk módu se stejnými parametry.

Celková propustnost se však nerovná součtu propustností jednotlivých fyzických cest. Pokud se v rámci jedné komunikace směřující k jednoho uzlu k jinému uzlu (ve výchozím nastavení) posílá více PDU, všechny jsou posílány po téže fyzické cestě, tedy komunikaci nelze rozdělit.

3.6.2 Řízení toku

 Výchozí nastavení je takové, že spoje v rámci EtherChannelu jsou rozdělovány podle cílové MAC adresy, tedy PDU, které mají konkrétní cílovou MAC adresu, jsou všechny směřovány na tentýž spoj. To však ne vždy vyhovuje, proto je možné nastavit i jiné dělení:

- dvojici [zdrojová MAC, cílová MAC],
- pouze podle zdrojové MAC,
- podobně pro IP adresy a porty.

Rozdělení zátěže mezi linky v rámci EtherChannelu se provádí podle hashovacího vyvažovacího algoritmu, který třídí PDU podle zadaného kritéria (například podle cílové MAC adresy) a celou komunikaci dělí podle pravidla naznačeného v tabulce 3.4. Každému „proudu dat“ určenému podle zvoleného kritéria je hashováním přiděleno číslo z intervalu 0–7 (tj. jedno z 8 čísel), a to bez ohledu na skutečný počet spojů v EtherChannelu. Toto číslo pak určuje konkrétní fyzický spoj.

Například pokud máme EtherChannel nad třemi fyzickými spojeními (EC porty), čemuž odpovídá předposlední sloupec tabulky vpravo, budou na první spoj směřovány PDU se třemi konkrétními přiřazenými čísly (0, 1, 2), jiná tři čísla (3, 4, 5) budou směřována na druhý spoj a zbylá dvě (6, 7) na třetí spoj. Je zřejmé, že neoptimálněji budou spoje vyváženy, pokud do EtherChannelu sdružíme 2, 4 nebo 8 spojů.

		Počet EC portů						
		8	7	6	5	4	3	2
Rozdělení komunikačních proudů mezi EC porty	1	1	2	2	2	2	3	4
	1	1	1	2	2	2	3	4
	1	1	1	1	2	2	2	
	1	1	1	1	1	2		
	1	1	1	1	1			
	1	1	1	1				
	1	1	1					
	1	1						

Tabulka 3.4: Hashování komunikace v EtherChannelu

3.6.3 Vytvoření kanálu

EtherChannel lze buď vytvořit manuálně napevno, stav je označován *on*. Jiná možnost je nechat za nás pracovat vyjednávací protokol.



Pro EtherChannel (EC) existují dva *protokoly vyjednávací vlastnosti* EC spojení:

- *LACP* (Link Aggregation Control Protocol) definovaný IEEE 802.3ax¹,
- *PagP* (Port Aggregation Protocol), což je proprietární protokol společnosti Cisco.

První z těchto protokolů je obecně použitelný ve všech zařízeních podporujících EtherChannel, druhý se používá pouze při propojování dvou zařízení Cisco. Na všech spojkách mezi dvěma zařízeními být použit vždy jen jediný z těchto protokolů (a obě zařízení mu musí „rozumět“).



Oba protokoly umožňují nastavit zařízení do jednoho ze dvou stavů – aktivního (*active* u LACP, *desirable* u PagP), ve kterém zařízení aktivně vyjednává sestavení EtherChannelu, a pasivního (*passive* u LACP, *auto* u PagP), ve kterém zařízení pasivně vyčkává na druhou stranu.

- Pokud zvolíme protokol LACP, musí být všechny porty na jedné straně spoje v módu *active*, na druhé straně je u portů přípustný buď *passive* nebo taky *active*, jinak se neaktivuje EtherChannel.

¹Správně je IEEE 802.3ax, ale v literatuře se můžeme setkat s nesprávným (významově posunutým) IEEE 802.3ad, což je NIC Teaming.

- Pokud zvolíme protokol PAgP, musí být port na jedné straně nastaven do módu desirable, na druhé straně je přípustný buď auto, nebo taky desirable (jinak spoj nebude fungovat).
- Pokud jsou porty ve svazku na jednom zařízení nastaveny do módu on, na druhé straně spoje taky musí být on (tato situace znamená, že jsme pro vyjednání kanálu nepoužili žádný protokol, ale ručně jsme určili, že má jít o EtherChannel).

Porty na prvním zařízení	Porty na druhém zařízení	⇒ vybrali jsme protokol
active	passive nebo active	LACP
desirable	auto nebo desirable	PAgP
on	on	žádný

Tabulka 3.5: Módy pro vyjednání EtherChannelu

Sebemenší chybička nebo nekonzistence způsobí, že EtherChannel nebude fungovat. Musíme si pohlídat jak rychlost, duplex a správný mód, tak i například nastavení VLAN. Pokud jde o přístupové porty (což je pro spoj mezi dvěma switchi méně běžné), musí být na všech portech ve svazku a na obou stranách nastaveno totéž číslo VLAN. Pokud jde o trunkové porty (přeposílající rámce podle IEEE 802.1Q, to je u spoje mezi switchi nejběžnější), musí být trunk na obou stranách a taktéž musí sedět seznam povolených VLAN, které mohou do trunku být poslány.

**Poznámka:**

Uvědomte si, že se nám tady míchají dva protichůdné mechanismy: zatímco STP zajišťuje, že pokud je mezi dvěma zařízeními více než jedna fyzická linka, bude aktivní jenom jedna z nich, kdežto EtherChannel potřebuje, aby mezi dvěma zařízeními vedlo více aktivních fyzických linek (které na logické úrovni budou chápány jako jediná). Takže pokud propojíme dvě zařízení více než dvěma kabely, okamžitě najede STP, a až po správné konfiguraci EtherChannel budou opravdu aktivní všechny fyzické linky.

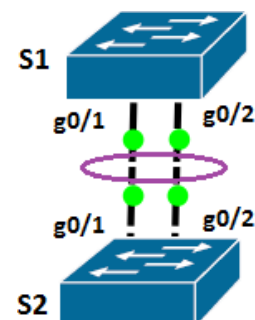
**Postup**

Obrázek 3.9 je podobný obrázku 3.7 s tím rozdílem, že zelená LED je na všech portech, tedy opravdu všechny porty mohou přeposílat, všechny fyzické linky jsou aktivní a STP nám tu nic neblokuje.

Podíváme se na zkrácený výpis nastavení etherchannelu na switchi S1:

```
S1#show interface etherchannel
GigabitEthernet0/1:
Port state      = 1
Channel group   = 1      Mode = Active      Gchange = -
Port-channel    = Po1    GC      = -      Pseudo port-channel = Po1
Port index      = 0      Load = 0x00     Protocol = LACP
... (zkráceno)

GigabitEthernet0/2:
Port state      = 1
Channel group   = 1      Mode = Active      Gchange = -
Port-channel    = Po1    GC      = -      Pseudo port-channel = Po1
Port index      = 0      Load = 0x00     Protocol = LACP
... (zkráceno)
```



Obrázek 3.9: Etherchannel mezi dvěma switchi

Takže na switchi S1 jsme na portech g0/1 a g0/2 nastavili mód active (tj. vybrali jsme protokol LACP pro dojednání parametrů), takže na portech druhého switchu musí být passive nebo active (lepší je passive, ať se „nehádají“, kdo začne vyjednávat). Tentýž výpis na switchi S2:

```
S2#show interface etherchannel
GigabitEthernet0/1:
Port state      = 1
Channel group   = 1      Mode = Passive      Gcchange = -
Port-channel    = Po1    GC      = -          Pseudo port-channel = Po1
Port index      = 0      Load = 0x00      Protocol = LACP
... (zkráceno)

GigabitEthernet0/2:
Port state      = 1
Channel group   = 1      Mode = Passive      Gcchange = -
Port-channel    = Po1    GC      = -          Pseudo port-channel = Po1
Port index      = 0      Load = 0x00      Protocol = LACP
... (zkráceno)
```

Na switchi S1 jsme postupovali takto:

```
S1>enable
S1#configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
S1(config)#vlan 10
... (vytvoříme všechny VLAN, které chceme přenášet)

S1(config)#interface range g0/1-2
S1(config-if-range)#shutdown
S1(config-if-range)#channel-group 1 mode active
S1(config-if-range)#interface port-channel 1
S1(config-if)#switchport mode trunk
S1(config-if)#switchport trunk native vlan 98
S1(config-if)#switchport trunk allowed vlan 10,20,30,98,99
S1(config)#interface range g0/1-2
S1(config-if-range)#no shutdown
```

Pokud to je možné, je dobré konfigurovat vypnuté porty. Abychom měli jistotu, že všechny porty v kanálu jsou nastaveny naprosto stejně, nejdřív vytvoříme etherchannel a pak nastavujeme (třeba VLAN trunk) na tomto kanálu.



Etherchannel sice podle názvu vypadá, že by se měl vztahovat k vrstvě L2 a Ethernetu, ale ve skutečnosti lze vytvořit Etherchannel kanály i na vrstvě L3.



Další informace:

- <http://www.samuraj-cz.com/clanek/cisco-ios-21-etherchannel-link-agregation-pagp-lACP-nic-teaming/>
- <http://www.fit.vutbr.cz/~ivesely/publikace/etherchannel.pdf>
- http://www.cisco.com/en/US/tech/tk389/tk213/technologies_tech_note09186a0080094714.shtml
- <http://www.ciscozine.com/2008/11/04/configuring-link-aggregation-with-etherchannel/>



3.7 SAN síť

 *SAN* (Storage Area Network) jsou sítě určené k propojení datových zařízení (obvykle serverů) s aktivními prvky sítě, jde tedy o zajištění vysoké dostupnosti úložišť dat, obvykle v datových centrech. K nejznámějším SAN technologiím patří Fibre Channel (FC) a iSCSI.

K zařízením ukládajícím data se můžeme přes počítačovou síť dostat několika způsoby:

- souborový přístup přes vhodný aplikační protokol (SMB, FTP, NFS apod.), kdy komunikujeme se souborovým systémem na daném zařízení a vzdáleně pracujeme v podstatě podobně jako lokálně (stáhneme si k sobě soubor nebo naopak uploadujeme, pracujeme s oprávněními u souborů, atd.), nízkourovňová komunikace probíhá uvnitř úložného zařízení mezi disky a místním souborovým systémem,
- blokový přístup na nižší úrovni, tj. souborový systém na našem zařízení přes síť posílá ty příkazy, které v předchozím případě probíhaly lokálně na úložném zařízení.

S prvním způsobem se setkáváme na jednoduchých NAS zařízeních (Network Attached Storage), tedy malých serverech, které máme třeba doma nebo v malých firmách (popřípadě tak funguje běžná komunikace se souborovými servery přes Internet), druhý způsob jsou právě SAN sítě v datových centrech.

Víme, že disky a jiná úložná zařízení mají kromě hardwarového datového rozhraní (SATA, SAS apod.) také signální rozhraní, tj. sadu příkazů, pomocí kterých se s těmito zařízeními dá komunikovat. SAN sítě přenášejí tyto příkazy, a většinou jde o sadu SCSI vhodné verze, ať už se jedná o jakýkoliv typ SAN (FC nebo iSCSI).

3.7.1 Fibre Channel

 Fibre Channel (FC) je převážně optické rozhraní (tj. na optických kabelech) pro rychlé přesuny dat na (relativně) krátké vzdálenosti, které pracuje na stejných vrstvách jako Ethernet, většina logiky technologie je v MAC podvrstvě. Původně bylo určeno pro použití „uvnitř počítače“ (konkrétněji serveru) ke komunikaci mezi I/O zařízeními a procesory, ale postupně se prosadilo jako síťový standard. Pod pojmem *kanál* (channel) se zde rozumí přímé nebo přepínané predikovatelné propojení dvou zařízení.

Podle názvu by se mohlo zdát, že jako přenosový prostředek je možné použít pouze optický kabel, ale ve skutečnosti je možné použít také koaxiál nebo STP (stíněnou kroucenou dvojlinku), což ale má vliv na dosažitelné vzdálenosti. Zatímco optika může vést na stovky metrů až kilometry, STP je použitelný jen na vzdálenost 50 nebo 100 m (podle toho, jakou požadujeme rychlost).

Přenosová rychlost závisí na verzi, nejnovější verze dosahují rychlosti v jednotkách až desítkách Gb/s pro každý směr.

Jde o plně duplexní sériové rozhraní pro point-to-point spoje (i více point-to-point spojů s vhodným hardwarem, přepínačem). Přepínání je zde velmi jednoduché a rychlé, bývá realizováno hardwarově. Výhodou je tedy vysoká rychlost přepínání a tedy i dostupnosti dat.

 *Protokol FCP* (Fibre Channel Protocol) slouží k efektivnímu přenosu nízkourovňových příkazů pro práci s vnějšími paměťmi přes počítačovou síť. Stroj, ze kterého chceme přes FC síť přistupovat k některému datovému úložišti v datovém centru, potřebuje speciální síťové rozhraní *HBA* (Host Bus Adapter), díky kterému to vypadá, jako bychom přistupovali k diskům ve svém zařízení.

V FC jde obvykle o přenášení příkazů ze sady SCSI, ale novější specifikace si poradí také s komunikací NVMe u rychlejších SSD.

Takže v FC síť máme optické kabely (dnes obvykle se zakončením v SFC+ nebo podobném modulu, případně konektor typu LC), které se zapojují do:

- FC řadiče diskových polí v úložném zařízení,
- HBA karty na zařízení, které chce vzdáleně k úložnému zařízení přistupovat,
- FC switche.

Disky v diskovém poli rozhodně nemusí mít hardwarové rozhraní SCSI, resp. novější SAS, nebo přímo FC, jen by měly rozumět sadě příkazů SCSI (což není takový problém), případně NVMe. Jde o to, aby přenášená data nemusela jít přes vyšší vrstvy síťového zásobníku, zlepšuje se tím latence komunikace.

 Fibre Channel může používat některou z těchto topologií:

- point-to-point – na jedné straně je datové úložiště, na druhé straně server s HBA kartou,
- přepínaná FC síť (switched fabric) – používáme DCE, což jsou switche s podporou FC,
- FC over Ethernet (FCoE) – používáme běžné switche, které „neumí“ protokol FCP.

FCoE (Fibre Channel over Ethernet) je pokus o sloučení jinak rozdílných sítí Fibre Channel a optický Ethernet. FC rámce se v tomto případě ještě musejí zapouzdřit do ethernetových rámců, které bez problémů projdou přes ethernetové switche. Je to levnější alternativa, protože FC switche bývají většinou dražší.

3.7.2 iSCSI

 iSCSI (Internet Small Computer System Interface) je obdobou FCoE určenou spíše pro menší SAN síť. Jedná se vlastně o možnost posílání SCSI příkazů přes IP síť. Klient (iniciator) potřebuje přes síť přistupovat k serveru (disku), z pohledu klienta je funkčnost stejná jako kdyby šlo o interní disk v počítači (typu SCSI, SAS apod.), jen nízkourovňové příkazy jsou posílány na síť. Na serverové straně jsou tyto příkazy rozbaleny z paketu a předány příslušnému úložišti.

Zatímco FC pracuje na vrstvách L1 a L2 ISO/OSI, iSCSI pracuje až na síťové vrstvě.

 Technologie iSCSI může být implementována softwarově nebo hardwarově. Softwarová implementace je levnější (konkrétně – zdarma), příslušný software si lze stáhnout od výrobce/distributora každého operačního systému nebo je dokonce součástí instalace systému (jen je obvykle nutné aktivovat příslušnou službu či démona).

Hardwarová implementace znamená podporu na NIC (za tu se připlácí). Výhodou hardwarové implementace je vyšší rychlost, protože pomocný procesor na NIC odlehčuje hlavním procesorům při zpracovávání (nejen) iSCSI požadavků. Obecně platí, že do menší SAN stačí softwarová implementace, do větší nebo hodně vytížené SAN je dobré investovat do síťových karet s hardwarovou implementací iSCSI (samozřejmě pokud máme SCSI disky).



Další informace:

- O SAN sítích (Fibre Channel a iSCSI) najdeme pěkný článek například na <https://www.samuraj-cz.com/clanek/storage-technologie-a-san-site-aneb-pripojeni-serveru-k-diskovemu-poli/>
- Článek o FCoE (první díl seriálu): <http://www.abclinuxu.cz/clanky/hardware/fcoe-fibre-channel-over-ethernet>
- Článek o iSCSI: <https://www.samuraj-cz.com/clanek/iscsi-san-sit-a-konfigurace-na-windows-server-2012/>



Další možná řešení SAN sítí jsou InfiniBand, ATAoE (ATA over Ethernet), apod.



Souhrn kapitoly

V této kapitole jsme se zabývali různými technologiemi doplňujícími či rozšiřujícími možnosti lokálních sítí, především virtuálními lokálními sítěmi, protokolem STP, sdružováním linek do logických spojů s vyšší propustností a sítěmi optimalizovanými pro zpřístupňování dat v datových centrech.



Rozlehlé sítě a telekomunikace

 **Rychlý náhled:** WAN (Wide Area Network) je rozlehlá síť zabírající obvykle rozlohu státu, kontinentu či celého světa. Propojuje nikoliv koncová zařízení, ale různé sítě typu LAN, MAN či menší WAN sítě. V této kapitole se budeme zabývat různými tématy souvisejícími s WAN sítěmi: spojovými protokoly, architekturou WAN sítí, komunikací v těchto sítích, a také různými přístupovými sítěmi. Na konci kapitoly se krátce budeme věnovat pobočkovým ústřednám.

 **Klíčová slova:** WAN, DCE, DTE, ABM mode, i-frame, s-frame, u-frame, X.25, Frame Relay, ATM, MPLS, label, FEC, protokol LDP, SONET/SDH, WDM síť, DWDM, PPP, telekomunikační síť, first-mile/last-mile síť nebo přístupová síť, xDSL, DOCSIS, FTTx, pobočková ústředna.

 **Cíle studia:** Cílem této kapitoly je porozumět principům WAN sítí – jejich architektuře, propojení, přenosům, souvisejícím protokolům, a orientovat se v problematice optické infrastruktury, protože se ta se u WAN sítí typicky používá. Cílem je také pochopit různé způsoby přístupu do těchto sítí.

4.1 WAN sítě

4.1.1 Struktura WAN sítí

-  Jakou topologii vlastně používají WAN sítě? Záleží, jaký pohled použijeme.
- WAN sítě jsou navzájem uspořádány do (převážně) hierarchické struktury velmi podobné hierarchii podle protokolu IP. Souvisí to s tím, že lokální a regionální poskytovatelé konektivity tuto konektivitu také odněkud musejí získat, od velkých nadnárodních poskytovatelů. Celková struktura má charakter *páteřní sítě*.
 - Když se podíváme na vnitřní strukturu WAN sítě, obvykle zjistíme, že používá architekturu *mesh*: redundantní spoje, které zajišťují vysokou dostupnost, odolnost proti výpadku.
-  Zatímco u lokálních sítí víceméně platí, že data po těchto sítích posílaná souvisejí s vlastníkem sítě, u MAN a zejména WAN sítí tomu tak není. Provozovatelé WAN sítí jsou prostě majiteli síťové infrastruktury, kterou pronajímají svým zákazníkům, je to páteř, na kterou jsou napojeny zákaznické sítě.

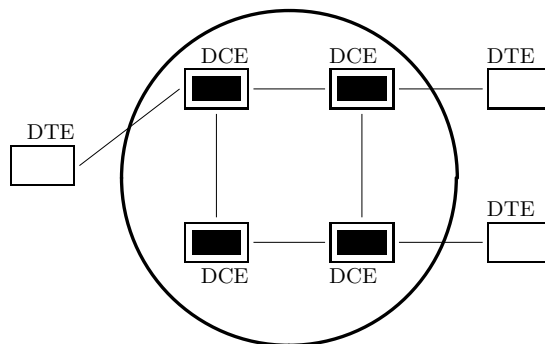
Síťová konektivita (případně v určité definované kvalitě) je jedinou službou, kterou WAN poskytují, funkčnost je podřízena také tomu, aby bylo možné poskytnuté služby tarifkovat. Dokonce může být

problém se směrováním na jiné než unicast adresy, protože pro WAN je typická spojovaná komunikace přes virtuální okruhy.

 Ve WAN sítích obvykle rozlišujeme tyto typy zařízení:

- *DCE (Data Circuit Equipment)* jsou zařízení uvnitř WAN sítě (jádro). Obvykle přímo nekomunikují s ničím, co by do WAN nepatřilo, prostě dokážou přes sebe navázat komunikační okruh (circuit) a pak přes něj přenášet data. Pro zákazníka jsou „neviditelné“. Jejich typickou vlastností je rychlost přepínání dat (obvykle s hardwarovou podporou), je to obdoba switchů v LAN sítích. DCE se nezdržují prohlížením přenášených PDU, zajímá je jen „lokalizační část“ záhlaví.
- *DTE (Data Terminal Equipment)* jsou zařízení na rozhraní mezi WAN sítí a tím, co je k ní připojeno, jsou to tedy hraniční zařízení. Potřebují „rozumět“ jak protokolům pracujícím ve WAN, tak i protokolům z připojené sítě (vzhledem k připojené síti obvykle potřebují funkcionalitu vrstvy L3, IP). Pracují jako překladatelé mezi dvěma druhy sítí.

Některé WAN sítě rozlišují i další typy svých zařízení, ale nám budou tyto pojmy stačit.



Obrázek 4.1: Zařízení v síti Frame Relay



Poznámka:

V některých WAN se používá trochu jiná terminologie:

- zařízení P (Provider) odpovídají zařízením DCE, jsou tedy uvnitř WAN sítě,
- zařízení PE (Provider Edge – na hranici poskytovatele) jsou na hranicích WAN sítě,
- zařízení CE (Customer Edge – na hranici zákazníka) jsou na straně zákazníka a zprostředkovávají pro něj přístup k WAN síti.

Zařízení typu PE a CE tedy mohou dohromady tvořit DTE.



4.1.2 Protokoly vrstvy L2 pro WAN sítě

Pro protokoly pracující na vrstvě L2 se vžil název *spojové protokoly*. Kromě těch, které byly zmíněny dříve (například LLC), zde najdeme také protokoly pro WAN sítě (protože většina WAN sítí je implementována především na této vrstvě).

Obvykle platí, že každá WAN síť má „svůj“ spojový protokol, nicméně u těchto protokolů můžeme vysledovat určité společné vlastnosti (protože ve skutečnosti jsou mezi nimi vztahy předek-potomek).

Předkem všech novějších spojových protokolů je protokol HDLC (High Data Link Control). Jeho potomkem je LAPB pro síť X.25, další generace je LAPD pro ISDN, dále LAPF pro Frame Relay, a další. Do rodiny patří také protokol PPP a jeho odvozeniny.

 **Uzly v síti a přenosové režimy.** Spojové protokoly obvykle rozlišují primární uzel (začíná komunikaci, určuje parametry spojení) a sekundární uzly (pouze odpovídají). Pokud se jedná o point-to-point komunikaci, pak máme jeden primární a jeden sekundární uzel, když je to point-to-multipoint komunikace, je třeba jeden z uzlů prohlásit za „nejdůležitější“ – primární.

Uzly mezi sebou komunikují v konkrétním režimu:

- *normální režim odpovědi* (NRM, normal response mode) – sekundární uzel musí počkat na vyzvání od primárního, aby mohl komunikovat,
- *asynchronní režim odpovědi* (ARM, asynchronous response mode) – sekundární uzel může začít komunikovat s primárním bez povolení,
- *asynchronní vyvážený režim* (ABM, asynchronous balanced mode) – každý uzel může začít komunikaci, ten se pro tuto komunikaci stává primárním uzlem.

Novější spojové protokoly typicky fungují v režimu ABM.

 **Rámec spojového protokolu.** Rozlišujeme několik typů rámců:

- *I-frame* (Information Frame, informační, datový) – nese informace z vyšší vrstvy a případně řídicí informace, tyto rámce podporují řazení, řízení toku, detekci chyb, zotavení. Určení: přenos dat.
- *S-frame* (Supervisory Frame, také služební či dohlížecí rámec) – obsahuje řídicí informace, je používán pro požadavek zahájení nebo ukončení spojení, hlášení o stavu, potvrzení přijetí I-rámce.
- *U-frame* (Unnumbered Frame, nečíslovaný) – pro řídicí informace (dohoda režimu provozu), na rozdíl od předchozího nepodporují číslování rámců do posloupnosti, a může také obsahovat data z vyšší vrstvy. Obecně: pro to, co se vejde do jediného rámce.

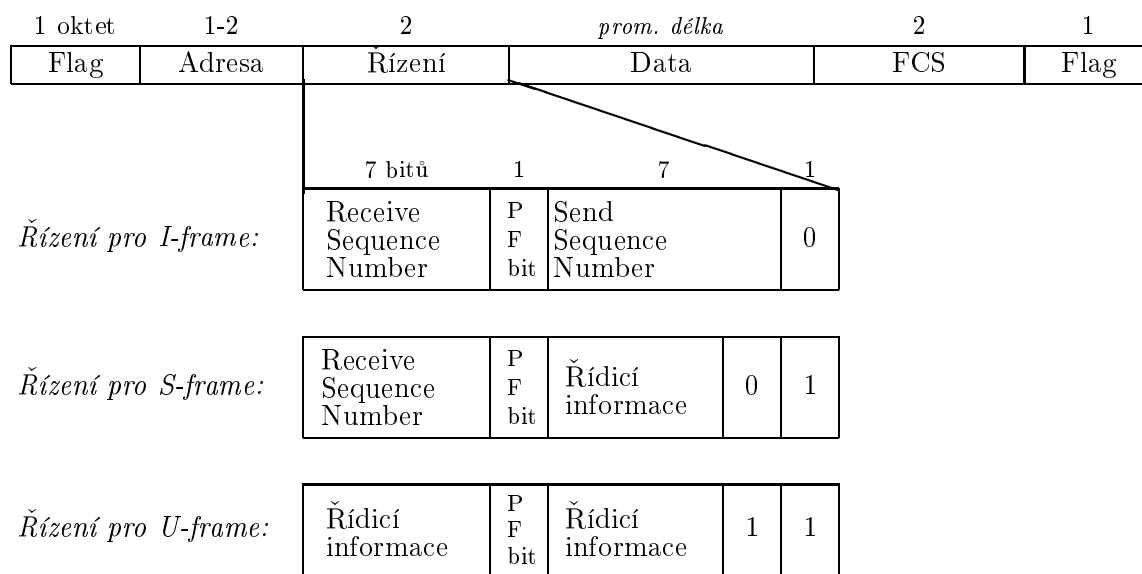
Rámce typu I-frame používají čísla Sequence Number pro oba směry (podobný mechanismus jako u TCP), S-frame mají čísla pouze pro jeden směr (přenášený obsah má totiž smysl pouze pro jeden ze směrů komunikace), U-frame nepoužívá žádná taková čísla (je nečíslovaný, nedělí komunikaci na části, které by bylo třeba kompletovat).

Čím méně místa zabírají čísla Sequence Number, tím více místa zbývá na speciální významové bity (typ rámce, určení konkrétního požadavku nebo odpovědi na požadavek, atd.).

 **Obecná (zjednodušená) struktura WAN rámce je tato (viz obrázek 4.2):**

- flag (příznak) na začátku rámce – synchronizační informace ve tvaru 01111110,
- adresa – formát podle toho, jak se v konkrétní WAN adresují okruhy,
- pole pro řízení – Sequence Numbers a speciální významové bity podle typu rámce,
- data (pokud jsou v tomto rámci přenášena),
- FCS – kontrolní součet,
- flag stejně jako na začátku rámce.

 V řídicím poli je také P/F bit (poll/final). Primární uzel sděluje sekundárnímu nastavením tohoto bitu na 1, že vyžaduje okamžitou odpověď (tj. posílá se dotaz), v opačném směru sekundární uzel primárnímu, že tento rámec je poslední ze zasílaného sledu rámců.



Obrázek 4.2: Typický formát WAN rámce, zde pro protokol HDLC

**Úkol**

Rámce HDLC mají ve skutečnosti několik variant. Jednou z nich je Cisco HDLC, který Cisco i v současnosti používá na sériových linkách. Podívejte se na <https://www.freecnastudyguide.com/study-guides/ccna/ch11/11-3-point-point-wans-layer-2/> a srovnajte Cisco HDLC s původním formátem HDLC.



Protokol IEEE 802.2 (LLC) známe už z kapitoly o Ethernetu, také se jedná o spojový protokol, takže se k němu krátce vrátíme i zde. Je určen pro point-to-point spoje, pracuje v asynchronním vyváženém režimu (ABM). Opět se rozlišují tři typy rámců – I-frame, S-frame a U-frame. Formát rámce je celkově jednodušší, protože LLC rámec se vždy zapouzdřuje do některého MAC rámce. Používá se jiná adresace (přístupové body služeb).

LLC nabízí tři typy služeb:

- *nepotvrzovaná služba bez spojení* (typ 1) – nejpoužívanější, funguje podobně jako doručování datagramů (paket je vybaven všemi dostupnými informacemi včetně adres zdroje a cíle a pořadí paketu ve zprávě), žádné ošetření chyb, v tomto ohledu spoléhá na protokoly vyšších vrstev,
- *služba se spojením* (typ 2) – pro virtuální okruhy, zahrnuje navázání spojení, potvrzování po doručení omezené skupiny paketů, atd., pouze první paket obsahuje informace o adresách,
- *potvrzovaná služba bez spojení* (typ 3) – pakety jsou po skupinách potvrzovány, tato služba je jen málo využívána.

Koncové stanice podporují vždy nejméně typ 1, a pak obvykle ještě jednu ze zbývajících služeb.

4.2 Nejznámější WAN sítě

Předkové WAN sítě jsou protokoly SDLC a HDLC používané pro vzdálený přístup k mainframům (SDLC) a LAN sítím (HDLC), například přes telefonní linky.

4.2.1 X.25 – přepínání paketů

Nejstarší skutečnou WAN sítí je X.25 standardizovaná ITU (později byl vytvořen i ISO standard). Tato síť byla typická extrémní spolehlivostí (protože byla projektována pro velmi nespolehlivé linky) a nízkou rychlostí. Implementovala přibližně vrstvy L1–L3 a pracovala na principu přepínání paketů.

X.25 používala spojový protokol LAPB (Link Access Procedure Ballanced), který pracoval v módu ABM na point-to-point linkách.

Síť X.25 byla charakteristická velmi propracovanými řídicími mechanismy pro zajištění spolehlivosti, včetně potvrzování, aby vůbec bylo možné ji provozovat na nespolehlivých linkách. Protože se spolehlivost linek postupně zlepšovala, stala se tato technologie přežitkem a byla nahrazena jinými, méně spolehlivými, ale zato rychlejšími.

4.2.2 Frame Relay – přepínání rámců

Specifikace Frame Relay původně vznikla v rámci specifikace ISDN, ale protože se ukázala dobrá životaschopnost tohoto řešení, v r. 1989 se osamostatnila.

Typickým použitím Frame Relay bylo z principu propojování lokálních sítí, implementace páteřní sítě. Počítá se s přenosem po kvalitním a rychlém vedení, tedy Frame Relay poskytuje sice službu se spojením, ale nespolehlivou (poskytuje detekci chyb, ale bez možnosti opravy).

 V síti Frame Relay se nacházejí tato zařízení:

- *DTE* –vlastníkem může být zákazník nebo jde o pronájem, jsou to v podstatě hraniční zařízení.
- *DCE* (Data circuit-terminating equipment) – poskytují služby synchronizace a přepínání provozu (obvykle FR switche), vnitřní zařízení sítě.

U obou existuje komponenta fyzické vrstvy a komponenta linkové (spojové) vrstvy. Na fyzické jde většinou o některé sériové rozhraní, na linkové pracuje protokol LAPF.

FRAD (Frame Relay Access Device, Frame Relay Assembler/Disassembler) je zařízení, které slouží k přístupu do sítě Frame Relay. Může to být buď samostatné zařízení, anebo je to součást směrovače, přes který je lokální síť připojena do Frame Relay sítě. *FRAD* je tedy představitel *DTE* nebo je součástí *DTE*.

Komunikuje se po virtuálních okruzích, obvykle přepínaných (*SVC*). Pro *SVC* se používají operace navázání spojení (vytvoření okruhu), přenos dat, *Idle* a ukončení spojení (zrušení okruhu). U *PVC* se používá jen operace přenosu dat a *Idle*.

 **Garance informačního toku.** *Parametr CIR* (Committed Information Rate, zaručený informační tok) určuje propustnost, kterou má spoj zaručovat. Hodnota *CIR* je pro okruhy *PVC* dohodnuta mezi poskytovatelem (telekomunikační společností) a uživatelem (vlastníkem lokální sítě, která má být připojena), pro okruhy *SVC* je dohodnuta při navazování spojení. Rámce posílané nad hodnotu *CIR* mohou být při větším zatížení sítě zahazovány.

Pro uživatele má hodnota *CIR* tento význam:

- čím vyšší *CIR*, tím je služba dražší,
- čím menší *CIR*, tím větší pravděpodobnost zahazování paketů a potenciálně pomalejší přenos.

Nárazový informační tok (také *EIR*, Excess Information Rate) představuje maximální informační tok, který dokáže přenosová cesta zvládnout (přesněji to, co je navíc nad *CIR*).

 Rámce zasílané nad hodnotu CIR, které nepřekračují hranici EIR, jsou označeny jako *vhodné k zahození* (DE, Discard-Eligible). Jsou předávány přepínači Frame Relay (tj. do FR sítě) tak dlouho, dokud není k dispozici dostatečná šířka pásma (dokud se nevejdou do CIR). Označení rámce jako vhodného k zahození obvykle neznámá, že rámec nebude doručen, ale že se jeho doručení může (i dost hodně) zdržet.

 *SLA* (Service Level Agreement) = dohoda o úrovni poskytovaných služeb (tento pojem se používá obecně u tohoto typu služeb, nejen pro Frame Relay). Jedná se o dohodu o parametrech služby, zde také zahrnuje CIR a EIR.

 **Adresace a přepínací tabulky.** Po navázání okruhu se používá lokální adresace pomocí *čísel DLCI* (Data Link Connection Identifier). Ve skutečnosti nejde o identifikátory zařízení, jsou to identifikátory okruhů. Hodnoty DLCI pro DTE v naší LAN získáme od poskytovatele služby Frame Relay (typicky telekomunikační společnost), samozřejmě pokud si Frame Relay implementujeme sami, tak si DLCI taky sami přidělíme.

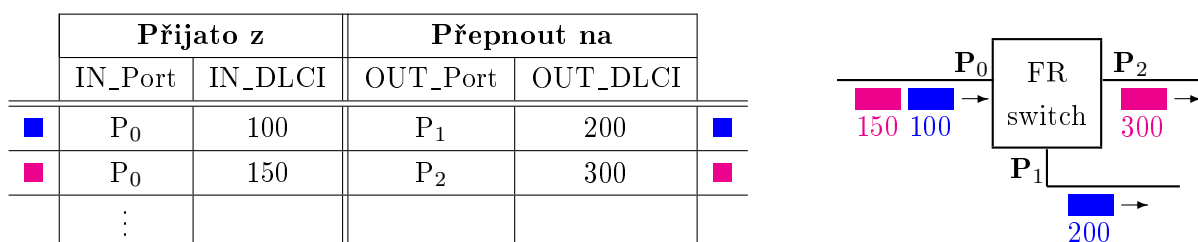
V celé síti WAN (propojující různé lokální sítě) však může existovat více okruhů se stejným DLCI, navíc se hodnota DLCI při průchodu přes přepínače mění, jak dále uvidíme.

Některé hodnoty DLCI jsou vyhrazeny pro účely signalizace stavů, problémů apod., především z rozsahu 0–16. Dále rozmezí 1019–1022 je určeno pro skupinové vysílání.

 Zákazník potřebuje propojit své pobočky. Mezi nimi povede virtuální okruh (obvykle SVC). Poskytovatel FR služby musí být schopen vytvořit okruhy tak, aby se neprolínaly cesty různých zákazníků (třebaže části těchto cest vedou přes sdílenou infrastrukturu, tytéž kabely), také musí být odlišeny rámce posílané do různých poboček téhož zákazníka.

Takže na začátku cesty (na hranici FR sítě) je paket přicházející z LAN sítě zákazníka na hraničním DTE zapouzdřen do FR rámce s takovou hodnotou DLCI, která sdělí sousedovi (což už bude vnitřní uzel DCE ve FR síti), přes který okruh má tento rámec jít. Soused má ve své přepínací tabulce informaci o tom, jak s takovým rámcem zacházet: číslo DLCI v záhlaví může být změněno na jiné, které dalšímu DCE na cestě určí pokračování okruhu, také je zde informace o portu, na který má být rámec odeslán.

 Na jednotlivých zařízeních jsou rámce přepínány podle *přepínacích tabulek*. Uvnitř FR sítě tedy najdeme FR switche, jejich tabulka je jednoduchá, jak vidíme na obrázku 4.3 (případně by byl ještě další sloupec – informace o tom, zda je daná cesta aktivní, apod.). Přepínání je implementováno v hardwaru.



Obrázek 4.3: Ukázka jednoduché přepínací tabulky Frame Relay switche

Na hranici sítě jsou FR routery (v roli DTE), které na jedné straně komunikují se zařízeními lokální sítě zákazníka, na druhé straně komunikují s FR switchi. Tato zařízení obsahují dvě tabulky:

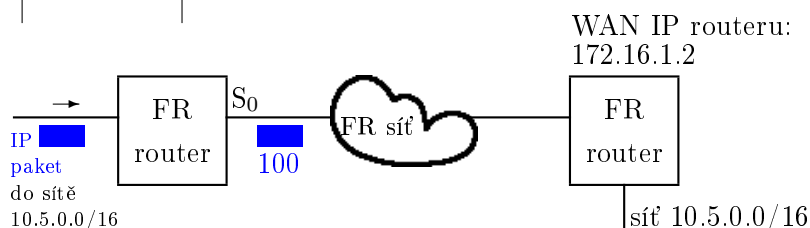
- směrovací tabulku určující, kam se rámec směřující do sítě s danou IP adresou má poslat,
- tabulku mapování DLCI.

Routing Table:

Network	Next Router	Interface
10.5.0.0/16	172.16.1.2	S ₀
10.27.0.0/16	172.16.1.8	S ₁
⋮		

FR Map:

Next router	DLCI
172.16.1.2	100
172.16.1.8	200
⋮	



Obrázek 4.4: Směrování na FR routeru (DTE), paket přichází do FR sítě

Obsah a význam obou tabulek vidíme na obrázku 4.4. Na třetí vrstvě je paket nasměrován na daný DCE (na druhé straně FR sítě) a příslušné rozhraní, na druhé vrstvě je nastaveno DLCI.

Úkol

Podívejte se na oba předchozí obrázky a zjistěte, co se stane, když paket s cílovou adresou 10.5.24.7 přijde na router, jehož tabulky jsou na obrázku 4.4. Které číslo DLCI bude přiděleno pro první část cesty přes FR síť? Jak bude tento rámec zpracován na switchi, jehož přepínací tabulka je na obrázku 4.3?

 **Linková vrstva.** Na linkové vrstvě pracuje *protokol LAPF* (Link Access Procedure – Frame). Do (datového) rámce protokolu LAPF se zapouzdřuje IP paket.

Adresové pole má 2 oktety (může být rozšířeno až na 4) a zcela změněnou strukturu, řídicí pole bylo vyřazeno. V adresovém poli je uloženo číslo *DLCI*. V řídicím poli najdeme především bity pro řízení toku dat a kontrolu chyb:

- *C/R* – command/response bit, určuje, zda jde o příkaz (command) nebo odpověď či reakce na dříve zasláný příkaz (response),
- *bit FECN* (Forward-explicit congestion notification, dopředné oznámení o přetížení) – DCE na cestě oznamuje, že tento rámec byl poslán po přetíženém spojení, a tedy DTE, které je cílem rámců komunikace, by mělo snížit frekvenci příjmu rámců (pokud tak dokáže vyšší vrstva reagovat),
- *bit BECN* (Backward-explicit congestion notification, zpětné oznámení) – oznámení zdroji vysílání o přetížení linky, zdrojové zařízení by mělo snížit objem vysílání,
- *bit DE* (Discard-Eligible) – v rámcích, které mohou být při přetížení zahozeny.

Když DCE zjistí přetížení, nastaví v posílaném rámcu bit FECN. Takto přijímající DTE zjistí, že na lince došlo k zahlcení spoje. Naopak když DCE přepíná rámec s nastaveným FECN, v nejbližším rámcu posílaném v okruhu v opačném směru nastaví bit BECN. Takto vysílající DTE zjistí, že na lince došlo k zahlcení spoje.

 **LMI** (Local Management Interface) je rozšíření specifikace Frame Relay, které určuje komunikaci mezi zařízeními DTE a DCE, tedy určuje, jakým způsobem se má k FR přistupovat z lokální sítě.

**Další informace:**

- <http://www.cs.vsb.cz/grygarek/TPS/FrameRelay/frame.html>
- <http://books.google.cz/books?id=tROBDX-OBREc&printsec=frontcover>

**4.2.3 ATM – přepínání buněk**

Sítě ATM (Asynchronous Transfer Mode) byly původně považovány za velmi nadějnou technologii, nicméně v praxi byly postupně vytlačeny jednoduššími a pružnějšími sítěmi MPLS. „Asynchronní“ v názvu neznamena asynchronní přenos, ale nepravidelné zasílání „plných“ buněk na cestě, buňky jsou obsazovány daty podle potřeby, ne podle algoritmu.

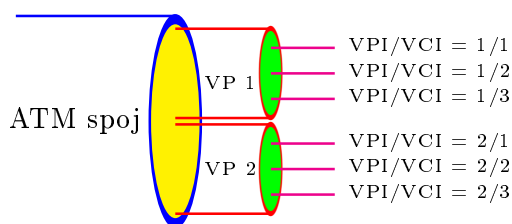
ATM pracuje na principu komunikace se spojením, nespolehlivá služba (tj. bez oznámení chyb), na plném duplexu, v statistickém multiplexu. Důležitou vlastností je garance kvality služeb pro různé typy dat – vlastně jednou z nejdůležitějších charakteristik ATM jsou právě široké možnosti řízení kvality služeb.

**Poznámka:**

Pozor na záměnu zkratk: ATM je nejen WAN technologie, ale také se jedná o zkratku používanou pro bankomaty – v americké angličtině se jim říká Automated Teller Machine (ATM). Nicméně není vyloučeno, že ATM zařízení by mohla být připojena přes ATM síť. . .



Adresace okruhů je na rozdíl od FR dvouhodnotová. *Virtuální cesta* a *virtuální kanál* jsou obdobou DLCI. V jedné fyzické přenosové cestě může vést více virtuálních cest, v jedné virtuální cestě vede více virtuálních kanálů. Číslo virtuální cesty označujeme VPI (Virtual Path Identifier), číslo virtuálního kanálu označujeme VCI (Virtual Channel Identifier). *Hodnota VPI/VCI* plně identifikuje přenosovou cestu mezi dvěma sousedními uzly, na každém ATM přepínači se mění.



Obrázek 4.5: Virtuální cesty a virtuální kanály v ATM (zjednodušeně)

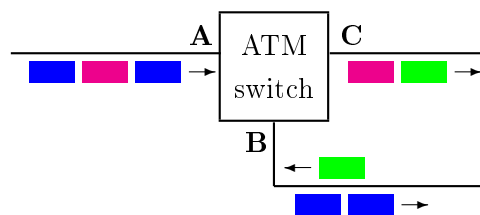
ATM switch si vede *tabulku připojení*, ve které jsou přidruženy příchozí a odchozí VPI/VCI, záznam se tvoří během navázání spojení (vytvoření okruhu).

Na obrázku 4.6 je zjednodušená a zkrácená přepínací tabulka ATM switchu. Jsou v ní tři virtuální cesty. Každý řádek obsahuje jednu přepínací informaci, například v prvním řádku zjistíme, že buňka přicházející z portu A po cestě 1 v kanálu 29 je přepnuta na port B cestu 2 kanál 42. Porty jsou označeny písmeny spíše pro snadnější odlišení, většinou se setkáme s označením čísly nebo řetězci, podle výrobce.



PDU se nazývají *buňky*, jsou velmi malé s konstantní délkou 53 oktetů (z důvodu jednoduchého a rychlého přepínání). Aby se do buněk vůbec vešla nějaká data, je záhlaví velmi malé, pouze 5 B, na

	Přijato z			Přepnout na			
	Port	VPI	VCI	Port	VPI	VCI	
■	A	1	29	B	2	42	■
■	A	3	42	C	1	8	■
■	B	1	18	C	2	36	■
	⋮						



Obrázek 4.6: Ukázka jednoduché přepínací tabulky ATM switche

data zbývá 48 oktetů. V záhlaví máme kromě adres VPI/VCI také například jednobitový příznak CLT (Cell Loss Priority, priorita ztráty buňky, 1 bit), který plní stejnou roli jako u FR bit DE.

Protože se komunikuje v statistickém multiplexu, jsou v přenosu definovány sloty, do kterých se „sázejí“ buňky. Zde je právě výhodou, že buňky mají konstantní velikost odpovídající velikosti slotů.



Poznámka:

O ATM bychom mohli ještě dlouho pokračovat (zejména architektura této sítě je velmi složitá – kromě vrstev se zde totiž vyskytují i roviny a několik různých druhů přizpůsobovacích vrstev k napojení na jiné protokoly), ale vzhledem k tomu, že dnes se s ATM nesetkáme, by to nemělo smysl.



4.2.4 MPLS – přepínání značek

Dosud jsme se zabývali přepínáním paketů (X.25, no vlastně tuto síť jsme přeskočili, už se nikde nepoužívá), rámců (Frame Relay) a buněk (ATM), teď se podíváme na přepínání značek.



Technologie byla představena firmou Ipsilon Networks pod názvem *IP Switching*. Později společnost Cisco vytvořila proprietární standard *Tag Switching*, který síť zbavil závislosti na ATM, podobný, ale otevřený standard později vydalo sdružení IETF.



MPLS (MultiProtocol Label Switching) je síť založená na *přepínání značek* (label, také návěští). Účelem je přesunout co nejvíce režie směrování (včetně administrace, QoS apod.) na okraj sítě tak, aby vnitřní oblast sítě byla co nejrychlejší. MPLS dokáže velmi rychle přenášet nejen běžná data, ale také hlas a video, a to se zajištěním QoS. Další výhodou je snadnější implementace virtuálních sítí.

Obrovskou výhodou MPLS je mnohotvárnost. Nemá přímo definovanu adresaci ani směrování, dokáže spolupracovat s více odlišnými protokoly. Původně MPLS pracovala jako nástavba nad ATM, ale v současné době pracuje také nad Ethernetem, Frame Relay, SONET/SDH, DWDM a dalšími. Takže výhody MPLS můžeme shrnout takto:

- rychlost přepínání,
- zajišťování řízení provozu (vyvažování zátěže) a QoS (odlišné zacházení s různými PDU),
- podpora VPN,
- schopnost spolupráce s mnoha různými protokoly, pod MPLS mohou fungovat různé technologie.

V ISO/OSI modelu bychom mohli MPLS zařadit někam mezi druhou (spojovou) a třetí (síťovou) vrstvou, také se označuje jako vrstva 2.5 nebo 2+. Je to proto, že MPLS přidává přes L3 záhlaví (obvykle IP) svoje záhlaví, a následně je to celé zapouzdřeno do L2 záhlaví některého protokolu, nad kterým MPLS pracuje.

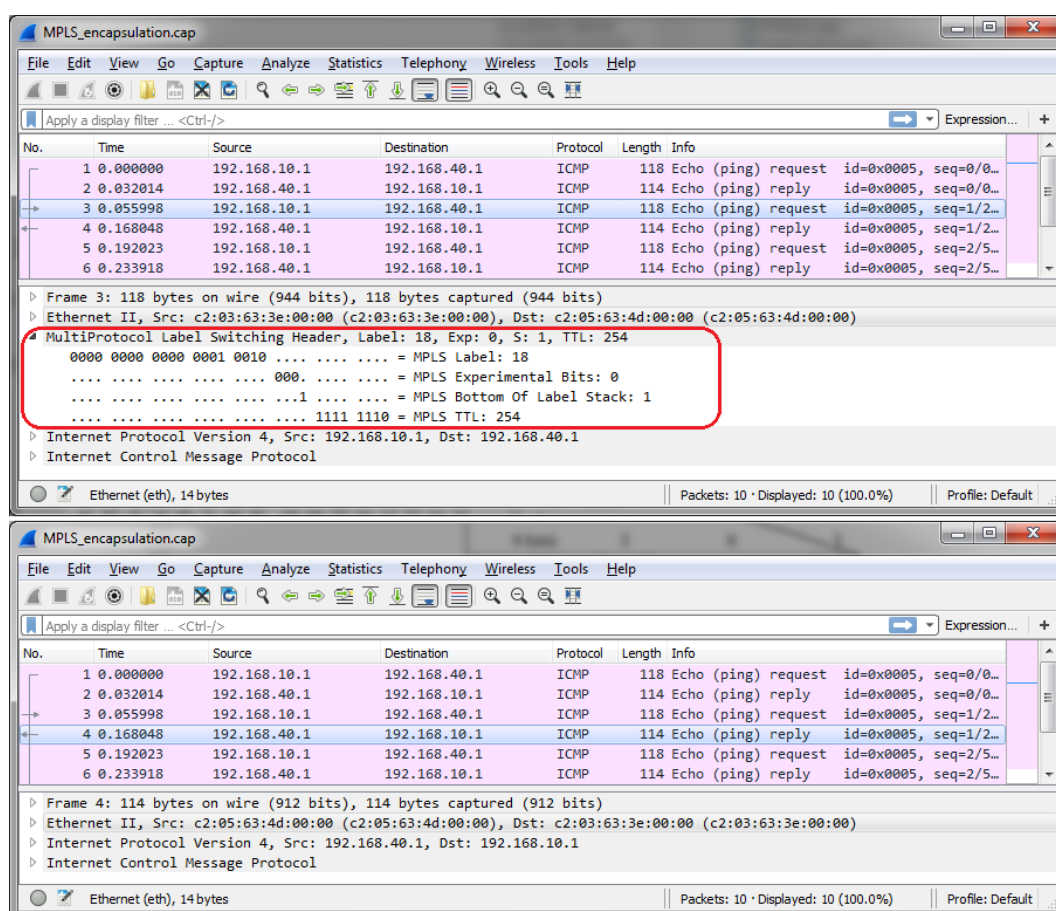
 Každý paket, který vstupuje do MPLS sítě, je na okrajovém směrovači opatřen jedním nebo více *MPLS záhlavími* uspořádanými do *zásobníku* (stack), a to na rozhraní mezi záhlavími druhé a třetí vrstvy. Záhlaví protokolu MPLS má tuto strukturu:

- samotná značka (návěští, label, 20 bitů),
- QoS informace (3 bity), v případě MPLS se setkáme s názvem CoS (Category nebo Class of Service) nebo také ve významu Experimental (Exp.), do tohoto pole se případně mapuje QoS informace ze zapouzdřeného paketu,
- příznak konce zásobníku (1 bit, end of stack, bottom of label stack); pokud nenásleduje další záhlaví, je nastaven na 1 (to znamená, že následuje už přímo záhlaví IP paketu),
- TTL (Time to Live, 8 bitů) platné pro síť MPLS, může se mapovat z IP paketu.

Značka uložená v MPLS záhlaví jednoznačně určuje směrování paketu na cestě k následujícímu směrovači, v podobném smyslu jako DLCI u FR nebo VPI/VCI u ATM.

Příklad

Na obrázku 4.7 nahoře vidíme zprávu ICMP Echo request zapouzdřenou do IPv4 paketu, ten je uvnitř MPLS rámce, který je zapouzdřen v ethernetovém rámci.



Obrázek 4.7: Nahoře ping request posílaný přes MPLS spoj, dole odpověď došla po běžném ethernetovém spoji¹

¹Obrázek podle rámce v cap souboru z packetlife.net

Tento MPLS rámec má v zásobníku pouze jedno záhlaví (může mít více), label je v něm nastaven na 18, QoS (experimental) bity nejsou použity (jsou nastaveny na 0), je nastaven bit označující poslední záhlaví v zásobníku (bottom of label stack) a TTL je 254.



Záhlaví tedy může být v rámci více. *Vnější záhlaví* na zásobníku je určeno právě ke stanovení cest. Nenabízí v podstatě o moc víc než IP záhlaví (s jedním důležitým rozdílem – rychleji se zpracovává). Další záhlaví hlouběji v zásobníku mají trochu jiný, specifický, účel, například záhlaví VPN pro určení virtuální sítě, záhlaví pro QoS nebo záhlaví pro řízení provozu.

Záhlaví rámce 2. vrstvy	MPLS Label 1, ES bit = 0	MPLS Label 2, ES bit = 0	MPLS Label 3, ES bit = 1	Záhlaví paketu 3. vrstvy (např. IP)	Datová část paketu
-------------------------------	--------------------------------	--------------------------------	--------------------------------	--	--------------------

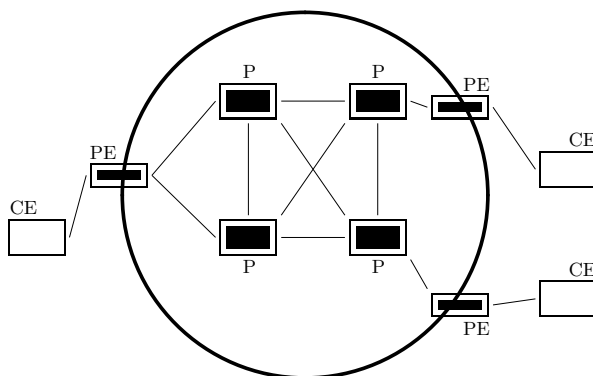
Tabulka 4.1: Zásobník značek v MPLS

Přidělování značek (ve všech záhlavích ve stacku) probíhá formou třídění datových jednotek do *tříd* (FEC – Forward Equivalence Class), datové jednotky patřící do téže třídy jsou z uzlu odesílány se stejným záhlavím. Třída je stanovena na základě několika kritérií – především podle prefixu IP adresy (VLAN) a dále například podle některých charakteristik VPN.

Směrování v MPLS. Směrovače (spíše přepínače) v síti MPLS jsou nazývány *LSR* (*Label Switching Router*). LSR na okraji sítě (tj. komunikující také s uzly v lokálních sítích, které jsou sítí MPLS propojeny), jsou označovány jako *ELSR* (*Edge LSR*, hranové směrovače), a jsou buď vstupní (*Ingress ELSR*) nebo výstupní (*Egress ELSR*) podle směru provozu. U některých výrobců se můžeme setkat s odlišnou terminologií – LSR jsou nazývány provider/core router, ELSR se nazývají provider edge router (ale v mnoha dokumentech se používá „původní“ terminologie).

Velice často se také setkáváme s trochu jiným označením:

- P (Provider) – přepínače uvnitř MPLS sítě, ve vlastnictví poskytovatele (nebo jeho poskytovatele),
- PE (Provider Edge) – na hranici MPLS sítě,
- CE (Customer Edge) – na hranici sítě zákazníka, buď v jeho vlastnictví nebo pronajat od poskytovatele (komunikuje s PE).



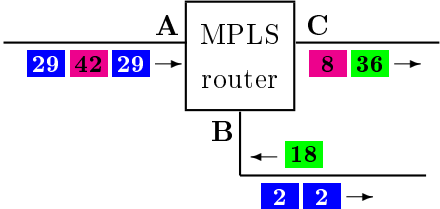
Obrázek 4.8: Struktura MPLS sítě

Na obrázku 4.8 je naznačena obvyklá struktura MPLS sítě s uzly P, PE a CE. CE je součástí lokální sítě zákazníka, k jednomu PE může být připojeno více CE (tj. více LAN).

 Směrování (přepínání značek) probíhá podobně jako v ATM, ale ještě jednodušeji. Směrovače si vedou jednoduchou *tabulku* (také se nazývá Label Forwarding Database – LFD), která je podobná tabulce pro ATM, ale místo dvojice VPI/VCI je zde pouze hodnota značky (datagram z portu PPP opatřený značkou XXX je opatřen značkou YYY a poslán na port QQQ, apod.). Jedná se jen o určení vazby mezi příchozí a odchozí značkou, cesty jsou jednosměrné.

Na obrázku 4.9 je ukázka tabulky na LSR uvnitř sítě. Pokud by se jednalo o vstupní Ingress ELSR, hodnoty ze sloupce se vstupní značkou a portem by nebyly používány, obdobně u Egress ELSR by nebyly používány hodnoty ze sloupce s výstupní značkou a portem.

	In port	In label	Prefix adresy	Out label	Out port	
	A	29	128.95.0.0/16	2	B	
	A	42	128.101.0.0/16	8	C	
	B	18	141.62.0.0/16	36	C	
	⋮					



Obrázek 4.9: Ukázka jednoduché přepínací tabulky MPLS směrovače

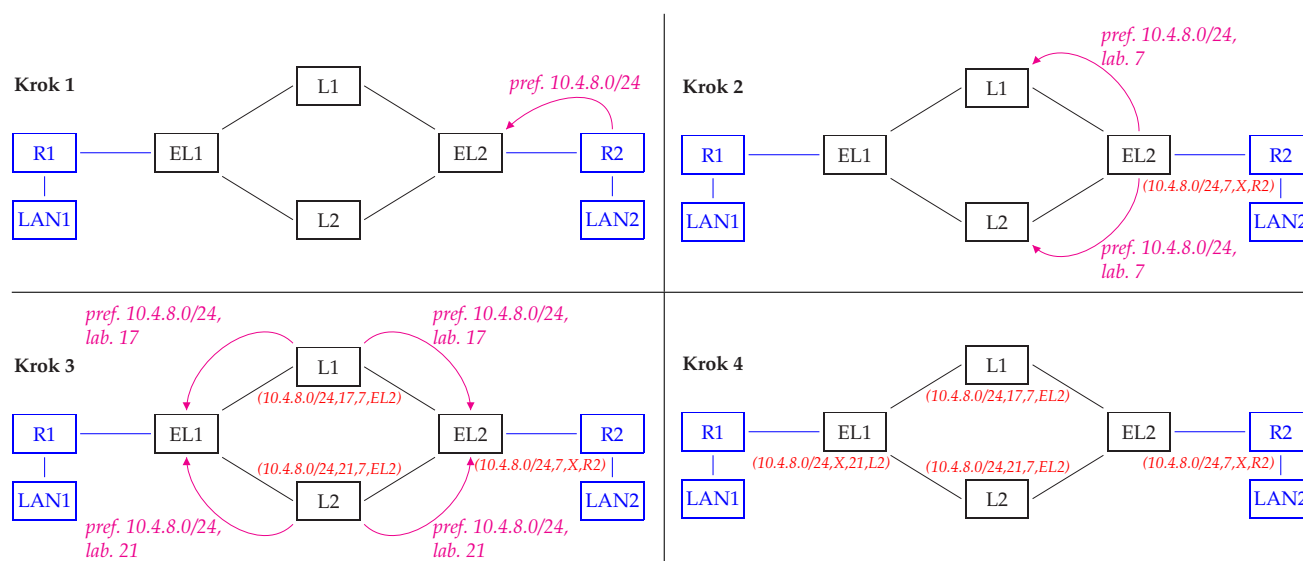
 Značka uložená v MPLS záhlaví se při průchodu směrovači neustále mění, podle toho, jak to bylo nastaveno při navazování spojení, podle obsahu tabulek na průchozích LSR. Tento proces se nazývá *Label Swapping* (výměna značek/návěští).

Konkrétní obsah tabulky na LSR závisí také na protokolu, se kterým MPLS spolupracuje (tj. na síti, nad kterou pracuje). Jestliže například MPLS pracuje nad ATM, jsou MPLS značky používány v ATM jako čísla VPI/VCI.

 **Vytváření cest.** *Protokol LDP* (Label Distribution Protocol) slouží k výměně informací o značkách mezi směrovači, výměna informací o prefixech adres probíhá podle *některého směrovacího protokolu* obecně označovaného zkratkou IGP (Interior Gateway Protocol), což může být prakticky kterýkoliv, velmi často OSPF (nyní spíše OSPFv2). U MPLS se však setkáváme i s dalšími protokoly, například při implementaci virtuálních sítí (VPN) se používá BGP. O směrovacích protokolech se budeme učit později.

 Postup přidávání záznamů do tabulek v základní variantě MPLS je naznačen na obrázku 4.10. Oznamování probíhá „proti směru“ cesty, každý uzel pro příchozí požadavek stanoví číslo (značku), které má právě volné, a pošle požadavek na všechny sousední uzly (i na ten, ze kterého požadavek přišel). Například podle obrázku 4.10 je *výstupním* LSR směrovač EL2:

- EL2 připojí LAN směrovač R2, prefix adresy je 10.4.8.0/24, přiřadí značku 7, do tabulky přidá záznam o tom, že cokoliv přijde na tuto značku a prefix, odešle směrovači R2,
- EL2 odešle sousedním MPLS směrovačům (L1 a L2) informaci „cestu k síti s prefixem 10.4.8.0/24 najdete na značce 7“,
- L1 obdrží tuto informaci, najde volnou značku 17, přidá do tabulky záznam o tom, že cokoliv přijde na značku 17 a daný prefix, v tom změni značku na 7 a pošle na směrovač EL2, odešle informaci směrovačům EL1 a EL2,



Obrázek 4.10: Připojení nové LAN do MPLS

- podobně L2, ale našel volnou značku 21,
- EL1 obdrží dvě informace, ale informace z L2 přišla dřív, tedy uloží výstupní značku 21.

**Další informace:**

- https://www.juniper.net/documentation/en_US/junos/topics/topic-map/mpls-overview.html
- <http://www.ietf.org/dyn/wg/charter/mpls-charter.html>
- http://www.cisco.com/en/US/products/ps6557/prod_white_papers_list.html
- http://www.cs.vsb.cz/grygarek/TPS/MPLS/Introduction_to_MPLS.htm
- http://www.cs.vsb.cz/grygarek/TPS/MPLS/Introduction_to_MPLS_II.htm



GMPLS (Generalized MPLS, také G-MPLS) je rozšíření MPLS, které zobecňuje MPLS na další vrstvy ISO/OSI a další rozhraní. Účelem je odstranit co nejvíce „komunikačních mezikroků“ a umožnit nasazení MPLS technologie i nad takovými řešeními, nad kterými to dřív nebylo možné (nejen nad řešeními používajícími pakety a rámce). Návěští už nemusí být jen číslo, ale může to být například vlnová délka v optické síti, fyzický port, časový slot v TDM apod., cokoliv, podle čeho lze rozšířit různé komunikace, tedy MPLS lze nasadit například i nad DWDM.

Podobné vlastnosti jako GMPLS mají síť *ASON* (Automatically Switched Optical Network), také běží nad optickými cestami či SDH, také zahrnují podporu QoS, VPN a rychlého transportu.

**Další informace:**

- <http://www.faqs.org/rfcs/rfc6002.html>
- FARREL, A., BRYSKIN, I. *GMPLS: Architecture and Applications*. San Francisco, Elsevier, 2006. Většina stran dostupná na <http://books.google.cz/books?id=cCWS75MIUpC&printsec=frontcover>
- [http://wh.cs.vsb.cz/mil051/index.php/Any_Transport_over_MPLS_\(AToM\)](http://wh.cs.vsb.cz/mil051/index.php/Any_Transport_over_MPLS_(AToM))



4.3 Infrastruktura optických sítí

4.3.1 SONET/SDH

 **SONET** (Synchronous Optical Network, ANSI T.105) je digitální optický (na optických vláknech) přenosový systém využívající časový multiplex, centrálně synchronizovaný atomovými hodinami. Je standardizován a používán v USA, Kanadě a Japonsku.

 **SDH** (Synchronous Digital Hierarchy, ITU-T G.707) je obdoba sítě SONET ve zbytku světa, rozdíly jsou jen malé (mj. v řídicích protokolech). Často se setkáme s označením SONET/SDH, odkazujícím na technologii, kterou mají obě řešení společnou.

SONET/SDH byla vytvořena pro hlasovou komunikaci a pro tento účel je také optimalizována, sama o sobě není efektivním řešením pro přenos dat. Proto nad SDH obvykle pracuje ještě jiný typ WAN/MAN sítě (nad SDH může pracovat například ATM, POS – Packet over SONET, MPLS, EoS – Gigabit Ethernet over SONET/SDH).

Podobně jako některé předchozí WAN, i zde je používána PDU o konstantní délce – *blok*. Síť pracuje v časovém multiplexu, tedy přenášené bloky dat jsou umísťovány do rezervovaných *slotů*.

 Každý blok má délku 810 oktetů. Bloky jsou posílány v pravidelných intervalech 125 μ s.

Zákazník je nucen zvolit si z několika málo typů služby podle propustnosti (2 Mb/s, 34 Mb/s, 130 Mb/s), za tuto propustnost také platí, i když ji zrovna nevyužívá (pásmo nelze pronajmout jinému zákazníkovi).

4.3.2 WDM

 **Wavelength Division Multiplexing (WDM)** (zvaná „normální“ WDM) je technologie multiplexování světla v optických spojích do vlnových délek, první řešení jsou už ze 70. let 20. století. Bez použití WDM je možné jedním optickým spojem vést jen jeden signál, s použitím WDM více původních signálů zároveň na různých vlnových délkách, na každé vlnové délce podle jiného protokolu a o jiné rychlosti.

WDM je záležitost fyzické vrstvy (L1) a v jednom optickém vlákně používá dvě vlnové délky 1310 a 1550 nm.

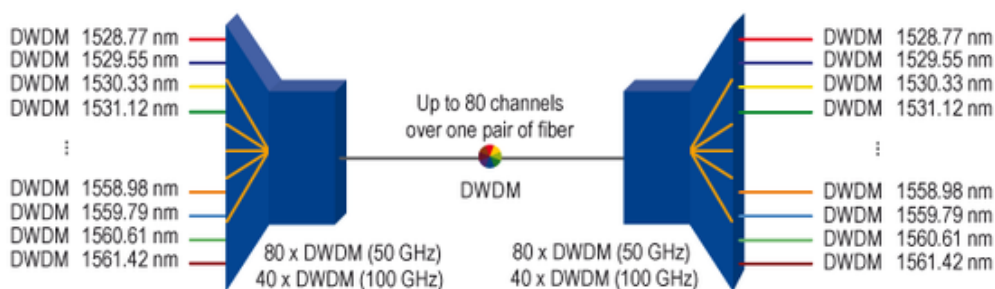
 **CWDM** (Coarse WDM, standard ITU-T G.694.2 a dále pro kabely novější standardy G.652.C a G.652.D) je řešení multiplexování do optického vlákna používající poměrně široké kanály (šířka 20 nm). Typický dosah je v desítkách kilometrů a slučuje až 16 různých vlnových délek do jednoho signálu (záleží na konkrétním standardu), což je dostačující pro metropolitní síť a hlavně celkem levné.

CWDM se dnes používá například v sítích poskytovatelů kabelové televize (přenáší se jak televizní, tak datový signál), v Ethernetu se s CWDM setkáme na fyzické vrstvě v 10GBase-LX4, také v sítích FTTx (Fibre to the Node/Building/Home/...) nebo v průmyslu na páteři při sesítování velkých areálů.

 **DWDM** (Dense Wavelength Division Multiplexing) umožňuje přenášet kanály na vlnových délkách s velmi malým odstupem (hustěji), i pod 1 nm. Čím menší odstup mezi vlnovými délkami, tím více na jednu stranu lze přenášet kanálů paralelně, ale na druhou stranu se může zhoršit kvalita přenosu, proto obvykle není využíván plně. Do jednoho optického vlákna lze poskládat až 128 kanálů, ale ve frekventovaných datových páteřních sítích se spíše setkáme s 32 kanály v jednom vlákně (jmenovitě – CESNET2). Počet

kanálů také závisí na použitých optických multiplexerech (které provádějí samotné mapování signálu do kanálu a zpět).

Obvyklá rychlost na jeden kanál (jednu vlnovou délku) je 10 Gb/s, 40 Gb/s nebo vyšší, podle použitého standardu, a celková propustnost jednoho vlákna je (téměř) součtem propustností jednotlivých kanálů.



Obrázek 4.11: Skládání signálu v DWDM²

Nad DWDM (mezi DWDM a TCP/IP) jsou pak implementovány další technologie, které mají urychlit a zefektivnit přenos, v současné době především IP/MPLS a na okrajích rozlehlé sítě pak například EoMPLS (Ethernet over MPLS), který velmi dobře spolupracuje s připojenými ethernetovými LAN.

OTN (Optical Transport Network) v sobě spojuje výhody SONET/SDH a DWDM. Z SDH bere výhodu transparentnosti, pokročilého managementu (včetně monitorování, detekce chyb) a možnost vytvoření point-to-multipoint spojů, z DWDM má efektivní multiplexování velkého množství původních signálů do optického vlákna.

OTN se dá podsunout pod Ethernet, MPLS, IP a další.

Pro implementaci existuje více různých řešení fyzické vrstvy lišících se propustností a tím, se kterými protokoly vyšších vrstev je třeba spolupracovat. Například při spolupráci s Ethernetem existují různá řešení pro 1000Base-X, 10GBase-W, 40Gb Ethernet, 100Gb Ethernet a další. Princip je podobný skládání linek v SONET/SDH, například ODU4 (pro 100Gb Ethernet) se dá složit z až dvou ODU3 nebo deseti ODU2 nebo čtyřiceti ODU1 nebo osmdesáti ODU0 modulů.

Úkol

Zařízení, která multiplexují signál do barevného spektra pro přenos optickým vláknem (a naopak), se nazývají optické multiplexery.

Najděte v některém internetovém obchodě CWDM multiplexer (například 8kanálový CWDM multiplexer od Fiber24 nebo FS, jako jsou na obrázku 4.12) a zjistěte jeho parametry. Je zajímavé, že pokud jde opravdu o čistě optické zařízení, tak může fungovat čistě pasivně bez nutnosti napájení.



Obrázek 4.12: Příklady 8kanálových CWDM multiplexerů³

²Zdroj: <https://www.pandacomdirekt.com/en/technologies/wdm/what-is-dwdm.html>

³Zdroje: <https://shop.fiber24.net>, <https://www.fs.com/>

Zařízení posílající a přijímající data do/z těchto multiplexerů potřebují optické konvertory (převádějící elektrický signál na optický a naopak), obvykle s moduly SFP/SFP+/XFP/... Pokuste se najít vhodný konvertor na tomtéž webu, kde jste našli multiplexer.



Další informace:

- <https://www.optcore.net/introduction-to-dwdm/>
- <https://www.itu.int/itu-t/recommendations/rec.aspx?rec=13076>
- <https://pon.fibrain.com/artykuly-techniczne/optical-transmission-lexicon-catv-cwdm-dwdm-ftth,44.html>
- <http://www.fiber-optic-cable-sale.com/dwdm-vs-otn-whats-difference.html>



4.4 Telekomunikační síť

4.4.1 Využití telekomunikačních sítí pro přenos dat

Telekomunikační síť je síť skládající se z mezilehlých zařízení přenášející různé typy zpráv. Řadíme zde jak sítě telefonních operátorů (včetně mobilních), tak i počítačové sítě, zvláště MAN a WAN sítě.

 Typická architektura telekomunikační sítě je založena na *rovinách* – planes (protínajících různé vrstvy ISO/OSI):

- *data plane* (user plane) – tato rovina je viditelná i pro uživatele, zde se přenášejí uživatelská data,
- *control plane* (signalling) – pro přenos řídicích informací,
- *management plane* – pro administraci, management sítě.

Vysvětlíme si několik pojmů z oblasti telekomunikací.

 *Telefonní síť* má hierarchickou strukturu:

- *účastnické zařízení* (telefon, modem apod.), připojuje se přes *účastnickou přípojku* (outlet),
- (veřejná) *místní telefonní ústředna* (MTO) – v pravidelných intervalech na ulicích, v angličtině se označuje většinou jako *street cabinet*,
- *uzlová telefonní ústředna* (UTO),
- *tranzitní telefonní ústředna* (TTO).

V rámci firemního areálu může také fungovat *pobočková ústředna* (PBX, Private Branch Exchange), která bývá připojena na některou veřejnou místní telefonní ústřednu (tj. funguje jako brána).

V telefonní síti platí: pro přenos hlasu plně postačují frekvence z intervalu 300–3400 Hz, ořezání okolních frekvencí prakticky neovlivní kvalitu hovoru. Přenos musí být multiplexován, v analogové síti je volen *frekvenční multiplex* s šířkou pásma pro jeden přenos zaokrouhlenou nahoru (s rezervou) – 4000 Hz.

Z toho vyplývá, že vyšší frekvence se v telekomunikačním kabelu dají využít pro přenos dat. Pokud bychom netrvali na přenosu analogového hlasu (na čemž v poslední době většinou netrváme), mohli bychom pro data využít i nižší frekvence.

 Claude Elwood Shannon je zakladatel moderní teorie informace. *Shannonův teorém* udává strop pro maximální rychlost přenosu.

„Aby bylo možné rekonstruovat (na cílové stanici) spojitý frekvenčně omezený analogový signál ze vzorků, musí být vzorkován s frekvencí alespoň dvakrát vyšší než je maximální frekvence tohoto signálu při přenosu.“

Důsledkem je ohraničení maximální dosažitelné rychlosti přenosu, a to

- šířkou pásma (počet stavů v rámci pásma nelze zvyšovat do nekonečna, ztratili bychom možnost jeho rozlišení),
- kvalitou signálu (vyplývá z fyzikálních vlastností linky, zhoršená šumem, také „odstup signálu od šumu“),

ale nezávisí na použité technologii (včetně použité modulace). Vzorec:

$$\max(\text{rychlost přenosu}) = \text{šířka pásma} * \log_2(1 + \text{signál/šum})$$

Podle Shannonova teorému můžeme rychlost přenosu (nebo dosah) zvýšit buď tím, že rozšíříme frekvenční pásmo (směrem nahoru), nebo zvýšíme kvalitu signálu – kabely asi vyměňovat nebudeme, ale existují různé možnosti ovlivňování šumu. V každém případě pro přenos dat je schůdnější spíše první možnost.

 **PSTN** (Public Switched Telephone Network) – běžná telefonní síť dimenzovaná pro přenos hlasu, původně čistě analogová s frekvenčním multiplexem, po digitalizaci se používá časový multiplex s rezervací komunikačního pásma. Pozor, digitalizace se týká především telekomunikačních ústředn, část spoje u zákazníka je analogová.

Po digitalizaci bylo nutno (v principu analogový) hlas v MTO digitalizovat, používal se některý kodek, tedy CODEC (COder-DECoder). Podobný princip se dnes používá i ve VoIP přímo v koncovém zařízení.

 **POTS** (Plain Old Telephone Service) je technologie využívající PSTN pro přenos analogového hlasu a digitálních dat. Protože PSTN je zkrájí analogová a až od ústředny digitální, je třeba digitální data modulovat na analogový signál, používá se MODEM (MODulator-DEModulator).

 **Modem** je tedy zařízení, které moduluje digitální signál na analogový a v cíli ho zpětně demoduluje do digitální podoby. Modemy byly dříve analogové, dnes používáme digitální modemy (pro ADSL, VDSL apod.) – telekomunikační linka na straně zákazníka je prostě pořád analogová.

 **ISDN** (Integrated Services Digital Network) byla prvním pokusem o digitalizaci přístupových sítí, ale tak trochu na půl cesty. Šlo o vytáčené připojení (tarifikace podle času), protože správa sítě stavěla na původní POTS, vlastně to byla jen digitální nástavba POTS.

 **xDSL** je skupina technologií, které se oprostily od většiny součástí POTS, z telefonní sítě se využívá pouze spoj mezi účastnickou přípojkou a nejbližší ústřednou MTO. Pak komunikační cesta odbočuje do plně digitální počítačové WAN sítě poskytovatele.

4.4.2 Slučování linek

Předchůdcem sítě SDH, se kterou jsme se seznámili v předchozím textu, je síť PDH (Plesiochronous Digital Hierarchy). Zatímco SDH je „synchronní“ (komunikace je synchronizována v celé síti, s časovým multiplexem), PDH je „téměř synchronní“ – plesiochronní, tedy komunikace je téměř synchronizována (tolerují se malé odlišnosti v rychlosti).

V PDH se začalo využívat *slučování (agregace) linek* za účelem zvýšení kapacity přenosu. Značení a některé další prvky se odlišují v normách platných pro Evropu od norem platných v Americe a Japonsku. Zatímco v Americe a Japonsku se setkáme s T-linkami (T-carrier, normy podle ITU-T), v Evropě se používají E-linky (E-carrier, normy podle CEPT – Evropské rady pro správu pošt a telekomunikací).

Princip vychází ze sdružování linek do svazků, které mají logicky vyšší přenosovou kapacitu a tím i rychlost přenosu. V obou normách jsou linky řádu 0 (nula) tvořeny jedinou linkou o rychlosti 64 kb/s. V řádu 1 se sdružuje buď 24 nebo 30 linek datových a 2 linky signální (nepřenášejí data, ale řídicí signály), vzniká T1-linka (Amerika a Japonsko) nebo E1-linka (Evropa). Sdružením čtyř T1 linek vznikne T2 linka, podobně u E-linek. Postup a navyšování přenosové rychlosti najdeme v tabulce 4.2.

Značení	Dat. kanálů	Přen. rychlost	Značení	Dat. kanálů	Přen. rychlost
T1	24	1,54 Mb/s	E1	30	2,05 Mb/s
T2	96	6,31 Mb/s	E2	120	8,45 Mb/s
T3	672	44,38 Mb/s	E3	480	34,37 Mb/s
T4	4032	274,18 Mb/s	E4	1920	139,27 Mb/s
			E5	7680	565,15 Mb/s

Tabulka 4.2: Srovnání T-carrier (Amerika, Japonsko) a E-carrier (Evropa) linek

SONET	SDH	Přen. rychlost	SONET	SDH	Přen. rychlost
STS-1	STM-0	52 Mb/s	STS-192	STM-64	9 953 Mb/s
STS-3	STM-1	156 Mb/s	STS-768	STM-256	39 813 Mb/s
STS-12	STM-4	622 Mb/s	STS-3072	STM-1024	159 252 Mb/s
STS-48	STM-16	2 488 Mb/s			

Tabulka 4.3: Úrovně a rychlosti v sítích SONET a SDH

Fyzicky bývají sdružovány spoje na kroucených dvoulinkách (UTP), což početně souhlasí: $30 + 2 = 32$, což je dělitelné číslem 4 (počtem drátů v UTP).

Rámec E1 se skládá z 32 timeslotů („oken“ o stejné délce – 125 μ s), přenášených paralelně. Jeden timeslot obsahuje prostor pro jeden oktet (8 bitů), od toho se odvíjí přenosová kapacita.

SDH se od PDH liší především využitím optických kabelů (jednovídrových) a dokonalejší synchronizací, která se provádí přes celou síť pomocí atomových hodin. Velikost timeslotu je stejná (125 μ s), ale multiplexování funguje trochu jinak (protože už se neodvíjí od UTP kabelu).

Na slučování linek stála i technologie ISDN – přenosová kapacita je v této technologii rozdělena na B-kanály a D-kanál, existují dva druhy rozhraní

- BRI (u nás euroISDN2) pro domácnosti a malé kanceláře, slučuje dvě datové a jednu signální linku,
- PRI (u nás euroISDN30) pro větší firmy, dražší, slučuje 30 datových a jednu signální linku.

V BRI je bitový tok rozdělen na časové rámce po 48 bitech (0,25 ms) – 2 oktety pro každý B-kanál, 4 bity pro D-kanál a 12 doplňkových bitů, rychlost u BRI je maximálně 160 kb/s. B-kanály pracují s přepojováním okruhů (proto tarifikaace podle času – platí se za dobu existence přepínaného okruhu SVC) a mají vyšší přenosovou rychlost, D-kanály pracují s pakety, a to se spojením i bez spojení.

**Poznámka:**

Zde sice o agregaci linek hovoříme vzhledem k telekomunikačním sítím, ale linky je možno agregovat i lokálně. V předchozím textu jsme se setkali například s technologií EtherChannel pro sloučení několika ethernetových linek, další termíny jsou Link Aggregation a NIC Teaming. Používají se pro navýšení přenosové kapacity spoje mezi serverem (nebo jiným zařízením) a switchem. Ovšem je třeba mít na daném zařízení více ethernetových portů a podporu příslušné technologie na obou stranách spoje.



<http://www.samuraj-cz.com/clanek/pripojeni-rychlejsi-a-spolehlivejsi/>



4.5 Přístupové sítě



Poskytovatelé Internetu (ISP, Internet Service Providers) provozují svou vlastní WAN síť, přes kterou přenášejí data svých zákazníků, ale tato data se nějak do dané sítě musejí dostat. Tedy je potřeba nějaká síť přemostující vzdálenost mezi sítí zákazníka a sítí ISP. Takové sítě se označují jako *přístupové sítě*, nebo také *sítě první míle* (first-mile networks) nebo *poslední míle* (last-mile networks), protože jsou na začátku cesty dat do Internetu nebo naopak na jejím konci.

4.5.1 Protokol PPP

PPP je jeden ze spojových protokolů původně pro sériové spoje, ale na rozdíl od jiných je poněkud univerzálnější (ve smyslu spolupráce s jinými protokoly) a jednodušší (je osekáný jen na ten základ, který je při jeho používání potřebný). S protokolem PPP a jeho odvozeninami se dnes setkáváme zejména v přístupových sítích, plní tam podobnou roli jako Ethernet.



Protokol *PPP* (Point-to-Point Protocol) je určen pro duplexní point-to-point spoje. Nerozlišuje primární a sekundární uzly, podporuje synchronní i asynchronní přenos. Dokáže provádět automatickou konfiguraci při navazování spojení, přičemž se testuje kvalita spoje. Součástí rámce je i určení protokolu vyšší vrstvy, proto dokáže nejen zapouzdřovat IP pakety, ale i jiné typy paketů vyšší vrstvy. Důležitou vlastností je i možnost autentizace.

Řízení kvality linky funguje tak, že na lince je nastaveno „procento kvality“, a pokud jsou přenášené rámce natolik poškozeny a ztráceny, že kvalita přenosu klesne pod tuto úroveň, linka je vypnuta (a pak by měla najet záložní linka nebo by se měla použít jiná cesta pro data).



Specifikace PPP se skládá ze dvou částí (přesněji – obsahuje dvě vrstvy):

- *protokoly řízení sítě* (NCP, Network Control Protocols) – protokoly zajišťující komunikaci s vyšší vrstvou; existuje více těchto protokolů pro různé síťové architektury (TCP/IP, TCP/IP_{v6}, NetWare, AppleTalk apod.), tato vrstva zasahuje částečně i do síťové vrstvy v ISO/OSI modelu,
- *protokol řízení spoje* (LCP, Link Control Protocol) – navazování a udržování spojení, dohodnutí konfigurace na začátku spojení (negociace), může zajišťovat také testování spoje a autentizaci.

Účelem je maximálně omezit množství funkcí vlastního spojového protokolu (tj. LCP) a takto zefektivnit jeho fungování.



Autentizace je volitelnou vlastností, lze ji provádět některým z těchto protokolů:

- *PAP* (Password Authentication Protocol) – textové heslo se posílá po spoji, tedy nejde o bezpečnou metodu autentizace,

- *CHAP* (Challenge Handshake Authentication Protocol, RFC 1994) – používá se asymetrická šifra (klíč), iniciující stanice odešle žádost, obdrží náhodně vygenerovanou sekvenci znaků, kterou zpracuje pomocí tajného klíče a odešle zpět, po ověření příjemcem může začít komunikace,
- *EAP* (Extensible Authentication Protocol, RFC 2284) – zcela odděluje fázi navazování spojení od fáze autentizace,
- MS-CHAP je drobným vylepšením CHAP,
- lze propojit s AAA mechanismy s využitím autentizačních serverů (Radius, Diameter, TACACS+).

PPP také podporuje šifrování a komprimaci.



Další informace:

O autentizaci EAP v PPP je informace například na

<https://packetlife.net/blog/2010/jun/14/ppp-authentication-eap/>.



Rámec protokolu PPP je do značné míry podobný rámcům jiných spojových protokolů:

- příznak (opět binárně 01111110),
- adresa (1 B) – toto pole existuje, ale má vždy konstantní hodnotu (binárně samé 1),
- řídicí pole (1 B) – obsahuje vždy hodnotu 0000011 (indikuje nečíslovaný rámec ve službě bez spojení),
- protokol (2 B, toto pole se u jiných protokolů nevyskytuje) – obsahuje identifikaci protokolu vyšší vrstvy (například pro IPv4 je to 0×0021, pro IPv6 0×0057, pro IPX je zde 0×002B), k nalezení v RFC dokumentech,
- data – maximálně 1500 B,
- FCS,
- příznak.

To byl rámec obsahující data pro přenos. Používají se také služební LCP pakety, které je třeba zapouzdřit do PPP rámce následovně:

- v poli protokolu je hodnota 0×C021,
- v poli data je samotný LCP paket obsahující tyto informace:
 - kód – identifikuje funkci LCP paketu (požadavky na změnu parametrů, potvrzení přijetí požadavků na změnu, odmítnutí požadavků, ukončení spojení, atd.),
 - identifikátor – u příkazového paketu obsahuje náhodně vygenerovanou hodnotu, u odpovědi musí obsahovat tutéž hodnotu jako paket s příkazem, na který odpovídá,
 - délka celého LCP paketu v oktetech,
 - data – obvykle posloupnost uspořádaných dvojic (volba, hodnota), například typ autentizačního protokolu, vyšší max. délka dat, atd.



Rozšíření PPP. Pro rozšíření protokolu PPP platí v podstatě totéž co pro PPP (vrstva L2, přibližně účel, šifrování, většinou i formát rámce). Tato rozšíření jsou navíc přizpůsobena konkrétnímu způsobu využití.

Multilink PPP (PPP po více spojích, také se označuje MP, MPPP, MLP; RFC 1990) je varianta PPP pro vícebodové spoje. Není tím míněno ani tak více uzlů, jako spíše více sériových spojů pro jedinou komunikaci. Lze zkombinovat více sériových spojů do jediného logického kanálu s vyšší propustností

(agregace), protokol dokáže data rozdělit, vhodně přeuspořádat a pak na druhém konci linky opět uvést do původního stavu.

Tunneling PPP (PPTP, Point-to-point Tunneling Protocol, RFC 1171) je určen pro virtuální privátní síť a k připojení vzdáleného klienta k firemní síti (obecně LAN). Ovšem dnes je podporován prakticky už jen ve Windows a jeho podpora se bude omezovat i zde, doporučuje se přechod na jeho následovníka *L2TP* (Layer 2 Tunneling Protocol, RFC 3931).

L2TP: do rámce protokolu L2TP se zapouzdřují rámce PPP. Sám o sobě nemá tento protokol zabezpečovací funkce, proto se často kombinuje se zabezpečením pomocí IPSec nebo jiných metod. K L2TP se vrátíme v kapitole o bezpečnosti, je to jeden z protokolů pro vytváření VPN (tunelů).

Připojení k Internetu se řeší napojením PPP (iniciační fáze připojení) na protokoly WAN. Využívá se také v technologiích xDSL pro přenos přes WAN poskytovatele, a to v těchto variantách:

- PPPoA (PPP over ATM, RFC 2364) – PPP paket je zapouzdřen do ATM buněk,
- PPPoE (PPP over Ethernet, RFC 2516) – PPP paket je zapouzdřen do ethernetového rámce, dnes běžné.

Tyto varianty umožňují navázat autentizované spojení s ISP a dynamicky získat IP adresu. Jsou potřeba u typů komunikace, kde se navazuje spojení, nevyužívají se u trvalého připojení k Internetu (resp. využívají se tehdy, když se navazuje spojení; když už je navázáno, nejsou potřeba). Autentizace bývá jen na straně zákazníka (tj. zákaznické zařízení prokazuje svou totožnost poskytovateli služby).



Další informace:

- http://www.tcpipguide.com/free/t_PointtoPointProtocolPPP.htm
- <https://www.ciscopress.com/articles/article.asp?p=2202412&seqNum=5>
- <https://www.geeksforgeeks.org/point-to-point-protocol-ppp-frame-format/>



4.5.2 Digitální přístupové telekomunikační síť

Pro přístupové síť (první míle, poslední míle) existují různé technologie – může jít o bezdrátové síť (Wi-fi v kombinaci s dalšími protokoly, WiMAX) nebo telekomunikační síť (většinou ADSL/VDSL nebo podobná technologie), koaxiální linky pro tzv. kabelový internet, optické přístupové síť (FTTx), v některých případech se přístupová síť dokonce jeví jako speciální součást WAN sítě (mobilní síť, třeba LTE).

Specifikem přístupové sítě je, že spojuje velice odlišné technologie vzhledem ke způsobu komunikace (v LAN sítích máme typicky paketový přenos dat, kdežto ve WAN sítích se komunikuje přes přepínané okruhy nebo něco podobného), a také vzhledem k fyzické vrstvě a zacházení se signálem. Potřebujeme mít kontrolu nad tokem dat jednotlivých zákazníků.

V této sekci se zaměříme spíše na komunikaci přes telekomunikační linku, na bezdrátové a mobilní síť se podíváme v další kapitole.



xDSL (Digital Subscriber Line) je skupina širokopásmových technologií, která v sobě sdružuje několik různých technologií (ADSL, SDSL, HDSL, HDSL-2, G.SHDSL, IDSL, VDSL). Jedná se o vyhrazenou službu, point-to-point, spojově orientovanou, tarifkace není závislá na době připojení (na rozdíl od staré ISDN nebo ještě staršího spojení přes analogové modemy).

 **Asynchronous Digital Subscriber Line (ADSL).** ADSL je asymetrická služba. Asymetrická proto, že downstream (stahování dat do DTE) je rychlejší než upstream (odesílání dat do sítě).

Teoreticky lze dosáhnout rychlostí až 24 Mb/s (v novějším standardu dokonce až 48 Mb/s). U nás nabízená rychlost je 8 Mb/s nebo 16 Mb/s *pro downstream*, ale reálná rychlost může být nižší. Upstream bývá na rychlosti 1 Mb/s.

Rychlost je ovlivněna více různými kritérii. Je to nejen použité zařízení a přenosové protokoly, ale také délka spoje – čím větší vzdálenost k místní ústředně, tím pomalejší spojení.

Rozsah	Šířka	Účel
0 kHz – 4 kHz	4 kHz	přenos hlasu (telefon)
4 kHz – 26 kHz	22 kHz	nárazníkové pásmo
26 kHz – 138 kHz	112 kHz	upstream
138 kHz – 1100 kHz	962 kHz	downstream

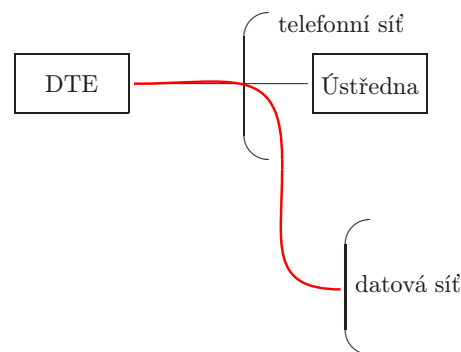
Tabulka 4.4: Využití přenosového pásma v ADSL

V ADSL se používá modulace *DMT* (Discrete MultiTone). DMT dělí celé pásmo na kanály přibližně po 4 kHz (celkem kolem 255 kanálů), kanály 7–32 pro upstream, od 33 pro downstream. Průběžně se sleduje kvalita přenosu v jednotlivých kanálech, přenos se adaptivně rozkládá na ty kanály, které jsou považovány za právě nejkvalitnější. Na rozdíl od CAP (která byla použita v první verzi ADSL a vychází z QAM) se používá více nosných kmitočtů (to jsou ty tóny v názvu, také se nazývá Multicarrier Modulation).

V technologii ADSL (i v dalších xDSL technologiích) se navýšení rychlosti dosahuje kromě jiného i odkloněním datového toku z telefonních linek na datovou síť poskytovatele (WAN síť na optických kabelech). Princip je naznačen na obrázku 4.13.

Agregace znamená sdílení přípojky ADSL. Každý ISP rozdělí kapacitu linky, kterou má k dispozici, formou časového multiplexu. Vzorec pro agregační poměr je následující:

$$\text{agregační poměr} = \frac{\text{rychlost}(ISP)}{\sum_i \text{rychlost}(i)}$$



Obrázek 4.13: Provedení ADSL

Typické hodnoty agregace mohou být například 1 : 50, 1 : 20.

FUP (*Fair User Policy*) je omezení toku dat (většinou se odvozuje od množství přenášených dat). V současné době není u ADSL FUP uplatňováno.

FUP v ADSL obvykle souvisí s protokolem PPPoA. Proto pokud zjistíme, že FUP je na naší lince uplatňováno, měli bychom projít konfiguraci ADSL modemu a zjistit, jestli je možné nastavit místo PPPoA (PPP over ATM) protokol PPPoE (PPP over Ethernet). U starších ADSL modemů to nejde, ale u novějších by to neměl být problém.

 **Zařízení v síti ADSL:** když nepočítáme koncová zařízení v síti zákazníka (která ve skutečnosti do sítě ADSL vůbec nepatří), v ADSL se používají tato zařízení:

- *ADSL modem* (MODulator/DEModulator) – moduluje odchozí digitální signál na analogový a přijíma demoduluje, připojují se k němu datová koncová zařízení,

EchoLife HG520i

- Status
 - System information
 - Service information
 - Statistics
- Basic
- Advanced
- Tools

Service Information

LAN Interface

IP Address	Subnet	MAC Address	Status
	255.255.255.0		✓

Port	Speed	Duplex	Status

WAN Interface

PVC	VPI/VCI	IP Address	Subnet	Gateway	Encapsulation	Status
PVC-0	8 /48				Route -PPPoE	✓
PVC-1	8 /35				Bridge -RFC2684	✓

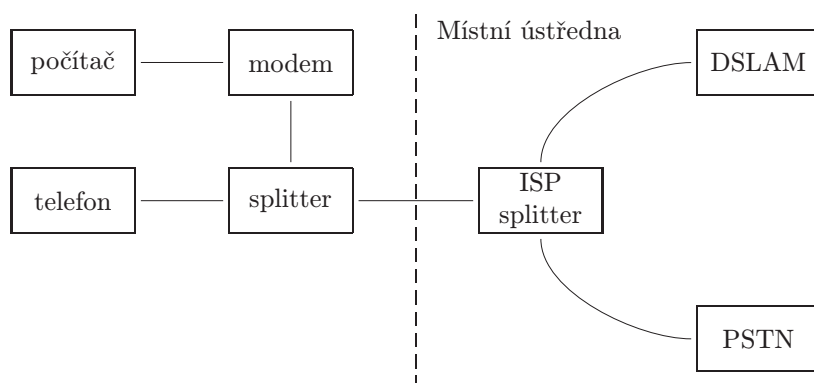
WAN Setting

PVC	0
VPI	8
VCI	48
Active	Yes
Mode	Routing
Encapsulation	PPPoE
Multiplex	LLC

Obrázek 4.14: Zapouzdření do PPPoE

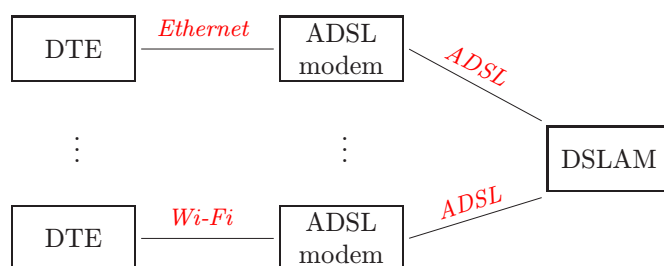
- *splitter* – slučuje analogový hlasový přenos (z analogového hlasového koncového zařízení) a modulovaná data (z ADSL modemu) a naopak v opačném směru, dnes se už prakticky nepoužívá, protože málokde najdeme analogové telefony,
- *DSLAM* (DSL Access Multiplexer) – na straně ISP, sdružuje spojení z ISP splitterů nebo ADSL modemů pro jednotlivá připojení.

Je výhodnější mít splitter zvlášť, pokud vůbec. Pokud nepoužíváme žádná analogová zařízení (analogový telefon), měli bychom se splitteru zbavit, protože může být zdrojem mírných posunů signálu a tím i poruch.



Obrázek 4.15: Zařízení v síti ADSL

V SOHO (Small Office, Home Office) je ADSL modem (nebo novější VDSL) jen modulem v poněkud složitějším zařízení, například v tzv. ADSL Wi-fi routeru (takže ADSL modem + Wi-fi access point + integrovaný router, ať je uvnitř i funkcionality vrstvy L3, pak třeba firewall, atd.). Takové zařízení se připojuje přes konektor RJ-11 (telefonní kabel) k lince poskytovatele dané služby. Funguje jak pro

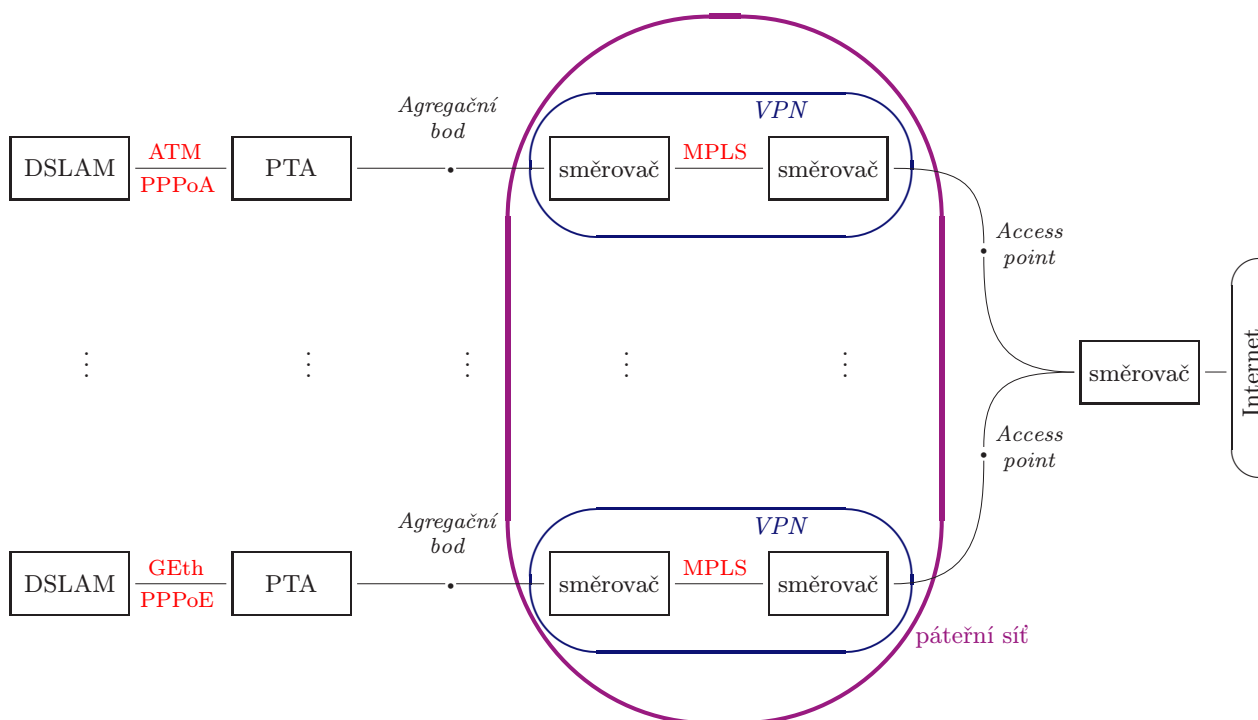


Obrázek 4.16: ADSL je technologie první míle

navázání spojení (když je zařízení zapínáno či resetováno), tak i pro samotnou komunikaci. Zapouzdřuje pakety (třeba IPv4, IPv6) do protokolových datových jednotek PPPoE, PPPoA nebo IPoA.

 *DSLAM* je brána do datové sítě poskytovatele. Existují menší DSLAMy pro připojení 24 nebo 48 DSL linek (a obvykle stejný počet pro POTS – analogový telefon), a pak větší pro tisíce linek. Obvykle podporuje protokoly PPPoE i PPPoA.

DSLAM má rozhraní RJ-11 směrem k zákazníkovi a pak další rozhraní – optika do WAN sítě provozovatele, případně rozhraní pro správu, atd. Administrace se provádí přes webové rozhraní, anebo přes konzoli, záleží na konkrétním výrobci.

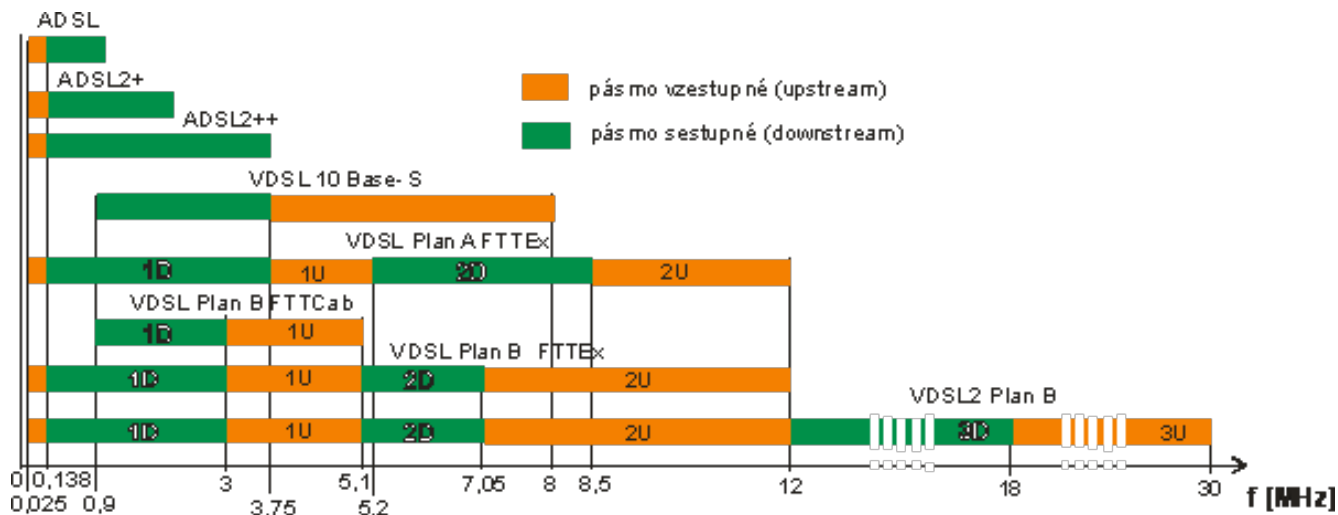


Obrázek 4.17: ADSL – na straně ISP

Na obrázku 4.17 je zjednodušený náčrtek struktury sítě na straně poskytovatele Internetu. Pojmy k tomuto obrázku:

- *PTA* je širokopásmový server, provádí se zde konfigurace IP adres, autentizace uživatelů, autorizace, účtování (např. RADIUS),
- *AP* (*Access Point*) je přístupový bod pro konkrétního ISP, přes který lze přistupovat z virtuální sítě tohoto ISP k směrovači na Internet.

Very High Bit-Rate Digital Subscriber Line (VDSL). VDSL je následníkem ADSL. Je také asymetrická (může být nastavena na symetrický přenos), používá širší pásmo než ADSL (protáhnutí na vyšší frekvence), a z toho důvodu je rychlejší.



Obrázek 4.18: Porovnání frekvencí ADSL a VDSL⁴

Z obrázku 4.18 je patrné, že oproti ADSL technologii VDSL využívá širší frekvenční pásmo, tedy je k dispozici více komunikačních kanálů. Zatímco u ADSL je na nižších frekvencích upstream a na vyšších downstream, u VDSL jsou kanály pro oba směry rozloženy celkem rovnoměrně v celém spektru (důsledkem je, že zatímco u ADSL při přechodu na novou verzi využívající vyšší frekvence byl navýšen jen downstream, protože upstream z důvodů kompatibility měl pásmo pevně stanovené, u VDSL mohou být navyšovány rychlosti v obou směrech).

Zatímco ADSL provádí přenos na vzdálenost až 5 km (při vyšších rychlostech méně) – k nejbližšímu MTO nebo předsunutému DSLAMu, VDSL pouze něco přes 1 km, což je daň za rozšíření pásma. Ovšem pokud chceme opravdu využít vysoké rychlosti nabízené VDSL, tak bychom měli být maximálně 500 m od DSLAMu.

VDSL verze 1 (ITU G.993.1) používá multiplexování QAM nebo DMT, VDSL verze 2 (ITU G.993.2) používá pouze DMT.

Vectoring je technika, jejímž úkolem je snížit rušení mezi telefonními linkami nebo mezi různými páry v rámci jednoho kabelu, a má smysl zejména u VDSL 2, protože zde se vysílá na vyšších frekvencích než u starších technologií. Vectoring funguje tak, že signál na lince je neustále monitorován a upravován podobným způsobem, jak funguje například technologie pro potlačení šumu ve sluchátkách. Vectoring má smysl zejména na krátkých spojit do několika stovek metrů, na delších VDSL linkách už nemá tak výrazný efekt.

Aby vectoring fungoval, musí být podporován oběma stranami (tedy musí být podpora v DSLAMu i ve VDSL routerech všech připojených klientů), což už dnes nebývá problém.

Verze 2 nabízí rychlosti až 100 Mb/s v obou směrech paralelně, ale pouze do maximální vzdálenosti 300 m k připojenému DSLAMu. Typické rychlosti jsou 40 Mb/s, s dalšími doprovodnými technologiemi více.

⁴Zdroj: <http://access.feld.cvut.cz/view.php?cisloclanku=2004120302>



Úkol

Najděte na Internetu jakýkoliv středně velký DSLAM nebo MSAN. Pokud nemůžete nic najít, zkuste se zaměřit například na MSAN „Zyxel 2U 5-slot Temperature-Hardened MSAN IES4105 Series“. Zjistěte, jaké má parametry.



Další xDSL technologie. Kromě ADSL a VDSL existují i další xDSL technologie. Liší se v počtu potřebných RJ-11 rozhraní, (a)symetrii a dalších parametrech.

High bit-rate Digital Subscriber Line (HDSL) není určena pro účastnické přípojky, používají ji ISP pro propojení pobočkových ústředí, případně ji mohou využít velké firmy pro své vnitřní potřeby. Vyžaduje dva nebo tři páry telefonních vodičů, což je považováno za jednu z nevýhod.

Je to symetrická technologie (stejný rozsah pásem pro downstream a upstream), ale může pracovat i asymetricky. Podporuje pouze přenos dat, neobsahuje podporu telefonních hovorů.

HDSL-2 (také SHDSL – SinglePair HDSL) je varianta HDSL, která používá pouze jeden pár telefonních vodičů, tedy při zavádění technologie HDSL-2 na samotném vedení telefonních linek není třeba nic měnit.

Symmetric Digital Subscriber Line (SDSL) je rozsahem funkcí podobná technologii ADSL nebo VDSL, ale je symetrická, neobsahuje podporu telefonních hovorů. Princip přenosu je podobný HDSL. Používá se pro účastnické přípojky, u kterých se preferuje symetričnost přenosu.

4.5.3 DOCSIS

Data Over Cable Service Interface Specification (DOCSIS) je standard pro tzv. „kabelový internet“, který představuje širokopásmové (broadband) spojení přes existující televizní sítě postavené na koaxiální kabeláži (CATV – Community Antenna TV). Evropská varianta je EuroDOCSIS. DOCSIS verze 3 je standardizována jako ITU-T J.222.



Momentálně jde spíše o hybridní infrastrukturu – kombinaci optických linek (v síti poskytovatele) a koaxiálních linek (vedoucích k zákazníkům, kvůli zpětné kompatibilitě a ceně). Tento princip se nazývá infrastruktura *hybrid fiber-coaxial (HFC)*.

DOCSIS je asymetrická, stejně jako ADSL/VDSL. Typické rychlosti jsou vyšší než u xDSL, DOCSIS verze 3.1 z roku 2013 (D3.1) a verze 4 z roku 2017 (D4) mohou přenášet data rychlostí až 10 Gb/s na downstreamu, upstream je pomalejší.

Na fyzické vrstvě je na jednotlivých kanálech obvykle používána modulace QAM (v závislosti na kvalitě linky to je 4096-QAM nebo pomalejší, může spadnout i na 256-QAM), nebo jiná modulace podle toho, co vše je přenášeno, například signál DVB-C. Kanály jsou kombinovány metodou OFDM nebo OFDMA.



Podobně jako xDSL technologie používají na straně poskytovatele služby zařízení DSLAM, DOCSIS používá *Cable Modem Termination System (CMTS)*. Role je stejná – CMTS je umístěn někde ve skříní (obdobě MTO) na ulici a k němu jsou připojeni zákazníci, obvykle koaxiálem.



Další informace:

- https://www.rohde-schwarz.com/es/tecnologias/tv-por-cable/docsis/tecnologia-docsis/tecnologia-docsis_55513.html
- <https://community.cisco.com/t5/networking-documents/docsis/ta-p/3115673>

- https://www.ieee802.org/3/efm/public/jul01/presentations/gummalla_1_0701.pdf



Úkol

Najděte na Internetu nějaký DOCSIS CMTS. Pokud nemůžete nic najít, hledejte například zařízení „Casa CMTS/CCAP C40G“.

Zjistěte parametry nalezeného zařízení (například pro výše uvedené zařízení na <https://www.casa-systems.com/assets/Casa-Datasheet-C40G.pdf>). Jaké verze DOCSIS jsou podporovány? Který typ portu je používán pro management? Je možné se napojit přes SSH? Podporuje DHCP? Kolik portů pro downstream (k zákazníkům) je přítomno a jaký typ konektoru?

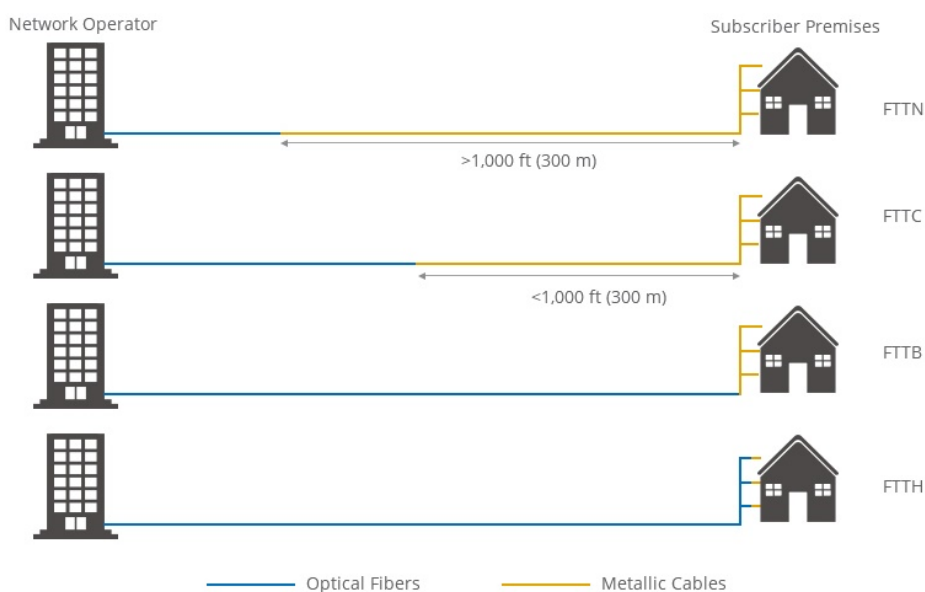


4.5.4 FTTx

Fiber-to-the-x (FTTx) je optická širokopásmová technologie. Písmeno „x“ určuje většinou to, jak daleko vede optika:

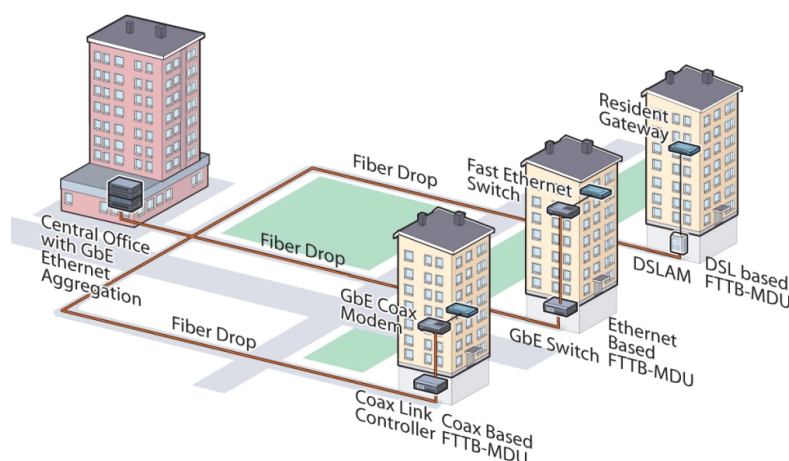
- FTTN (fiber-to-the-node or neighborhood): optický kabel je ukončen ve skříni na ulici (street cabinetu), odkud pokračuje jiná technologie (xDSL, DOCSIS apod.),
- FTTC (fiber-to-the-cabinet or curb): podobné jako FTTN, ale skříň na rozhraní s uživatelskými linkami je blíže zákazníkům (nejvýše stovky metrů),
- FTTB (fiber-to-the-building): konec optické sítě je někde v budově, kde se nacházejí zákazníci (obvykle někde v přízemí ve sdílené místnosti, kde se řeší „elektrina“, odtud vedou spoje podle jiné technologie do jednotlivých bytů či kanceláří,
- FTTH (fiber-to-the-home): optika až do bytu, což je nejkomfortnější, ale také nejdražší varianta.

První dva typy se taky označují jako FTTP (Fiber-to-the-Premises).



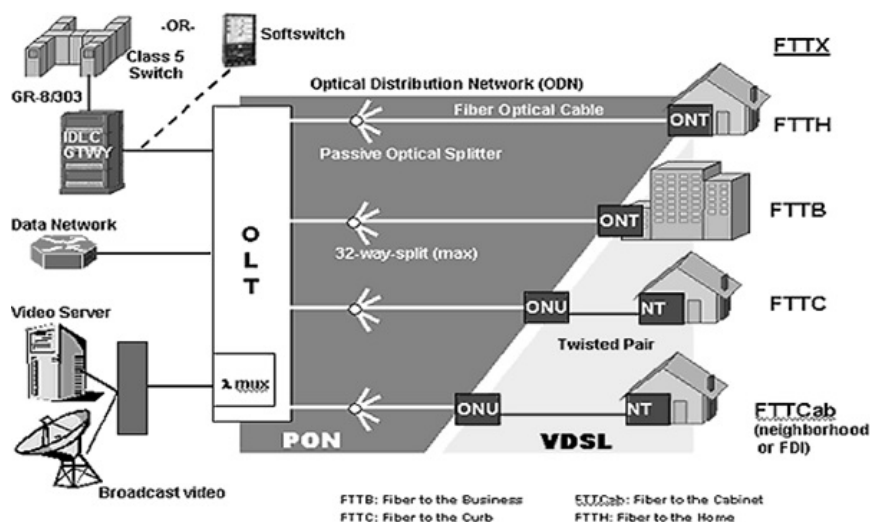
Obrázek 4.19: Možnosti řešení FTTx⁵

⁵Zdroj: <https://www.fs.com/blog/comprehensive-understanding-of-fttx-network-8176.html>

Obrázek 4.20: Tři způsoby využití FTTx⁶

 FTTx síť je hierarchická, podobně jako telekomunikační síť. Spoje jsou implementovány buď jako *aktivní optická síť* (AON, active optical network) nebo jako *pasivní optická síť* (PON, passive optical network). AON používá analogii switchů a routerů (umístěných ve skříních na ulicích) a každý zákazník dostává pouze svůj vlastní provoz, nikoliv provoz jiných zákazníků. *Plně dedikované optické linky* jsou další možností (tj. celá cesta je vyhrazena jednomu zákazníkovi), ale tomu odpovídá i cena.

PON neposkytuje možnost switchování provozu (nebo nanejvýš jen v některých uzlech sítě), provoz jednotlivých zákazníků je oddělován jiným způsobem. Používá se časový multiplex TDMA. Tato varianta je nejlevnější.

Obrázek 4.21: FTTx OLT, ONT a ONU⁷

 Jak je výše uvedeno, FTTx se kombinuje s jiným typem technologie. Z toho vyplývá, že je obvykle nezbytné konvertovat optický signál na elektrický (například při cestě do koaxiálu pro DOCSIS nebo kroucené dvojlinky pro xDSL), a podle potřeby překódovat či jinak modifikovat signál. Zařízení, které

⁶Zdroj: <https://broadbandworldforum.wordpress.com/2013/09/23/alternative-access-options-in-building-coax-networks-vs-vdsl-for-ftth/>

⁷Zdroj: https://www.tutorialspoint.com/ftth/ftth_quick_guide.htm

takto zpracovává signál, se nazývá *Optical Network Terminal* (ONT, ITU terminologie pro FTTH a FTTB) nebo *Optical Network Unit* (ONU, IEEE terminologie pro další typy FTTx). Na druhé straně linky (na straně poskytovatele služby) to je *Optical Line Terminal* (OLT).



Úkol

Najděte jakékoliv zařízení GPON OLT (například od Ubiquiti) s nejméně 8 GPON porty, projděte si jeho parametry. Najděte nějaké zařízení FTTH ONT, také si projděte jeho parametry.



Další informace:

- https://www.tutorialspoint.com/ftth/ftth_quick_guide.htm
- <https://get.ospinsight.com/fiber-to-the-home-the-ultimate-guide>
- <https://www.linkedin.com/pulse/abc-pon-understanding-olt-onu-ont-odn-gharbi-skander?trk=seokp-post-cta>
- <https://broadbandworldforum.wordpress.com/2013/09/23/alternative-access-options-in-building-coax-networks-vs-vdsl-for-ftth/>



4.5.5 Obecně k zařízením v přístupových sítích

Pro telekomunikační společnosti může být výhodné mít v MTO místo DSLAMu, OLT, CMTS apod. takové zařízení, které podporuje více než jen jednu z těchto technologií.

 **MSAN** (Multi-Service Access Node) je zařízení instalované v MTO, které poskytuje několik různých typů služeb pro přístupové sítě (různé xDSL, VoIP, FTTx, Ethernet, atd.).

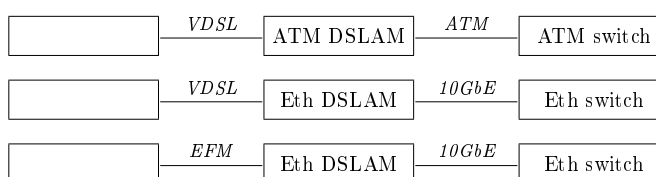
Následující platí pro všechny typy zařízení instalovaných do MTO: vzhledem k tomu, že takové zařízení nebývá vždy „po ruce“, když se vyskytne závada, bývá standardem možnost vzdálené správy.

 Zatímco MSAN, DSLAM, CMTS nebo jiné podobné zařízení je na jednom konci linky přístupové sítě (na straně poskytovatele služby), na druhém konci této linky (na straně zákazníka) je zařízení *CPE* (Customer-Premises Equipment, koncové účastnické zařízení).

4.5.6 EFM

EFM (Ethernet in the First Mile, také ETTH – Ethernet to the Home) (IEEE 802.3ah, rok 2004, roku 2008 přidán do hlavního standardu IEEE 802.3-2008) je možnost využití Ethernetu na celé cestě přes WAN až ke koncovým počítačům. Počítá se jak s metalickými, tak i optickými spoji. Protože 10G a rychlejší Ethernet je použitelný (a také používán) v rozlehlých sítích, bylo by praktické použít Ethernet také na poslední (první) míli připojení k internetu („first mile“).

EFM se v současné době vede přes VDSL linky, případně přes aktivní nebo pasivní optickou síť, samotná síť funguje podobně jako u VDSL. Nicméně posílají se ethernetové rámce. Podobně jako u VDSL, i pro EFM existují speciální zařízení DSLAM, a na straně klienta EFM modemy. EFM



Obrázek 4.22: Postup zavádění EFM

dokáže využívat více fyzických linek najednou, linky se dají slučovat stejně jako u dalších podobných technologií.

Výhodou je, že takto posíláme ethernetové rámce stejně jako v místní síti, a přitom se jedná o dedikovanou (vyhrazenou, nesdílenou) linku. Technologie EFM je chápána jako kompromis: levnější broadbandové technologie znamenají sdílení pásma s jinými uživateli (agregace apod.), kdežto jiné technologie nabízející dedikované spojení jsou dražší.

4.6 Pobočkové ústředny

Víme, že telefonní ústředny jsou uspořádány hierarchicky – na nejvyšších uzlech hierarchie jsou TTO (tranzitní telefonní ústředny), níže jsou UTO (uzlové telefonní ústředny), pod nimi MTO (místní telefonní ústředny), ke kterým se již připojují účastnické přípojky nebo v případě větších firem pobočkové ústředny (PBX).

 *Pobočkové ústředny* se používají k zajišťování vnitřních hovorů a také ke spojování hovorů mezi vnitřní a vnější telekomunikační sítí. Jsou vlastně přípojným bodem do sítě poskytovatele této služby (DTE).

VoIP je ve své podstatě datovou službou, tedy nevyžaduje připojení do telekomunikační sítě (pokud máme jen čistě VoIP ústřednu), stačí připojení do datové sítě. Tomuto typu služeb se budeme věnovat v jedné z následujících kapitol.

Ústředna může být buď přímo hardwarové zařízení, anebo software nainstalovaný na vhodně připojeném počítači (popřípadě virtuální, může být poskytována jako služba). Pobočkovou ústřednu zpravidla připojíme do sítě poskytovatele služby (telekomunikačního operátora), a obvykle i k datové síti.

 **Asterisk** je populární open-source software pro vytváření VoIP ústředn. Konfiguruje se přes webové rozhraní, existuje také několik GUI (nutno zvlášť doinstalovat, např. FreePBX), dále lze k aplikaci přistupovat programově přes API z aplikací napsaných v běžných programovacích jazycích. Asterisk najdeme také v mnohých prodávaných hardwarových ústřednách.



Další informace:

- <http://www.asterisk.org/>, <http://www.asteriskguru.com/tutorials/>
- <http://www.abclinuxu.cz/serialy/asterisk-voip-ustredna>
- <http://www.root.cz/serialy/pobockova-voip-ustredna-asterisk/>
- <http://www.freepbx.org/>



Souhrn kapitoly

V této kapitole jsme se zabývali základy WAN sítí. Ujasnili jsme si, na jakém principu funguje taková síť obecně, uvedli jsme si několik nejznámějších WAN sítí, podívali jsme se na infrastrukturu optických sítí, následovaly telekomunikační sítě, které původně byly základem přístupových sítí. Prošli jsme si nejběžnější technologie současných přístupových sítí (xDSL, DOCSIS, FTTx) a na konci kapitoly jsme krátce probrali technologii pobočkových ústředn.



Bezdrátové a mobilní sítě

 **Rychlý náhled:** V této kapitole se zabýváme bezdrátovými sítěmi různého typu: připomeneme si běžné bezdrátové sítě typu Wi-fi, kde se zaměříme spíše na fyzickou vrstvu, dále si zmíníme možnosti pokrytí rozsáhlejšího prostoru Wi-fi signálem, a následně se zaměříme na mobilní bezdrátové sítě. Na konci kapitoly se podíváme na další typy bezdrátových sítí určené zejména pro Internet věcí.

 **Klíčová slova:** Bezdrátová síť, mobilní síť, rádiové vlny, Wi-fi, IEEE 802.11, AP, frekvenční spektrum, interference, anténa, diverzita, MIMO, MU-MIMO, SSID, BSSID, ESSID, CSMA/CA, beacon rámec, AAA, WEP, WPA, WPA2, WPA3, IEEE 802.1X, Radius, Diameter, Tacacs+, WPS, wireless controller (WLC), Lightweight AP, základnová stanice (BTS, Node B), GSM, CDMA, GPRS, EDGE, UMTS, HSDPA, HSUPA, HSPA, LTE, LTE Advanced, 5G, Internet věcí (IoT), Bluetooth, ZigBee, Z-Wave, NFC, LoRaWAN.

 **Cíle studia:** Cílem této kapitoly je zorientovat se v problematice bezdrátového přenosu dat v různě rozlehlých sítích, umět rozlišit různé druhy bezdrátových technologií podle velikosti pokrytého prostoru, způsobu komunikace, zabezpečení, typického určení a dalších charakteristik.

5.1 Bezdrátové technologie

Jaké přenosové médium tedy u bezdrátových technologií používáme? Mohli bychom říci, že vzduch, více méně je to pravda, ale skutečným nosným médiem pro signál je to, na co signál modulujeme. Jsou tyto možnosti:

- *rádiové vlny* o určité frekvenci – Wi-fi, WiMAX, Bluetooth, NFC, ZigBee apod.,
- *zvukové vlny* (sonická bezdrátová technologie) – například ultrazvuk,
- *světlo* (optická bezdrátová technologie) – laser, infračervené záření (IR), mávání vlajkou...

Nás bude zajímat především první možnost. Sítě založené na přenosu po rádiových vlnách mohou být různě velké – od osobních (PAN, například Bluetooth nebo NFC) přes lokální (LAN, například Wi-fi) a metropolitní (MAN, například WiMAX) až k rozlehlým (WAN, mobilní sítě).

 **WLAN** (Wireless LAN) – touto zkratkou označujeme bezdrátové lokální sítě. Pozor, nepleťte si tuto zkratku s VLAN. taktéž existuje WWAN (Wireless WAN), případně WPAN (Wireless PAN).

 Bezdrátové sítě můžeme také rozdělit podle míry mobility:

- *mobilní* – připojená klientská zařízení i při relativně rychlém pohybu neztratí signál, sem řadíme především mobilní WAN sítě typu GPRS, EDGE, LTE apod.,
- *fixní* – připojená klientská zařízení mohou při rychlejším pohybu ztratit signál, což se stává například u Wi-fi.

Nic není jen černé nebo jen bílé. I pro Wi-fi existují standardy, které tuto technologii přibližují k těm mobilním, ale v základu se pořád jedná o fixní bezdrátovou technologii.

Dále budeme předpokládat, že všichni znají rozdíl mezi vysíláním v bezlicenčním pásmu (např. Wi-fi) a vysíláním v licencovaném pásmu.

5.2 Wi-fi

Technologie Wi-fi (Wireless Fidelity) je bezdrátovou technologií pracující v bezlicenčním pásmu, a to kolem 2,4 GHz, 5 GHz nebo i jiných (podle konkrétního podstandardu).

Základním standardem pro Wi-fi je IEEE 802.11, podstandardy (dodatky) pak k tomuto označení přidávají jedno- či dvoupísmenné zkratky (abeceda je krátká, jednopísmenné už nestačily). O standard a kompatibilitu zařízení (včetně certifikace) se stará Wi-fi Alliance.

IEEE 802.11 implementuje jen spodní podvrstvu L2, tedy MAC podvrstvu. Horní podvrstvu nechává na protokolu LLC (IEEE 802.2), což znamená, že používáme LLC rámce (jak bylo naznačeno v druhé kapitole od strany 33), které zapouzdřujeme do MAC rámce podle IEEE 802.11.

5.2.1 Access point a WDS

 Aktivním síťovým prvkem je *access point* (AP, přístupový bod). V základu toto zařízení pracuje na vrstvě L2, tedy protokol IEEE 802.11 se svými podstandardy implementuje vrstvy L1 a L2 (jako IEEE 802.3). Počítá se s podsunutím pod TCP/IP či jiný síťový zásobník. Nicméně, reálná zařízení mohou implementovat i jiné vrstvy, nebo naopak mohou být poněkud osekanejší. Například zařízení označované jako Wi-fi router přidává k funkcionalitě AP ještě funkcionalitu vrstvy L3, naopak Wi-fi repeater (extender) se zaměřuje jen na vrstvu L1 a pracuje pouze se signálem.

 AP vysílá signál, oblast pokrytou tímto signálem nazýváme *buňka* nebo také *základní oblast služeb* (BSA – Basic Service Area). Pokud propojíme víc AP do distribučního systému (WDS), pak oblast pokrytou signálem zapojených AP nazýváme *rozšířená oblast služeb* (ESA – Extended Service Area).

Distribuční systém (zde WDS, *Wireless Distribution System*) tedy funguje jako páteří síť propojující přístupové body. Obecně lze použít distribuční systém nejen pro propojení (bezdrátových) přístupových bodů, ale také třeba pro propojení dvou metalických lokálních sítí. Nejjednodušší distribuční systém lze vytvořit tak, že přístupový bod jedné BSA pracuje jako běžná stanice v BSA druhého přístupového bodu (opakováč), ale ne každý přístupový bod podporuje režim opakováče.

5.2.2 Fyzická vrstva

Dále budeme předpokládat, že všichni studenti zvládají technologii Wi-fi (různé standardy, práce se signálem apod.), tedy uvádíme především přehledovou tabulku (str. 113).

Specifikace	Frekvence	Max. propustnost	Ozn.
IEEE 802.11	2,4 GHz	2 Mb/s	
IEEE 802.11a	5 GHz	54 Mb/s	Wi-Fi 1
IEEE 802.11b	2,4 GHz	11 Mb/s	Wi-Fi 2
IEEE 802.11g	2,4 GHz	54 Mb/s	Wi-Fi 3
IEEE 802.11n	2,4 GHz, 5 GHz	250 Mb/s na anténu, celkem až 600 Mb/s	Wi-Fi 4
IEEE 802.11ac	5 GHz	433 Mb/s na anténu a stream, až 3,5 Gb/s	Wi-Fi 5
IEEE 802.11ad	60 MHz	cca 7 Gb/s, podle použitých kanálů	WiGig
IEEE 802.11ax	2,4 GHz, 5 GHz	9,6 Gb/s	Wi-Fi 6
IEEE 802.11ax-2021	2,4 GHz, 5 GHz, 6 GHz	9,6 Gb/s	Wi-Fi 6E
IEEE 802.11be	2,4 GHz, 5 GHz, 6 GHz	46,1 Gb/s	Wi-Fi 7

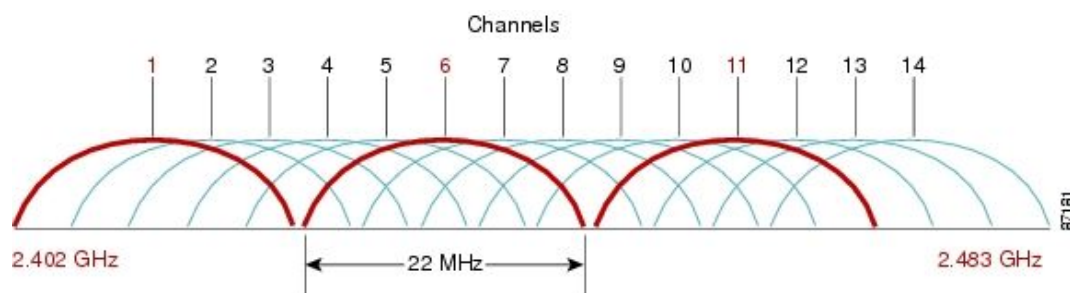
Tabulka 5.1: Podstandardy pro fyzickou vrstvu Wi-fi

**Další informace:**

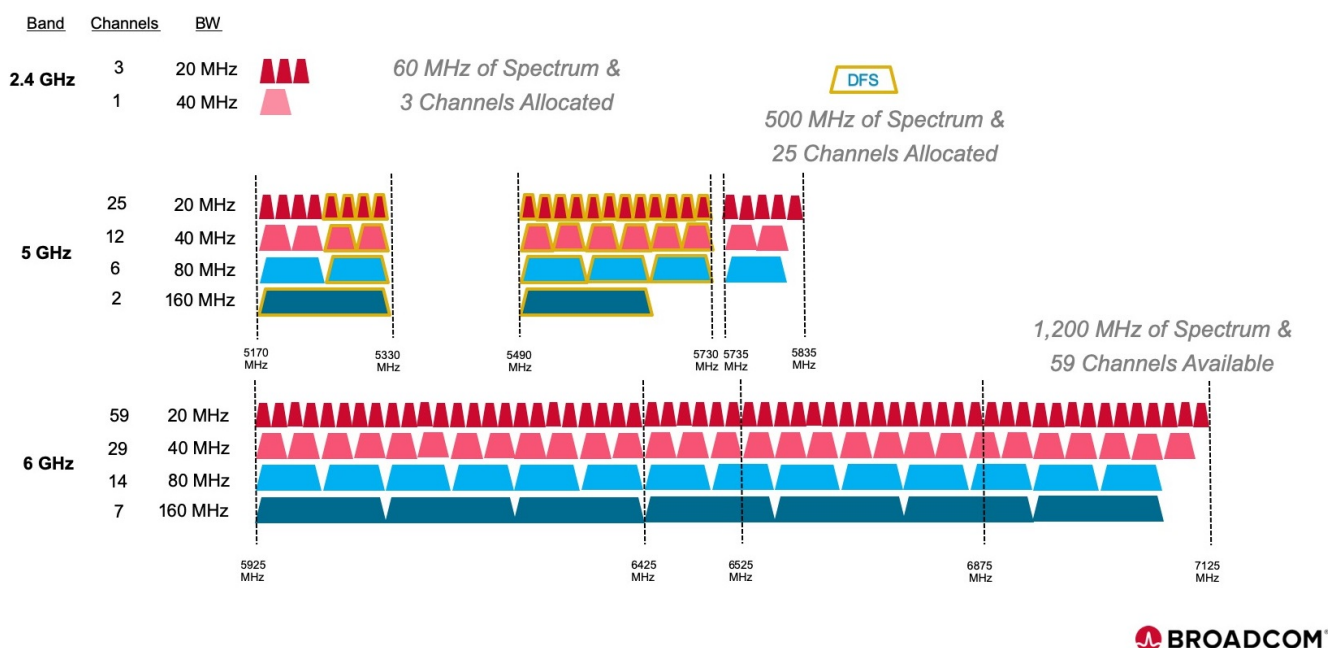
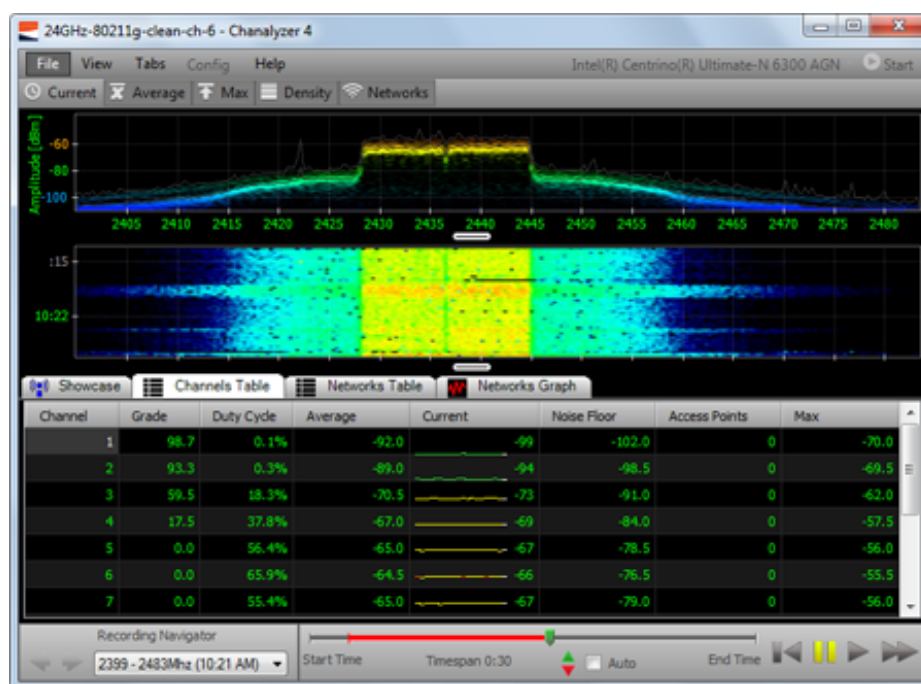
- <https://www.wi-fi.org/discover-wi-fi> (v menu vpravo jsou různé standardy a další témata)
- <https://www.cisco.com/c/en/us/products/collateral/wireless/white-paper-c11-740788.html> (srovnání IEEE 802.11ax a IEEE 802.11ac)
- <https://www.asus.com/cz/support/faq/1051272/> (Představení Wi-fi 7)
- <https://www.techtarget.com/searchnetworking/tip/Whats-the-difference-between-80211ax-vs-80211ac>

**5.2.3 Interference**

Interference (rušení) je problém hlavně v pásmu kolem 2,4 GHz. Ve stejném frekvenčním spektru pracují také bluetooth zařízení, mikrovlnné trouby (pokud si myslíte, že při provozu je mikrovlnka tak důkladně izolovaná, že ven se žádný signál nedostane, pak vezte, že jste na omylu) a většinou i dětské chůvičky a různá další zařízení (je nelicencované, a zrovna oblíbená frekvence), takže vzájemné rušení se ani zdaleka netýká pouze Wi-fi zařízení.

Obrázek 5.1: Frekvenční spektrum pro IEEE 802.11b¹

¹Zdroj: <http://www.cisco.com/en/US/docs/solutions/Enterprise/Mobility/emob30dg/RFDesign.html>

Obrázek 5.2: Frekvenční spektra používaná různými generacemi IEEE 802.11, nepřekrývající se kanály²Obrázek 5.3: Wi-Spy Chanalyzer³

Teoreticky by bylo řešením navolit na blízkých AP kanály tak, aby se dotyčné pětice co nejméně překrývaly (tj. 2 až 3 komunikující AP, které se „vidí“ a přitom se navzájem neruší), ale to je opravdu jen teorie. Ve skutečnosti (díky fyzikálním zákonům) běžná komunikace na kanálu „zamožuje“ nejen dva okolní kanály na každé straně, ale bohužel i vzdálenější, i když v menší míře. Na obrázku 5.1 vidíme tři

²Zdroj: <https://jp.broadcom.com/info/wifi6e>

³Zdroj: <http://www.wi-spy.com.au/>

zvýrazněné obloučky, které se nepřekrývají (1., 6. a 11. kanál). V reálu je spodní část obloučků značně roztáhnuta, šum zasahuje mnohem dále („do ztracena“), jak vidíme na obrázku 5.3 (výstup z aplikace Wi-Spy Chanalyzer 4 pro jeden zdroj signálu podle IEEE 802.11g, v horní části si všimněte amplitudy v různých kanálech, kanály jsou označeny frekvencí v MHz).

Tj. „vhodné“ rozdělení kanálů prakticky neexistuje, dva AP se budou navzájem rušit, i když budou pracovat na „dostatečně vzdálených“ kanálech. Pokud tedy nemáme jinou možnost, lze dva AP umístit tak, aby se jejich dosahy prolínaly, kanály zvolíme co nejdál od sebe (nicméně automaticky se to tak obvykle také děje), ovšem musíme počítat s jistým snížením rychlosti. Je zajímavé, že za určitých okolností je paradoxně lepší pro více AP zvolit tentýž nebo hodně blízký kanál.

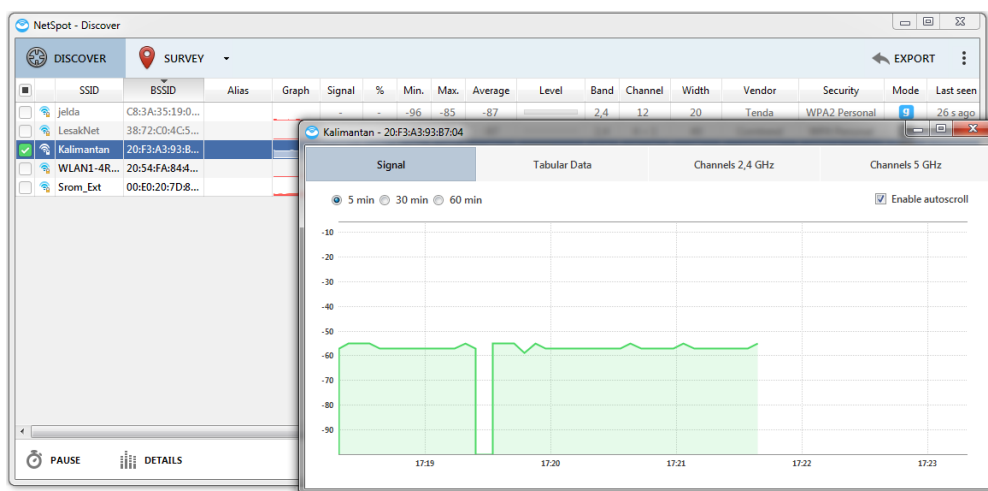
Výše popsané rušení je také jedním z důvodů, proč při použití WDS musíme počítat s nižší propustností sítě. Wi-fi síť může být rušena také z jiných zdrojů (čímkoliv na blízké frekvenci), například mikrovlnnou troubou, naštěstí tento zdroj rušení nebývá používán dlouhodobě.

5.2.4 Diagnostické nástroje pro fyzickou vrstvu

✂ Aplikace *inSSIDer* pro Windows a MacOS od společnosti MetaGeek je klasikou, ale má zdatnou konkurenci s velmi podobným rozhraním.

✂ Velice se InSSIDeru podobá nástroj *Acrylic WiFi*, který ve volně stažitelné verzi Home (pouze pro domácí použití) umí v podstatě totéž a také jeho rozhraní je podobné.

✂ Velmi nadějně vypadá volně šiřitelná aplikace *NetSpot* pro Windows a MacOS, která má více funkcí než InSSIDer (například mapování pokrytí signálem v oblasti či hledání problémů v síti).



Obrázek 5.4: Netspot

✂ Program *WifiInfoView* od Nirsoftu je volně šiřitelný nástroj pro zjišťování informací o dostupných bezdrátových sítích. Dá se spustit buď jako aplikace s oknem (vidíme seznam dostupných sítí, jejich vlastnosti a ve spodním podokně podrobnosti), tak i s různými dodatečnými parametry (na webu projektu jsou podrobné informace).

✂ Pro Linux existují obdobné nástroje jako InSSIDer, například *LinSSIDer* a *iwScanner*. Pro Android lze stáhnout aplikaci *Wifi Analyzer*. Sice se nedostaneme k tolika informacím jako obdobné aplikace pro

„nemobilní“ operační systémy (už z principu, to by ten Android musel být rootnutý), ale pro představu o frekvenčním spektru to bohatě stačí.

 Existují i další nástroje, které umožňují zkoumat frekvenční spektrum. Většinou jde o komerční nástroje anebo sice volně šiřitelné, ale vyžadující kompatibilní hardware (tj. záleží na štěstí, jestli takový máme). Z nejznámějších:

- *NetStumbler* – volně šiřitelný analyzátor spektra,
- *Wi-Spy* – komerční nástroj (analyzátor spektra včetně rozdělení na jednotlivé kanály, set se skládá z upravené Wi-fi síťové karty do USB konektoru a speciálního softwaru),
- *Chanalyzer* je sice komerční nástroj (jako součást Wi-Spy), ale existuje volně šiřitelná varianta; je to výkonný spektrální analyzátor (viz obrázek 5.3),
- *AirMagnet Wi-Fi Analyzer* je komerční nástroj funkcností podobný nástroji Wi-Spy, taktéž spektrální analyzátor (také vyžaduje speciální hardware),
- *Xirrus Wi-fi Inspector* – volně šiřitelný nástroj na zjišťování dostupných AP a jejich vlastností (firma Xirrus se stala součástí jiné firmy a tento nástroj už není na jejich stránkách, nicméně na download serverech je ještě dostupný),
- *Iperf* (taktéž volně šiřitelný, ovládá se v textovém shellu) slouží k ladění parametrů sítě a je to užitečný také jako pomocný nástroj při diagnostice sítě, dokáže odchytávat pakety (sniffer) a generovat „přirozený provoz“ na síti (užitečné při testování),
- *Jperf* je volně šiřitelná grafická nástavba programu Iperf (tj. většina lidí si s Iperfem instaluje Jperf).

Existují nástroje určené vysloveně pro mapování signálu, například *EKahau Heatmapper*, *EKahau Site Survey*, *Meraki Wi-fi Mapper*, *Aerohive Free Wi-fi Planner*, *Cisco Clean Air*, *VisiWave Site Survey*, *RF3D WifiPlanner*.



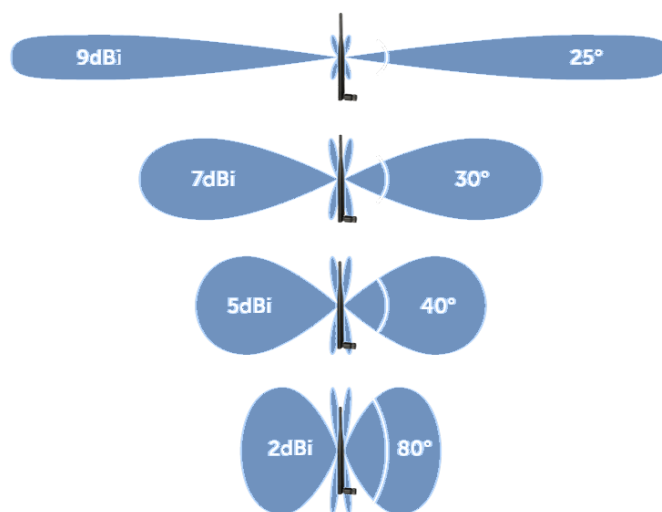
Další informace:

- <https://www.metageek.com/inssider/>
- <https://www.acrylicwifi.com/en/wlan-software/wlan-scanner-acrylic-wifi-free/>
- <https://www.netspotapp.com/>
- http://www.nirsoft.net/utils/wifi_information_view.html
- <http://alternativeto.net/software/inssider/>
- <http://www.stumbler.net>
- <https://shop.metageek.com/products/chanalyzer-essential/>
- Airagnet Wifi analyzer: <https://www.netally.com/products/airmagnet-wifi-analyzer/>
- <https://sourceforge.net/projects/iperf/>, <https://sourceforge.net/projects/jperf/>



5.2.5 Signál a antény

 Existují různé druhy antén – všesměrové s různými způsoby formování signálu či směrové určené pro point-to-point spoje. Antény mohou být buď *interní* (například v notebooku jsou antény obvykle vedeny podél displeje) nebo *externí* (vyměnitelné jsou obvykle připojeny koaxiálovým konektorem typu N nebo BNC).

Obrázek 5.5: Pokrytí signálem u všesměrové antény⁴

U běžných malých všesměrových antén je signál nejsilnější v rovině kolmé na osu antény, tedy nakloněním antény můžeme například ovlivnit to, zda bude signál dostatečně silný v okolních podlažích budovy. Pokud chceme mít signál jen v rámci jednoho patra, necháme anténu směřovat nahoru nebo dolů, kdežto když má signál dosáhnout do horního nebo spodního patra, nakloníme anténu do úhlu 45 stupňů. Pozor, ve směru, do kterého anténa míří, je signál nejslabší!

 Zařízení s jedinou anténou může fungovat jen v *polovičním duplexu*, tedy buď vysílat nebo přijímat, ale nikoliv obojí zároveň. Pokud ale má alespoň dvě antény, může pro každý směr použít jednu z nich, přičemž budou pracovat na různých kanálech, čímž implementujeme *plný duplex*.

 Ve specifikaci IEEE 802.11a/b/g se setkáváme s *diverzitou*, tedy možností využití více antén pro frekvenční pásmo 2,4 GHz. Diverzita funguje tak, že při komunikaci s konkrétním zařízením se nejdříve přijímá všemi, a po vyhodnocení kvality přijatého signálu ze všech antén je pak určena jedna, která bude pro dané zařízení využívána.

 Ve specifikaci pro IEEE 802.11n se objevuje nástupce diverzity – *MIMO* (Multiple Input, Multiple Output) [maimo] – určující využití více antén a algoritmus pro kombinování signálu z těchto antén. Na rozdíl od diverzity mohou (téměř) všechny antény fungovat zároveň.

Při použití MIMO jsou antény rozděleny na Tx (odesílající) a Rx (přijímající), abychom zajistili komunikaci v plném duplexu, a dalším parametrem je počet paralelních streamů (tedy kolik nezávislých kanálů se nám vejde do spektra). Zapisuje se takto:

$$\begin{array}{ccccc} \text{Tx} & & \text{Rx} & & \text{streamy} \\ 3 & \times & 3 & : & 2 \end{array}$$

což znamená 3 odesílající antény, 3 přijímající antény a 2 streamy.

Technologie MIMO má jednu slabinu – v jednom okamžiku může AP komunikovat maximálně s jedním klientem, třebaže více anténami.

 To řeší technologie *MU-MIMO* (Multi-User MIMO) používaná ve specifikaci IEEE 802.11ac, a také například u metropolitních sítí WiMAX. V jednom okamžiku může AP paralelně komunikovat s více

⁴Zdroj: <https://insmart.cz/zlepsit-wifi-signal-domaci-siti/>

různými klienty (maximum je 8 streamů, pro každého klienta 2, tedy maximálně 4 klienti).

V implementaci je bohužel trochu zpoždění, tedy když kupujeme zařízení podle IEEE 802.11ac podporující MU-MIMO, může jít o jednu ze tří fází, obvykle druhou:

- Wave 1 – SU-MIMO (Single-User MIMO), tedy žádná paralelní komunikace, max. 3 streamy,
- Wave 2 – MU-MIMO (více klientů paralelně), max. 3–4 streamy,
- Wave 3 – max. 8 streamů (s tím se obvykle u Wi-Fi 5 nesetkáme).

 V IEEE 802.11ax (Wi-Fi 6) se přechází z multiplexování OFDM na OFDMA. Co to znamená? Ve skutečnosti se liší i význam „konců“ těchto zkratk: OFDM je Orthogonal Frequency Division Multiplexing, OFDMA je Orthogonal Frequency-Division Multiple Access. Při OFDM je sice možné používat MU-MIMO, ale pouze ve formě, kdy různé kanály jsou přiřazeny různým koncovým zařízením. OFDMA jde dál: kanály dělí na subkanály (nazývají se RU = Resource Unit). Také v rámci subkanálů se používá ortogonální dělení, aby se jednotlivé subkanály daly jednoduše oddělit.



Další informace:

<https://www.extremenetworks.com/extreme-networks-blog/ofdm-and-ofdma-subcarriers-what-are-the-differences/>



5.2.6 Identifikátory a služby

 Klientská zařízení potřebují identifikovat síť, do které se připojují nebo v ní pracují. Od toho máme tyto *identifikátory*:

- *SSID* (Service Set ID) je identifikátor celé sítě, název sítě. Je to řetězec o délce maximálně 32 znaků, který musí znát každé zařízení, které se do sítě přihlašuje.
- *BSSID* (Basic Service Set ID) je identifikátor access pointu. Obvykle se jedná o MAC adresu tohoto AP, tedy délka je 6 oktetů.
- *ESSID* (Extended Service Set ID) je identifikátor WDS, pokud máme víc AP propojených do distribučního systému.

AP může poskytovat více SSID, přičemž každé z nich představuje jednu bezdrátovou síť. Například můžeme mít zvlášť SSID pro hosty, to se hodí nejen ve firmách, ale i v domácnosti. Ve firmách se na různá SSID napojují různé VLAN (čímž se také určují přístupová oprávnění těch, kdo v dané bezdrátové síti s konkrétním SSID komunikují), v domácnosti se na routeru obvykle také dá nastavit alespoň stupeň důvěryhodnosti.



Definice (Beacon rámeček)

Každý AP v pravidelných intervalech vysílá *Beacon rámeček* („majákový“ rámeček), ve kterém informuje o svých parametrech. Kromě jiného je v Beacon rámci řetězec SSID potřebný pro komunikaci s AP.



 AP poskytuje v bezdrátové síti tyto služby:

- *autentizace* – rozhoduje, zda do své buňky a tím i do sítě pustí konkrétní koncové zařízení,
- *asociace* (přidružení) – pokud je koncovému zařízení povolen přístup, je třeba vytvořit vazbu mezi AP a tímto zařízením, tedy technicky zajistit zařízení možnost komunikace v buňce,
- *de-asociace* – zrušení této vazby.

5.2.7 Přístupová metoda

Standard IEEE 802.11 definuje také přístupovou (kolizní) metodu, která je používána na podvrstvě MAC vrstvy L2. Na rozdíl od Ethernetu je v bezdrátových sítích používána neustále, protože při sdíleném médiu se nikdy nedá zcela vyhnout kolizím. Navíc se používají mechanismy snižující na minimum riziko kolizí a řešící problém skrytých stanic.

 *Problém skrytých stanic* nastává tehdy, když je buňka rozsáhlejší a dvě stanice patřící do téže buňky se navzájem „nevidí“ (ke každé z nich sice dosáhne signál z centrálního AP, ale nikoliv jejich signály navzájem). To je vcelku běžné, stanice mohou být od sebe příliš vzdáleny nebo navzájem zastíněny nábytkem či zdí, signál koncových zařízení taky bývá typicky slabší než signál AP.

Definice (Přístupová metoda CSMA/CA a mechanismus RTS/CTS)

Podle IEEE 802.11 se používá přístupová metoda *CSMA/CA*:

- CS (Carrier Sense) – nasloucháme na nosné,
- MA (Multiple Access) – vícenásobný přístup (sdílené přenosové médium),
- CA (Collision Avoidance) – kolizím se vyhýbáme, tedy jim předcházíme (nestačí je jen detekovat).

Každé zařízení, které chce vysílat, musí nejdřív požádat AP o povolení k vysílání (to je krátký správný rámec) a až ho získá, může odesílat data. CSMA/CA se dá implementovat více různými způsoby, jeden z mechanismů je například *RTS/CTS* (Request to Send, Clear to Send):

- DTE vyšle signál RTS (v krátkém řídicím rámci), čímž požádá o povolení k vysílání dat,
- pokud je možné vysílání povolit, AP odešle žádajícímu DTE signál CTS (v řídicím rámci),
- DTE může vyslat datový rámec (MAC podle IEEE 802.11 s LLC rámcem, v něm jsou data),
- zpět je doručeno potvrzení (ACK), také v řídicím rámci.



Jak vidíme, ve Wi-fi se v takovém případě používá potvrzovaný přenos bez navázání spojení (na spoji).

Mechanismus RTS/CTS víceméně řeší problém skrytých stanic – počet kolizí snižuje na minimum tím, že u krátkých řídicích rámců je mnohem menší pravděpodobnost kolize než u velkých datových, a AP taky nepošle signál CTS s povolením vysílání, pokud zrovna na dané frekvenci komunikuje jiný (žadateli skrytý) uzel v síti.

I přes takto určenou přístupovou metodu ke kolizím může docházet, protože přenosové médium je sdílené. Dochází k nim například tehdy, pokud dvě zařízení zároveň vyšlou signál RTS nebo se takto potkají rámce s RTS a začátkem vysílání datového rámce.

Metoda CSMA/CA je nedeterministická.

5.2.8 Wi-fi rámce

 Rozlišujeme tři typy MAC rámců podle IEEE 802.11:

- *datové rámce* (Data Frame) – obsahují payload v LLC rámci,
- *řídicí rámce* (Control Frame) – podtypy jsou RTS/CTS pro vyjednání povolení k vysílání, a ACK pro potvrzení přijatých rámců,
- *rámce pro správu* (Management Frame) – sem řadíme
 - beacon rámec (AP se ohlašuje),

- probe, probe response (zjišťování přítomnosti uzlů sítě a jejich možností),
- association request, association response (žádost o asociaci od stanice, odpověď),
- re-association request, response (přechod k jinému AP ve stejné ESS),
- disassociation (ukončení spojení k AP),
- authentication (žádost o autentizaci uzlu), de-authentication.

Uvnitř datového rámce najdeme LLC rámec (IEEE 802.2), ve kterém je většinou zapouzdřen IP paket. Nicméně například ve Wiresharku nemáme obvykle šanci tento fakt zjistit, protože vše za MAC záhlavím je šifrováno (včetně LLC záhlaví).

 Než se dostaneme k formě samotného záhlaví rámce, podíváme se na práci s adresami uvedenými v rámci. Údaje o adresách se v záhlaví mění podle toho, kterou částí Wi-fi sítě právě rámec prochází:

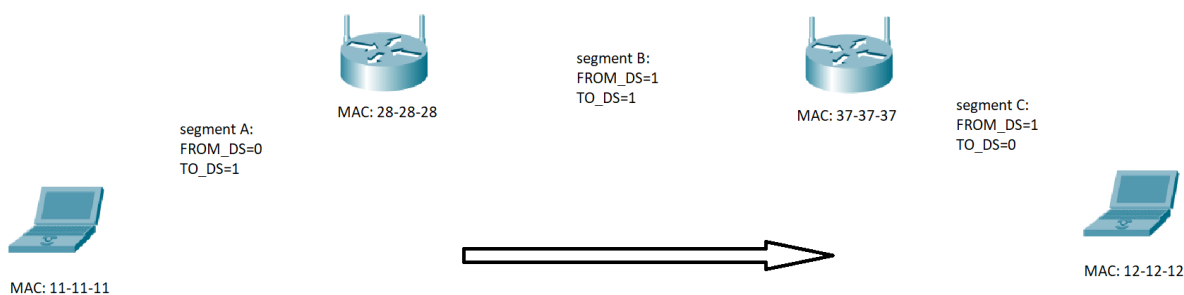
1. Jsme v ad-hoc síti, tedy spoje vedou vždy právě mezi dvěma klientskými zařízeními, žádný AP na cestě není (nebo se jedná o rámec, který v principu míří jen k sousedovi a ne dále).
2. Jsme v síti typu infrastruktura, rámec je na spoji mezi zdrojem (klientem) a některým AP.
3. Jsme v síti typu infrastruktura, rámec je na spoji mezi dvěma AP.
4. Jsme v síti typu infrastruktura, rámec je na spoji mezi některým AP a cílem rámce.

Takže záleží, na jakém typu spoje se právě rámec nachází, resp. mezi jakými dvěma typy uzlů v síti právě prochází (kdo je momentálním odesílatelem a příjemcem).

 V záhlaví rámce máme kromě samotných adres tyto dva příznaky (tj. jednobitové hodnoty):

- **FromDS** – je nastaven na 1, pokud na daném spoji je momentálním odesílatelem AP (resp. DCE), a je nastaven na 0, pokud je odesílatelem klientské zařízení (DTE).
- **ToDS** – je nastaven na 1, pokud na daném spoji je momentálním příjemcem AP (DCE), a je nastaven na 0, pokud je příjemcem DTE.

Zkratka „DS“ v názvech příznaků je *distribuční systém*.



Obrázek 5.6: Příznaky ve Wi-fi rámci

 Z toho vyplývá, že v nastíněných čtyřech případech jsou tyto příznaky nastaveny takto:

1. V ad-hoc síti jsou oba příznaky vždy nastaveny na 0, protože žádné AP na cestě ani být nemohou.
2. Na spoji mezi zdrojem rámce (DTE) a nejbližším AP je **FromDS=0** (odesílatel je DTE), **ToDS=1** (momentální příjemce je AP).
3. Na spoji mezi dvěma AP jsou oba příznaky nastaveny na 1: **FromDS=1**, **ToDS=1**. Na obou koncích spoje jsou AP, a tedy jsme uvnitř distribučního systému.
4. Na spoji mezi některým AP a cílem rámce (DTE), tedy v závěru cesty, je **FromDS=1**, **ToDS=0**.

 Nyní se zaměříme na *adresy*. V ethernetovém rámci máme vždy dvě adresy – MAC adresu cíle a MAC adresu zdroje. Ve Wi-fi rámci sice také používáme MAC adresy, ale jejich počet se dynamicky mění podle momentální pozice v síti, podobně jako zmíněné dva příznaky. Takže postupně:

1. V ad-hoc síti je to stejně jako u Ethernetu, máme MAC adresu cíle a MAC adresu zdroje, nic jiného nepotřebujeme.
2. Na spoji mezi zdrojem rámce (DTE) a nejbližším AP (tedy v první fázi cesty) jsou v záhlaví tři MAC adresy:
 - adresa momentálního příjemce na spoji (AP),
 - adresa momentálního odesílatele na spoji (zdroje rámce, DTE),
 - adresa finálního cíle rámce (DTE).

Všimněte si, že první dvě adresy se vztahují k právě aktuálnímu spoji, přičemž pořadí je zachovááno – nejdřív příjemce, pak odesílatel.

3. Na spoji mezi dvěma AP jsou v záhlaví dokonce čtyři adresy:
 - adresa momentálního příjemce na spoji (cílového AP),
 - adresa momentálního odesílatele na spoji (zdrojového AP),
 - adresa finálního cíle rámce (DTE),
 - adresa původního zdroje rámce (DTE).

Opět se první dvě adresy vztahují k aktuálnímu spoji, ale samozřejmě musí být v záhlaví uloženo i to, kdo je zdrojem a cílem rámce. V první i druhé dvojici je zachovááno dané pořadí.

4. Na spoji mezi některým AP a cílem rámce (DTE), tedy v závěru cesty, jsou v záhlaví opět „pouze“ tři adresy:
 - adresa momentálního příjemce na spoji (cíle rámce, DTE),
 - adresa momentálního odesílatele na spoji (AP),
 - adresa původního zdroje rámce (DTE).



Poznámka:

BSSID konkrétního AP je MAC adresa tohoto AP. Každé zařízení v buňce musí znát BSSID svého AP, protože ji umísťuje do záhlaví odesílaných rámců. Koncové zařízení ve skutečnosti přímo komunikuje pouze s jediným zařízením – AP v buňce, ke kterému je asociováno.



Struktura rámce se liší podle jeho typu.

 *Záhlaví datového rámce* je dlouhé 30 oktetů, což je poněkud více než u Ethernetu (navíc se také přidává zápatí s kontrolním součtem, stejně jako u Ethernetu). Najdeme tam následující položky (seznam je mírně zjednodušen):

- Frame Control (1 oktet, pole pro řízení) – bity tohoto pole se dělí do několika skupin:
 - version – verze,
 - Type – typ rámce (Data, Control, Management),
 - Subtype – podtyp, například u typu Management máme podtypy Beacon, Probe, atd.,
- Flags (1 oktet, příznaky), například:
 - příznaky ToDS, FromDS,

- příznak More Fragments – Wi-fi rámce mohou být fragmentovány podobně jako IP pakety, tento příznak slouží stejně jako příznak MF v záhlaví IP paketu (viz kapitolu o síťové vrstvě),
- Retry – tento příznak se nastaví, pokud se jedná o opakovaný přenos téhož rámce, atd.
- 4 adresy ve vhodném pořadí (Receiver, Destination, Transmitter, Source),
- BSSID platné v dané buňce (tj. MAC adresa AP),
- Sequence Number, Fragment Number – význam je podobný jako u obdobného pole v TCP segmentu, resp. IP paketu.

V záhlaví datového rámce toho ve skutečnosti najdeme více, toto byl jen „promazaný“ souhrn.

 *Řídící rámec* (Control Frame) víceméně zachovává strukturu, jakou jsme viděli u datového rámce, ale implementuje jen některá pole.

První pole samozřejmě používá (z něj se dá zjistit, že to je řídicí rámec a jeho podtyp – RTS, CTS nebo ACK), druhé pole se v něm sice také nachází, ale obvykle je celé nastavené na 0 (včetně příznaků ToDS a FromDS, protože je určen jen pro zařízení v buňce odesílajícího AP, přes žádné další zařízení nejde, nanejvýš může být nastaven příznak Retry).

Třebaže se to může zdát zvláštní, z adres tam najdeme buď dvě nebo pouze jedinou.

- U RTS rámce (žádost o povolení vysílání) jsou dvě adresy – cílový AP a žadatel.
- U CTS rámce (povolení vysílání) je jen jedna adresa – Receiver Address (příjemce, tj. zařízení, které žádalo o povolení), protože odesílající AP je „jasný“ (zařízení patří vždy jen do jedné buňky, buňka má jen jeden AP). Takže pokud klient odeslal žádost RTS, čeká na CTS obsahující svou adresu.
- U ACK rámce (AP potvrzuje koncové stanici doručení rámce) taktéž najdeme jen adresu klienta (Receiver Address). Pokud tedy po „povolovacím kolečku“ klient odeslal datový rámec, čeká na ACK rámec se svou adresou. Jestliže čeká marně, bude nutné datový rámec přenést znovu s příznakem Retry, třeba i několikrát.

Dále už nic nenásleduje (kromě kontrolního součtu), ani BSSID.

 *Správní rámec* (Management Frame) má celkovou strukturu rámce (vlastně rámců) naopak mnohem složitější. Ve skutečnosti se jedná o dva do sebe vnořené rámce:

- na podvrstvě LLC je to rámec IEEE 802.11 Wireless LAN Management Frame (místo 802.2),
- na podvrstvě MAC je to rámec IEEE 802.11 typu Management Frame, který v sobě zapouzdřuje ten z podvrstvy LLC.

Ve „vnějším“ rámci je struktura jako v datovém – nejdřív pole pro verzi, typ a podtyp tak, jak bylo dříve uvedeno, následuje pole příznaků (obvykle nulové, případně může být nastaven příznak Retry, pokud je nutné rámec posílat opakovaně). Další pole jsou stejná jako u datového – čtyři adresy, BSSID, Sequence Number, Fragment Number. V některých podtypech mohou být některé adresy broadcastové, dokonce i obě první adresy (to by u Ethernetu nešlo).

Za záhlavím vnějšího rámce následuje vnitřní rámec z podvrstvy LLC. Rámec odesílaný access pointem obsahuje:

- časové razítko (aby bylo jasné, kdy byl rámec odeslán, také kvůli synchronizaci),
- hodnotu beacon interval (jak často se posílají beacon rámce),
- pole příznaků (jednotlivé bity jsou nastaveny podle toho, co AP podporuje – typy multiplexování, zabezpečení apod.),

- sekvenci tagů (položek, kde u každé je typ položky, délka a konkrétní hodnota), podle toho, o jaký podtyp rámce se jedná.

Například u Beacon rámce nebo Probe Response (kde je odesílatelem AP) jsou tagy pro SSID, podporované rychlosti, seznam podporovaných šifrovacích algoritmů, apod.

V rámci Probe Request (odesílaném stanicí žádající o přidružení do buňky) je pouze pole tagů.



Poznámka:

Ke zkoušce nebude nutné znát podrobně všechna pole, je jen dobré mít hrubou představu o tom, co se v záhlavích rámců nachází. Například – pokud „skryjeme SSID“ (tedy zakážeme vysílání Beacon rámců), bude SSID snadno odposlechnutelné z rámců Probe Request a Probe Response.



Poznámka:

Ve Wi-fi síti lze místo rámců podle IEEE 802.11 používat rámce typu Ethernet II, stejně jako v ethernetových sítích. Proto se nedivte, že když zachytáváte provoz na bezdrátové síti, objevují se rámce Ethernet II. Co se týče kolizí, v případě vícestreamové komunikace (která je běžná jak ve firmách, tak i v mnohých domácnostech), to není až takový problém.

Jen bohužel rámce Ethernet II neumí šifrovat – u provozu šifrovaného na vyšších vrstvách (HTTPS nebo cokoliv jiného poskládaného s TLS či IPSec) to opět není problém, jenže část provozu není šifrovaná na vyšších vrstvách (například DNS překlady, pokud se nepoužívá nějaká forma zabezpečení DNS). Nicméně při oddělení streamů není velká pravděpodobnost odposlechu.



Úkol

Ve Wiresharku zkuste odposlechnout provoz Wi-fi sítě (jiné než Eduroam) v době, kdy se do buňky bude asociovat některý (jiný) klient. Pokud najdete rámce podle IEEE 802.11, prozkoumejte Beacon rámec, rámce Probe Request a Probe Response, některý datový rámec. Všimněte si rozdílů v záhlavích a také umístění SSID. Pokud váš AP komunikuje s využitím rámců Ethernet II, najdete nějaký vhodný provoz například na <https://wiki.wireshark.org/SampleCaptures> (tam najdete sekci o Wifi).



Další informace:

- <https://witestlab.poly.edu/blog/802-11-wireless-lan-2/>
- Různé L2 LAN rámce: https://flylib.com/books/en/2.221.1/local_area_network_lan_technologies.html



 **Diagnostické nástroje pro MAC podvrstvu** jsou především sniffery, tj. programy, které odchytávají pakety, umožňují je analyzovat či vizualizovat a případně nabízejí další doprovodné funkce. Vpodstatě se tyto nástroje používají souhrnně pro drátové i bezdrátové sítě. Síťové rozhraní, které je takto analyzováno, by mělo běžet v *promiskuitním módu*, to znamená, že by mělo přijímat všechny rámce včetně těch, ve kterých je cílová adresa jiná (určených jinému uzlu v síti). V případě nástrojů pro bezdrátové sítě se také setkáváme s nástroji pro „prolamování“ zapomenutých hesel/WEP, apod.

Nejznámějším a nejpopulárnějším volně šiřitelným snifferem je *Wireshark* (původně *Ethereal*). Taktéž oblíbený volně šiřitelný je *Kismet*⁵ – detektor sítí, sniffer a potenciálně IDS (budeme se učit později). Další jsou komerční *Eye P.A.* a *Cascade Pilot*.

5.3 Zabezpečení bezdrátové komunikace

5.3.1 AAA

 AAA (Authentication, Authorization, Accounting) je termín označující řízení přístupu, v našem případě přístupu do bezdrátové sítě. Jak zkratka napovídá, zahrnuje tři fáze:

- *Autentizace* (Authentication) – klient (a v některých případech i AP) má prokázat svou totožnost, jde tedy o kontrolu identity. Přístupové údaje se také označují anglickým termínem *credentials*.
- *Autorizace* (Authorization) – když je identita klienta zjištěna, je třeba mu přidělit přístupová oprávnění, vymezit povolenou činnost.
- *Účtování* (Accounting) – evidujeme stanovené činnosti klienta v síti (podle konfigurace), případně reagujeme při pokusech o překročení oprávnění přidělených ve fázi autorizace.

Každá z následujících možností zajišťuje šifrování jak v první fázi komunikace (autentizace a asociace), tak i při běžném provozu. Ovšem používají se různé postupy a různé (i vzhledem k bezpečnosti) šifrovací algoritmy.

 Krátká poznámka k šifrování. U jakéhokoli šifrování potřebujeme *šifrovací klíč*, který mají mít obě strany komunikace (odesílatel pro šifrování a příjemce pro dešifrování).

Existují dva základní druhy šifer – *symetrické* a *asymetrické*. Vzhledem k šifrovacím klíčům:

- Symetrické šifry používají tentýž klíč pro šifrování i dešifrování. Hlavním problémem (a taky nevýhodou) je, jak *bezpečně* dostat klíč druhému komunikujícímu.
- Asymetrické šifry používají pro každý směr komunikace pár klíčů – soukromý a veřejný, přičemž co zašifruji veřejným klíčem, musí být dešifrováno soukromým (a naopak – ve funkčnosti jsou zaměnitelné). Uživatel vygeneruje svůj pár klíčů, soukromý si uloží a veřejný předá (jakkoliv) druhé straně. Druhá strana udělá totéž se svým párem klíčů. Když chci něco odeslat, zašifruji data veřejným klíčem adresáta, pošlu a adresát si to dešifruje svým soukromým klíčem.

Nevýhodou symetrických šifer je tedy problém předání klíče protistraně, ale výhodou velká rychlost šifrování a dešifrování. U asymetrických šifer je to přesně naopak – výhodou je možnost transportu veřejného klíče i nezabezpečeným kanálem, nevýhodou velká výpočetní náročnost (hodně zdržuje přenos).

Proto se v praxi oba přístupy některým způsobem kombinují – používáme *hybridní šifrování*. Po navázání spojení se přenášená data šifrují symetricky, aby nebyla komunikace příliš zdržována, ale předem se symetrický klíč předá pomocí asymetrického šifrování (místo dat prostě zašifrujeme symetrický klíč, čímž ho dokážeme relativně bezpečně předat protistraně).



Další informace:

Pro ty, kteří cítí potřebu navýšit své znalosti o šifrování – různé šifrovací algoritmy jsou popsány například ve skriptech jiného předmětu: <http://vavreckova.zam.slu.cz/obsahy/analyzadat/analyzadat.pdf>.

Princip šifrovacích algoritmů je vcelku podrobně popsán v knize [11].



⁵<http://www.kismetwireless.net/>

Pro zvýšení zabezpečení se kromě (symetrického) statického šifrovacího klíče pro každý odeslaný rámec používá také dynamicky měnící se dodatek – *IV vektor* (inicializační vektor). Účelem je maximálně ztížit odposlech komunikace. Jenže když se něco neustále mění, musíme mít mechanismus, jak dát adresátovi vědět, jakou podobu má IV vektor pro ten konkrétní rámec (je pěkné, že hacker náš rámec nedešifruje, ale adresát by toho měl být schopen). Pro to jsou dva přístupy: buď mají obě strany algoritmus, kterým pro daný rámec IV vektor určí, nebo momentální IV vektor prostě adresátovi pošleme. První možnost je samozřejmě mnohem bezpečnější.

 **WEP (Wired Equivalent Privacy).** Jedná se o mechanismus zabezpečení, který už rozhodně bezpečným nazývat nelze. Jedinou výhodou je plná kompatibilita se vším, co má implementováno IEEE 802.11.

Používá se stejné šifrování pro první fázi komunikace i pro běžný provoz (pokud vůbec nějaké). Šifruje se pouze symetrickým algoritmem (nepoužívá hybridní postup), tedy klíč musí být nějak předán adresátovi (klientovi). IV vektor se v naprosté většině případů přenáší jako součást (nešifrovaného) záhlaví.

WEP nabízí tři režimy:

- otevřená síť (bez šifrování),
- symetrická šifra RC4 s 64bitovým klíčem, přičemž IV vektor tvoří jeho 24 bitů (takže 40 bitů je statická část klíče),
- symetrická šifra RC4 se 128bitovým klíčem, přičemž IV vektor taky tvoří jeho 24 bitů (104 bitů je statická část klíče).

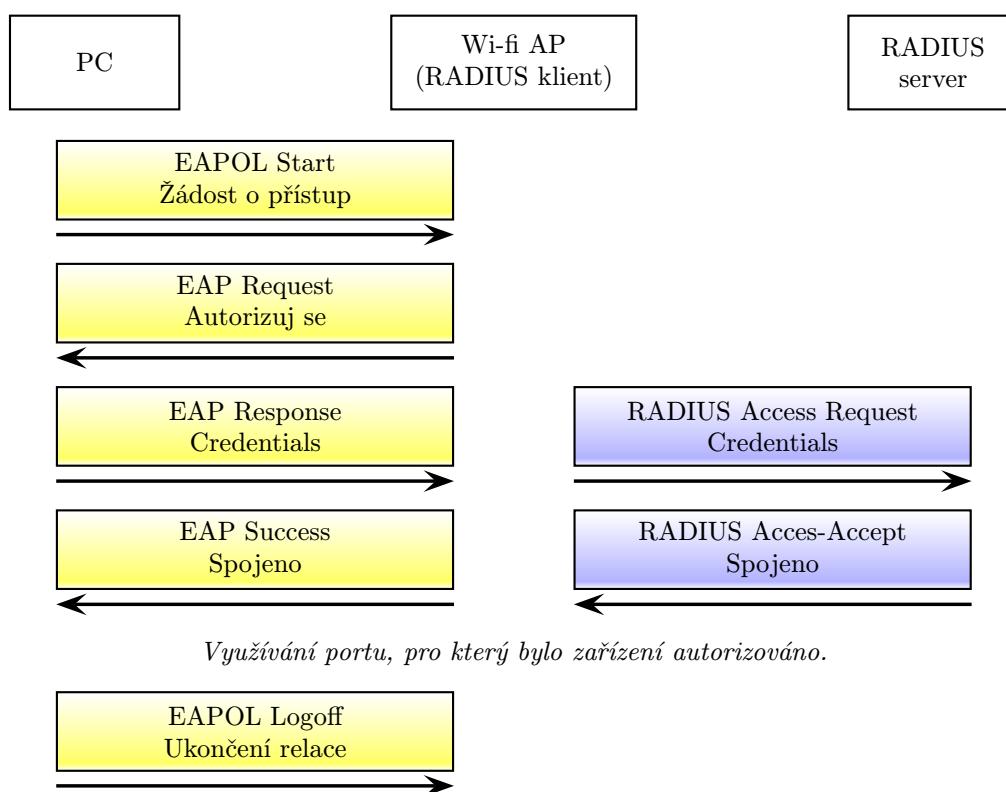
V každém případě pro cracknutí klíče stačí několik minut (s použitím programu AirCrack), takže WEP je považován za naprosto nevyhovující způsob zabezpečení.

 **WEP a IEEE 802.1X.** Mechanismus WEP byl brzy shledán nedostatečným, ale jiná možnost dlouho neexistovala. Proto se ve firemních sítích začala používat dodatečná možnost zabezpečení přihlašování ve formě implementace protokolu IEEE 802.1X (pozor, ne 11). Spočívá v zajištění autentizace a autorizace pomocí autentizačního serveru, který si vede databázi „povolených“ uživatelů s přístupovými údaji pro autentizaci (credentials) a údaji pro autorizaci.

 Jako autentizační servery se obvykle používají RADIUS, Diameter a TACACS+ (taky se označují jako AAA protokoly, protože každý z nich představuje také protokol postavený na IEEE 802.1X).

- Pro *RADIUS* (RFC 2865) existují volně dostupné implementace, nejoblíbenější je open-source projekt FreeRADIUS.
- *Diameter* (RFC 6733 a několik dalších dokumentů) má také volně dostupné implementace, například freeDiameter. Je novější než RADIUS a je považován za bezpečnější, například místo UDP komunikuje přes TCP nebo SCTP, jde o potvrzovaný přenos a odpovědi musí jít po stejné cestě jako dotazy, atd.
- *TACACS+* je proprietární protokol od Cisca. Navazuje spojení (používá TCP), potvrzuje komunikaci.

Princip si popíšeme na RADIUSu, u ostatních je podobný (až na konkrétní podobu komunikace, navazování spojení apod.). Uživateli nestačí jen vlastnit klíč, musí mít i své přihlašovací údaje (například jméno a heslo nebo třeba hardwarový token). Na obrázku 5.7 je naznačena komunikace s RADIUS serverem – klient se sice přihlašuje do některé buňky k příslušnému AP, ale místo toho, aby AP o připojení rozhodoval přímo, pouze přihlašovací údaje přepośle RADIUS serveru a při autentizaci a autorizaci vlastně funguje



Obrázek 5.7: Zjednodušená komunikace v síti při použití RADIUS serveru

jen jako zprostředkovatel. Kromě přihlašovacích údajů (credentials) si RADIUS server může navíc vyžádat další údaje, podle kterých pak rozhodne o autentizaci a konkrétních parametrech autorizace (přístupových oprávněních).

RADIUS síť obsahuje jen dva uzly – RADIUS server a RADIUS klienta, kterým je dotýčný AP, spojení je point-to-point. Toto spojení by mělo být dobře zabezpečeno (žádná Wi-fi, zde by měl být kabel, a taky šifrování). Samotný klient (žadatel, suplikant) není součástí RADIUS sítě.

WPA (Wi-fi Protected Access). Mechanismus WPA byl původně dočasným řešením, vznikl při čekání na standard WPA2. Dnes je považován za *něco mezi* nezabezpečeným WEP a bezpečným WPA2, přičemž na rozdíl od WPA2 je vhodný i pro starší zařízení, na kterých nelze WPA2 použít. Na takovém zařízení stačilo upgradovat firmware (tedy softwarová záležitost), WPA nevyžadoval úpravu hardwaru. Zjednodušeně můžeme WPA brát jako WPA2 osekáný o hardwarově závislé součásti.

Základ je podobný jako u WEP – šifra RC4 se 128bitovým klíčem, jehož IV vektor je 48bitový. Ovšem správa dynamických IV vektorů je prováděna protokolem TKIP (Temporal Key Integrity Protocol), takže IV vektor se už nedá odposlechnout z nešifrovaného záhlaví rámce. Navíc protokol TKIP přidává k rámcům 64bitový kontrolní součet (spíše digitální podpis rámce), který nazýváme MIC (Message Integrity Check, lidově „Michael“), podle kterého se pozná, jestli nebyl rámec cestou poškozen či záměrně upraven.

WPA pracuje ve dvou variantách:

- WPA-Enterprise – pro korporátní sféru, v podstatě pokračovatel WEP+RADIUS, v síti musí být autentizační server,
- WPA-Personal (WPA-PSK) pro domácnosti a malé firmy, pro autentizaci se používá sdílený klíč (PSK = Pre-Shared Key), nemáme autentizační server, ale uživatel i tak potřebuje autentizační

informace (heslo 8–64 znaků nebo 64 hexadecimálních číslic), ze kterých se dynamicky určuje klíč.

 **WPA2 (IEEE 802.11i).** Pro tento mechanismus existuje vlastní standard. Oproti WPA se místo protokolu TKIP používá protokol CCMP (Counter-Mode CBC MAC Protocol) s šifrováním AES, a právě použití šifrování AES vyžaduje hardwarovou podporu na zařízeních (dnes již většina procesorů obsahuje AES instrukce, takže obvykle nebývá problém se zprovozněním).

Stejně jako u WPA, i zde potřebuje uživatel přihlašovací údaje, které jsou využity buď pro autentizační server nebo pro protokol PSK, tedy opět máme dvě možnosti:

- WPA2-Enterprise – pro korporátní sféru, s autentizačním serverem,
- WPA2-Personal (WPA2-PSK) pro domácnosti a malé firmy.

V roce 2017 se objevila zpráva, že mechanismus WPA2 je náchylný na útok KRACK (Key Reinstallation Attacks), pomocí kterého lze obejít zabezpečení pomocí šifrování spojení.

 **WPA3** byl uveden roku 2018. Opět existují dva režimy – WPA3-Enterprise a WPA3-Personal. Pro Enterprise použití se předepisuje šifrování AES-256 v módu GCM s hashováním HMAC SHA-384. V Personal režimu lze použít AES-128 s protokolem CCMP jako u WPA2, ale výměna klíčů probíhá s využitím algoritmu SAE (Simultaneous Authentication of Equals), což je varianta výměny klíčů Dragonfly Key Exchange (funguje podobně jako Diffie-Hellman).

Příchod WPA3 byl uspišen objevením závažných bezpečnostních chyb ve WPA2, nicméně bezpečnostním chybám se nevyhnul ani samotný mechanismus WPA3. Například výše zmíněný algoritmus Dragonfly (ve WPA3 se používá při navázání spojení pro bezpečnou výměnu klíčů) má bezpečnostní slabinu DragonBlood, která je zneužitelná rovnou několika různými způsoby.



Další informace:

- <https://www.krackattacks.com/>
- <https://github.com/freeDiameter/freeDiameter>
- https://docs.oracle.com/cd/E86291_01/docs.81/20170711_120805_m_dsr_radius/concepts/c_dsr_radius_radius_versus_diam.html
- <https://www.root.cz/clanky/sifrovani-wpa2-bylo-prolomeno-wi-fi-site-je-mozne-odposlouchavat/>
- <https://www.root.cz/zpravicky/dragonblood-bezpecnostni-chyba-ve-wpa3-umoznuje-ziskat-heslo-k-wi-fi/>
- <https://www.wi-fi.org/discover-wi-fi/security>
- <https://medium.com/@reliancegcs/wpa3-explained-wi-fi-is-getting-major-security-update-2b6dca8f3aff>
- <https://www.security-portal.cz/clanky/wifi-sitě-jejich-slabiny>



5.3.2 WPS

Mechanismus WPS (Wi-fi Protected Setup) se objevil s tím, jak přibývaly různé typy zařízení připojovaných do Wi-fi sítě. Mnohá zařízení nemají klávesnici ani display, takže zadávání hesel by bylo problematické. WPS slouží ke zjednodušené autentizaci a asociaci takovýchto zařízení do Wi-fi sítě.

Existují dvě možnosti použití mechanismus WPS – WPS tlačítko nebo WPS PIN.

 **WPS tlačítko (PBC – Push Button).** Na AP zmáčkne tlačítko označené WPS (nebo QSS či PBC) nebo v administraci AP klepneme na „softwarové“ tlačítko. V průběhu několika desítek sekund

AP skenuje svou buňku a hledá nová zařízení. Pokud se v tomto časovém intervalu (kterékoliv) zařízení v buňce začne hlásit (čehož docílíme tak, že na na dotyčném zařízení taky zmáčkne tlačítko nebo provedeme něco podobného podle návodu), je považováno za důvěryhodné a AP s ním dojedná asociaci. Problémy jsou především tyto:

- V uvedeném časovém intervalu je naše síť prakticky nechráněná, cokoliv se může dostat dovnitř.
- Musíme stihnout na asociovaném zařízení včas zmáchnout tlačítko, takže v některých případech to je spíše práce pro dva.

 **WPS PIN.** Pro zařízení, které chceme takto dostat do sítě, musíme zjistit jeho PIN – 8místné číslo, které najdeme buď na nálepce přímo na zařízení nebo někde v dokumentaci. PIN pak nafukáme v administraci svého AP, načež si opět AP se zařízením vyjednají veškeré parametry a provede se asociace.

Ovšem problémy jsou i s touto metodou:

- AP s podporou WPS PIN *naslouchá neustále*.
- PIN je 8místné číslo, přičemž poslední číslice se vypočítává z předchozích, je obdobou kontrolního součtu. Pokud PIN do administrace zadáme špatně, AP vrací informaci, ze které se dá odvodit, která polovina je chybná – pokud je chybná první polovina, stačí uhodnout tři číslice, což je pro průměrného hackera časově nenáročná hra.

Pokud náš Wi-fi AP má tuto funkci, je lepší ji vypnout a zapínat opravdu jen tehdy, když chceme takto asociovat nové zařízení.

5.3.3 Jak tedy zabezpečit bezdrátovou síť

Takže co udělat, aby naše Wi-fi síť byla co nejlépe zabezpečená? Předně nastavit šifrování WPA2, pokud to všichni naši klienti zvládají. A co dál?

 **Credentials do administrace.** Přístupové údaje do administrace AP jsou z výroby nastaveny na určité standardní hodnoty (typicky admin-admin nebo něco takového). V žádném případě bychom to tak neměli nechat! Takže hned po zakoupení během počáteční konfigurace bychom jako první měli nastavit vlastní přístupové údaje do administrace AP.

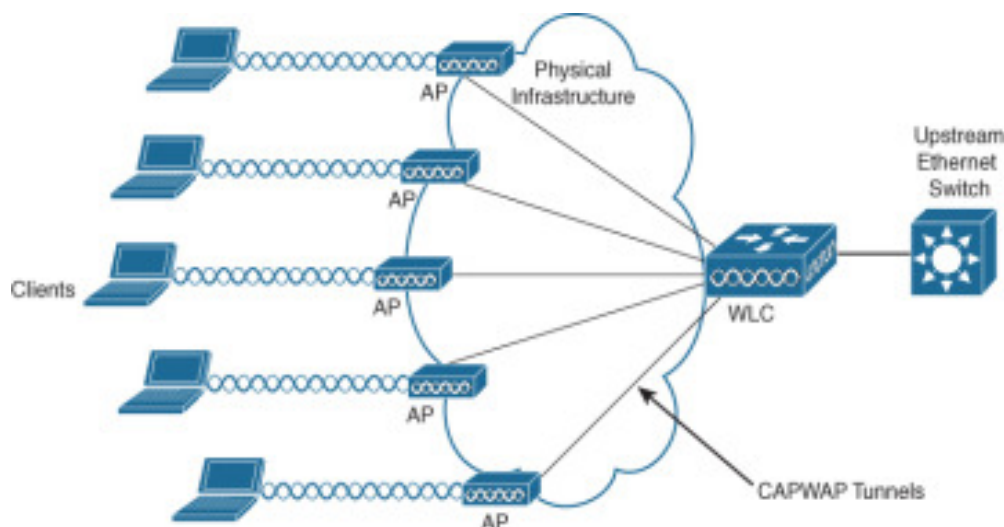
 **Skrytí SSID.** AP vysílá v pravidelných intervalech Beacon rámeček, kde je mimo jiné SSID (tedy název sítě). SSID je potřeba pro přihlášení do sítě, a tedy když AP nebude vysílat Beacon rámeček s SSID, teoreticky bychom tím zabránili pokusům o přihlášení osob, které v síti nemají co dělat. Prakticky se však SSID dá odposlechnout v autentizační komunikaci „povolených“ zařízení, takže stačí donutit některé takové zařízení k znovupřipojení (shodíme jeho komunikaci).

Navíc skrytí SSID může způsobovat problémy při hledání zdroje rušení jiné sítě – může nastat situace, kdy vidíme, že nám něco ruší síť, ale nevidíme, co to je, protože dotyčný má skryté SSID.

 **Filtrování MAC adres.** V administraci AP si můžeme vytvořit black list (seznam zakázaných) nebo white list (seznam povolených, nikdo jiný nebude asociován) MAC adres. Ale stačí si uvědomit, že MAC adresa se dá pozměnit (v Linuxu na to stačí jeden příkaz, ve Windows úprava registru nebo speciální program), takže opět bezzubé opatření.

5.4 Centrálně řízená bezdrátová síť

 V domácnostech a malých firmách se čím dál více objevují tzv. *Wi-fi mesh sítě*, což je ve skutečnosti jednoduchý WDS propojující několik AP v sadě. Jeden prvek ze sady bývá řídicí, jeden z jeho portů slouží jako WAN port, ostatní komunikují přes tento řídicí prvek. Mezi AP probíhá komunikace buď bezdrátově po vyhrazeném kanálu, nebo přes kabel v ethernetové síti, záleží na konkrétním řešení. Obvykle se jedná o proprietární řešení, tedy koupíme sadu od jednoho výrobce.



Obrázek 5.8: Wireless controller a odlehčené AP⁶

 Ve firemním prostředí se už mnohem déle používají *centrálně řízené sítě* (původně se nazývaly sítěmi čtvrté generace), kde je účelem spolehlivé pokrytí rozlehlých prostor signálem tak, aby se přes odlišnost v konceptu sítě mohli připojovat klienti s běžnými Wi-fi zařízeními. U různých výrobců se používá mírně odlišná terminologie, ale princip je podobný.

 Centrálním prvkem sítě je prvek, který se nazývá *Wireless controller* (Wireless LAN Controller, WLC). Tento prvek obsahuje konfiguraci sítě, zde se zajišťuje synchronizace sítě, zabezpečení, management, napojení „ven“, atd. Může se jednat o fyzické zařízení, ale existují i virtuální implementace. Další prvky v síti jsou *Lightweight AP* (LAP), „odlehčené“ přístupové body, které si můžeme představit jako „prodlouženou ruku“ controlleru, obdobu hodně dlouhých antén. LAP jsou ke controlleru připojeny pomocí ethernetových kabelů, a přes tyto kabely jsou i napájeny s použitím PoE.

Typické použití této technologie je ve velkých halách či jiných rozlehlých místnostech (továrny, sklady, hypermarkety, ale třeba i nemocnice nebo open-space prostory). Předpokládejme, že bychom použili klasické AP propojené do běžného WDS. Pokud bychom v takovém případě v rozlehlé místnosti s vysokým stropem umístili access pointy na strop, museli bychom řešit dilema: aby signál dosáhl až na zem, musel by být dostatečně silný, ale potom by AP musely být hodně daleko od sebe, aby se navzájem nerušily. Jenže pak by u země měl signál velmi proměnlivou dostupnost, vznikaly by „mapy“ míst bez signálu.

Řešení využívající controller a odlehčené AP umožňuje vytvořit hustou síť těchto AP, které se nebudou navzájem rušit, protože všechny mohou vysílat v totéž pásmu (resp. na více stejných kanálech) – je úplně jedno, se kterým z nich klient komunikuje. Handshake (předávání mezi AP) není nutné komplikovat,

⁶Zdroj: <https://www.ciscopress.com/articles/article.asp?p=3089360&seqNum=3>

veškerá správa je čistě na straně controlleru.

 LAP komunikují s WLC pomocí určitého protokolu – buď některého proprietárního, nebo se používá protokol CAPWAP (Control and Provisioning of Wireless Access Points Protocol) standardizovaný v RFC 5415.

V síti se obvykle používají VLAN sítě, pro každou na síťové vrstvě existuje vlastní podsít. Vzhledem k tomu, kde se takové sítě využívají, bývají součástí infrastruktury také autentizační servery (například Radius, TACACS+, Diameter). Mezi autentizačními servery a centrálním řídicím prvkem musí existovat fyzicky zabezpečené spoje.

Jako koncová zařízení se do sítě mohou připojit jakákoliv zařízení podporující IEEE 802.11 – notebooky, mobily, IoT zařízení apod. Jako řešení pro pokrytí rozsáhlých prostor to je výborné (čím vyšší hustota LAP, tím lepší signál), jedinou nevýhodou je cena.



Další informace:

- <https://www.cisco.com/c/en/us/products/wireless/wireless-lan-controller/what-is-wlan-controller.html>
- <https://www.samuraj-cz.com/clanek/obecne-o-konfiguraci-wlc-a-zvlaste-cisco-2500-wireless-controller/>
- <https://community.fs.com/blog/wireless-lan-controller-explained.html>
- http://www.intelek.cz/art_doc-D7A489A18B634F84C12575550053ECAE.html
- [http://www.intelek.cz/db/media.nsf/w/E4A2BCABAE379135C125732300589FD3/\\$file/extricom.pdf](http://www.intelek.cz/db/media.nsf/w/E4A2BCABAE379135C125732300589FD3/$file/extricom.pdf)



5.5 Mobilní technologie

Mobilní sítě jsou především charakteristické svou mobilitou, tedy stanice, která je součástí mobilní sítě, se může pohybovat se zachováním signálu, a to i mezi různými buňkami (oblastmi pokrytí signálu). V této sekci se budeme zabývat především *mobilními WAN*, třebaže existují řešení i pro menší sítě.

 Proces přecházení mezi buňkami bez přerušení spojení se zachováním všech parametrů a probíhající datových přenosů (relací) se nazývá *handover* (předávání, také handoff). U mobilních sítí je tato funkcionality velmi důležitá.

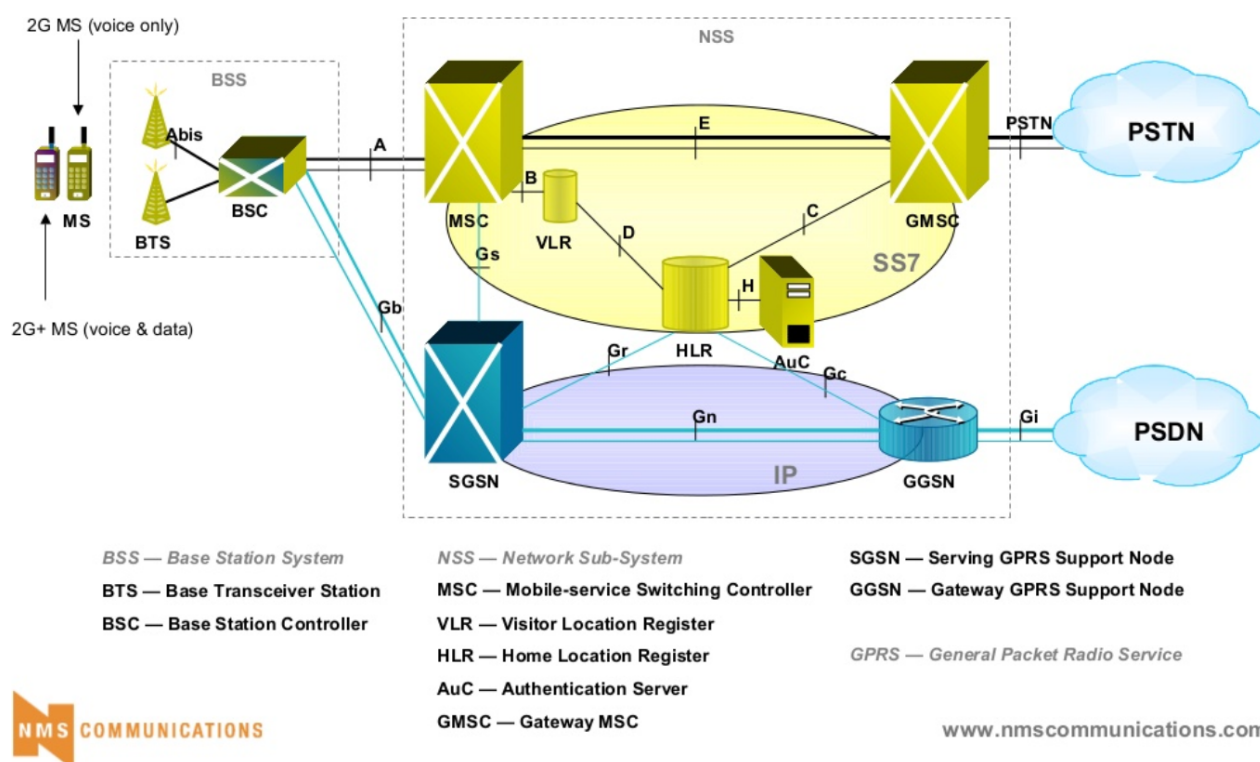
5.5.1 Struktura mobilní sítě

V mobilní síti potřebujeme určitou infrastrukturu sloužící jako pro přenos hlasu a dat, tak i pro řídicí a správní účely. Konkrétní názvy jednotlivých typů zařízení jsou v různých typech a generacích mobilních sítí odlišné, zde jen nastiňujeme základní princip a strukturu.

 Klientská zařízení přímo komunikují se *základnovou stanicí* (Base Station), oblast pokrytá signálem určité základnové stanice je její buňka (tj. klientská zařízení se nacházejí v buňce některé základnové stanice). Označení může být například BTS (Base Transceiver Station), Node B, atd.

Základnové stanice zvládají spíše jen samotnou komunikaci, logika sítě je trochu jinde. Vždy skupina základnových stanic je napojena na *radič pro hlasovou síť* a *radič pro datovou síť* (podle toho, co má být přenášeno), zde se dělí trasa hlasového signálu a dat.

To by byla *vrstva základnových stanic* a jejich správy.

Obrázek 5.9: Struktura mobilní sítě⁷

Trasy pro hlas a data mají společné body v dalších dvou vrstvách celé struktury:

- *administrativní podsystém* (operační) – zde najdeme například registr stanic a s ním související autentizační servery, tarifkaci apod.,
- *podsystém správy sítě* – zajišťuje fungování sítě včetně provozu bran do jiných sítí a jiných typů sítí (brány do sítí jiných poskytovatelů, brány na Internet, atd.).

To je jen hrubý a osekáný nástin, u každé technologie se používá jiná terminologie a také mohou být přidány další typy uzlů. Na obrázku 5.9 je ukázka struktury pro síť GSM kombinovanou s GPRS.



Další informace:

- <https://github.com/jreisinger/blog/blob/master/posts/gsm.md>
- <http://www.earchiv.cz/a008s200/a008s200.php3>



5.5.2 První generace (1G)

První generace mobilních datových sítí byla ještě analogová. Z toho důvodu se také využíval frekvenční multiplex (FDMA), každý připojený uživatel má přidělen jeden kanál, a to výhradně.

První komerční mobilní systém byl *NMT* (Nordic Mobile Telephone) fungující nejdříve v Saudské Arábii a následně v severských zemích (Norsko, Švédsko) od roku 1981. Do první generace také patří americký *AMPS* (Analog Mobile Phone System) z roku 1982 a *TACS* (Total Access Communications System) vytvořený původně pro Velkou Británii, ale později se rozšířil především v Asii.

⁷Zdroj: <https://www.slideshare.net/tbutomo/3-g-tutorial-32431714>

5.5.3 Druhá generace (2G)

Na začátku 90. let (1992) vznikla druhá generace, která je dodnes funkční (vedle dalších generací). Byl zaveden digitální přenos dat, ale síť je optimalizována především pro hlas, tedy propustnost je na velmi nízké úrovni.

Sítě druhé generace pořád pracují na principu *přepojování okruhů*, a to i pro posílání dat. Důsledkem je obvyklá tarifkace odvozená od doby připojení (jde o vytáčené připojení).

 **GSM** (Global System for Mobile Communications) je digitální síť, ale pořád využívá přepojování okruhů, a to s časovým (TDMA) i frekvenčním (FDMA) multiplexem – kombinace TDMA over FDMA (časové sloty se přenášejí vždy v rámci určité frekvence). Hovoříme o *kanálech* (to jsou vlastně SVC okruhy), uživateli je vždy po určitou dobu vyhrazen určitý kanál. GSM pracuje na třech různých frekvencích, hovoříme o GSM900 (frekvence 900 MHz), GSM1800 (1800 MHz) a GSM1900 (1900 MHz).

Struktura víceméně odpovídá nákresu výše, jen zde není „datová odbočka“ – v GSM se data přenášela přímo po okruzích.

Základnové stanice jsou rozmístěny tak, aby se buňky (přibližně kruhové, i když jsou zobrazovány jako šestiúhelníky – „plástve“) překrývaly. To však znamená, že dvojice buněk, které jsou přímo vedle sebe, nesmí používat tytéž komunikační kanály, tytéž kanály může používat až buňka natolik vzdálená, aby nedocházelo k inferencím. Rozdělení kanálů se provádí metodou podobnou obarvování, kterou známe z teorie grafů.

Při plném pokrytí signálem je možná plná mobilita stanice, *handover* – předávání stanice mezi jednotlivými buňkami.

Oproti první generaci je zde především *přechod na digitální technologii*, dále větší odolnost proti rušení a odposlechu, snadné napojení na pozemní telekomunikační systémy, a také vyšší rychlost. Teoretická rychlost je 13 kb/s pro každý směr, reálně 9.6 kb/s.

 **CDMA** (Code Division Multiple Access) má sice ve specifikaci přenos dat i hlasu, ale u nás je CDMA poskytována jen jako datová síť. Oproti GSM má výhodu nižších frekvencí (450 MHz), základnové stanice CDMA mohou být od sebe více vzdáleny. Teoretická rychlost uploadu je 1024 kb/s, té však lze dosáhnout jen v místech s výkonnějšími základnovými stanicemi (například v Praze nebo v Brně). Download je o něco rychlejší (až 3.1 Mb/s, v reálu také méně, několik set kb/s), jde o asymetrický přenos.

CDMA používá jiný typ multiplexu – kódový. Spočívá v tom, že v rámci jednoho sdíleného média se přenáší více digitálních signálů, které jsou rozlišovány pomocí kódování (tedy ne rozdělením do časových slotů ani pouhým rozvržením do frekvencí).



Poznámka:

U jednotlivých technologií jsou uváděny (teoretické) rychlosti. Tyto údaje berte s rezervou (může se na nich podepsat poměrně hodně parametrů), a taky nebudou vyžadovány na zkoušce. Jen je dobré mít představu o tom, které technologie jsou rychlejší a které pomalejší.



5.5.4 Přechodová generace (2.5G)

Zatímco mobilní sítě druhé generace byly založeny na přepínání okruhů, v přechodové generaci se již používá přepínání paketů, které je pro přenos dat výhodnější. Typická tarifkace je buď za přenesená data nebo paušál, netarifikuje se připojená doba (tj. nejedná se o vytáčené připojení).

 **GPRS** (General Packet Radio Service) je nástavba GSM (technicky ne nutně GSM, může být vázána na jiné typy sítí), která umožňuje jedné stanici využívat až 8 GSM kanálů (pokud jsou zrovna volné). Nejedná se ještě o širokopásmovou technologii, ale díky rozšíření pásma až na 8násobek jsou přenosové rychlosti vyšší. Jsou využívány ty z osmi GSM kanálů, které jsou právě volné, tedy přenosová kapacita sítě je využívána účelněji, kanály nejsou blokovány po celou dobu „připojení“.

Aby fungovalo napojení na GSM, bylo nutné do sítě GSM přidat ještě jednu vrstvu. Vrstva základnových stanic GSM zůstává, ale byla přidána dodatečná vrstva podpůrných uzlů sítě GPRS a další uzly byly přidány do vrstvy řízení sítě i do administrativní vrstvy (tam také, protože mechanismus tarifkace je odlišný). Přidané uzly mezi sebou komunikují jiným protokolem než ty původní – GTP (GPRS Tunelling Protocol), což je aplikační protokol nad protokoly TCP a UDP.

Přenosové rychlosti se odvozují od použitého kódování. Čím silnější kódování, tím nižší rychlost. Celková možná dosažená rychlost je 22.8 kb/s na jeden kanál, ale část rychlosti je pohlcena režii kódování. Při nejsilnějším kódování (CS-1) dosahujeme rychlosti maximálně kolem 9 kb/s na kanál, při nejslabším kódování (CS-4) maximálně kolem 21 kb/s na kanál. Volba kódování je zcela automatická, závisí na vzdálenosti stanice od základnové stanice (čím dál je stanice, tím silnější musí být kódování).

Dále je rychlost přenosu ovlivněna množstvím GSM kanálů, které stanice může využívat (maximálně 8 pro každý směr). V základním nastavení musí uživatel GPRS vzít zavděk tím, co „zbude“ po uživatelích GSM hlasových přenosů, ale v některých typech tarifů se můžeme setkat s garancí určitého počtu kanálů (vyhrazením pouze pro data).

GPRS zavádí základní možnost řízení kvality služeb (QoS) založenou na prioritě, propustnosti, zpoždění nebo spolehlivosti. Je samozřejmě na operátorovi, zda bude QoS nabízet svým zákazníkům.

 **EDGE.** Technologie EDGE (Enhanced Data rates for GSM Evolution) je sice také nástavbou sítě GSM a základní princip je podobný, ale navíc signál moduluje metodou 8-PSK, což v reálu znamená až trojnásobný nárůst rychlostí oproti GPRS.

5.5.5 Třetí generace (3G)

Třetí generace přináší možnost *souběžného využívání několika služeb* (u předchozích generací nebylo možné kombinovat hlas a data). Technologie jsou založeny na variantách metody CDMA, tedy se používá kódový multiplex. Na rozdíl od předchozích generací, třetí generace již není striktně optimalizována pro přenos hlasu, více se počítá s datovými přenosy.

 **UMTS** (Universal Mobile Telecommunication System) není pouhá nástavba, zavedení UMTS znamenalo poměrně velké zásahy do infrastruktury sítě. V tomto typu sítě je také nabízena vylepšená QoS.

Základem je WCDMA (Wideband CDMA), což je širokopásmová metoda přenosu. V metodě WCDMA má každý uživatel přiděleno frekvenční pásmo, které používá. V kódovém multiplexu může totéž pásmo využívat více uživatelů, jsou odlišeni jednoznačným binárním kódem.

Architektura sítě je podobná síti GPRS a zčásti na ní staví. Uzel s názvem *Node B* je obdobou základnové stanice. Zprostředkovává komunikaci mezi koncovou stanicí a samotnou UMTS sítí, zajišťuje

modulaci, kódování, ochranu proti chybám, tak jako původní základnové stanice. Také do dalších vrstev původní sítě byla přidána řada nových uzlů.

UMTS není u nás, ale ani v jiných zemích, moc rozšířená. Vzhledem k nutným úpravám infrastruktury se s UMTS typicky počítá pouze ve velkých městech. Teoretická rychlost je až 2 Mb/s, ve skutečnosti se však dosahuje pouze rychlosti několika desítek (místo i stovek) kb/s.

 **HSDPA** (High-Speed Downlink Packet Access) přináší další zvýšení rychlosti a snížení latence, jedná se tedy o jakési vylepšení UMTS pro download (slovo „Downlink“ v názvu). Změny je docíleno zvolenou modulací – QPSK (Quadrature Phase Shift Keying, používá se i v UMTS) nebo 16QAM, a dále jinými mechanismy plánování vysílání (plánování založené na QoS) a řízení chyb.

Stejně jako UMTS, také HSDPA používá WCDMA, ale navíc přidává nové mechanismy a nový typ čistě datového kanálu. Zatímco v UMTS jsou pakety odesílány každých 10 ms, v HSDPA jsou odesílány každé 2 ms. Pro zvýšení propustnosti je také možné využít technologii MIMO.

Přenosové rychlosti jsou obecně vyšší než v případě UMTS, v současné době kolem 1 Mb/s, ovšem stejně jako u jiných mobilních sítí s velkými odchylkami v obou směrech. Existují varianty HSDPA, které dosáhnou až na 14 Mb/s, ale ty u nás nejsou implementovány.

 **HSUPA** (High-Speed Uplink Packet Access) je obdobou předchozí technologie pro upload. Jejím účelem je tedy zvýšit rychlost a celkovou kvalitu (vč. latence) u uploadu.

Technologie HSDPA a HSUPA se souhrnně označují *HSPA*, ale operátor nemusí nutně použít obě (u nás O2 implementuje pouze UMTS/HSDPA).

 **LTE** (Long Term Evolution) je sice prezentována jako mobilní síť čtvrté generace, ale ve skutečnosti se jedná o třetí generaci (čtvrtá je až LTE Advanced). Je to čistě datová technologie, s přenosem hlasu se zde vůbec nepočítá. To znamená, že pokud chce uživatel použít hlasové služby, buď zvolí VoIP nebo obdobnou technologii (existuje také VoLTE), nebo bude LTE dočasně deaktivováno (což znamená, že když se přes LTE zrovna stahují data, bude jejich stahování přerušeno).

 Poznámka:

Někteří uživatelé si kupují mobilní zařízení „v Číně“. To není až takový problém, až na to, že standard LTE umožňuje vysílat ve více různých frekvenčních pásmech, přičemž jiná pásma se používají v Asii a jiná v Evropě. Nejběžnějším pásmem pro LTE u nás je pásmo kolem frekvence 800 MHz, ve větších městech mohou (nemusejí) být dostupné i některé další frekvence (1800 MHz a 2100 MHz) – s posunem se do budoucna moc nepočítá. V každém případě: pokud si koupíme LTE zařízení, které nepodporuje frekvenci 800 MHz, nebude nám toto zařízení tak užitečné jak bychom čekali.



LTE je první technologií 3G sítí, která je plně implementovaná jako síť postavená na IP (přepínání paketů), předchozí alespoň částečně stojí na přepínání okruhů. Na downlinku používá multiplex OFDMA, na uplinku SC-FDMA (speciální odnož OFDM optimalizovaná pro použití na uplinku).

Cílem je dosáhnout velmi vysokých rychlostí (až 100 Mb/s na downstreamu, 50 Mb/s na upstreamu na kanál) při velmi nízké latenci (10 ms, i méně). V současné době LTE reálně dosahuje na downstreamu téměř 80 Mb/s.

 Další informace:

- <http://home.pf.jcu.cz/~pepe/Diplomky/velicky.pdf>

- <http://access.feld.cvut.cz/search.php?rsvelikost=sab&rstext=all-phpRS-all&rstema=10&stromhlmenu=10>



5.5.6 Čtvrtá generace

Technologie 4G byla standardizována ITU-R pod označením IMT-Advanced (International Mobile Telecommunications Advanced). Je to souhrn podmínek, které musí splnit systém, aby mohl být označen jako systém 4G.

 **LTE Advanced** (také se přímo označuje jako IMT Advanced – International Mobile Telecommunications) vypadá pro uživatele velmi výhodně – přenosová rychlost může být teoreticky až 1 Gb/s (downlink), resp. zhruba poloviční rychlost na uplinku. Latence je přibližně 5 ms, což je polovina latence původního LTE (takže síť má lepší odezvu). Používá se stejné modulační schéma jako u původního LTE (OFDMA a SC-FDMA) v kombinaci s MU-MIMO (využívání více antén), rozdíl je v efektivitě a rozsahu využívání spektra a dále v metodách správy celé sítě.

LTE Advanced ve skutečnosti vznikla úpravou LTE, tedy trpí podobnými problémy. Například jde opravdu jen o datovou síť, proto lze hovory na zařízení provozovat jen buď přes některou starší technologii (tudíž přerušit případný probíhající datový přenos) nebo použít VoLTE (pokud to náš ISP podporuje).

 **Další 4G technologie** se sice taky objevily, ale buď se neprosadily, nebo se od nich postupně upustilo (což však nemusí být finální stav). Například:

- FLASH-OFDM (Fast Low-latency Access with Seamless Handoff-OFDM) propojuje OFDM s metodami CDMA a TDMA. Výsledkem jsou vyšší přenosové rychlosti a nízká latence. Další výhodou je pokročilá možnost řízení kvality služeb.
- IEEE 802.16m (WiMAX Advanced, Wireless MAN Advanced) je úpravou původní technologie WiMAX tak, aby splňovala kritéria IMT-Advanced.
- UMB (Ultra Mobile Broadband) měl být následník CDMA přizpůsobený požadavkům IMT-Advanced. Qualcomm se rozhodl ukončit vývoj ve prospěch LTE Advanced.



Další informace:

- <http://www.itu.int/pub/R-REP-M.2134-2008/en>
- <https://www.hindawi.com/journals/wcmc/2017/2315036/>
- <http://www.radio-electronics.com/info/cellulartelecomms/lte-long-term-evolution/4g-lte-advanced-carrier-channel-aggregation.php>



5.5.7 Pátá generace

5G síť (IMT-2020) ve skutečnosti nejsou reprezentovány jedinou technologií, je to balík různorodých technologií, který má oproti předchozí generaci zajistit vyšší propustnost (až 10 Gb/s) při nízké latenci (kolem 1 ms) a nízkém odběru energie, také při vysokém počtu připojených klientů na základnovou stanici.

V posledních generacích se teoretická rychlost navyšuje vždy cca desetkrát, ale nízká latence a nízká spotřeba jsou novinky důležité zejména pro zařízení Internetu věcí, M2M komunikace a v průmyslu.

 Základnové stanice jsou do centra sítě připojeny optickými kabely, což je dnes obvyklé vlastně u všech přístupových sítí (mobilní sítě jsou vlastně přístupové WAN sítě). Klienti se k základnovým stanicím připojují ve frekvenčních pásmech, která mohou být odlišná v různých zemích a pro různé poskytovatele konektivity. V ČR lze pro 5G využít jak kmitočty určené pro předchozí generace, které nejsou zrovna využívány jinými technologiemi, ale také jsou přidávána další pásma, včetně těch, která byla uvolněna z používání televizními multiplexy podle DVB-T. Použitelných frekvenčních pásem je sice hodně, ale dají se rozdělit do dvou skupin:

- Frequency Range 1, označení je Sub6GHz (tj. frekvence do 6 GHz), také C-band,
- Frequency Range 2, označení je mmWave (mikrovlny od 24 GHz výše), ovšem tyto frekvence vzhledem k horšímu dosahu a prostupnosti materiálů jsou vhodné jen na krátké vzdálenosti v malých buňkách, typicky v husté zástavbě.

Zatímco u předchozích generací převládalo multiplexování OFDM(A), v 5G lze použít celou řadu jak ortogonálních (OMA, Orthogonal Multiple Access), tak neortogonálních (NOMA, Non-Orthogonal...) technik s různými modulacemi. Taktéž se používá mnoho antén s MIMO (typicky mnohem více než u 4G) jak pro komunikaci s více zákazníky, tak i při Beamformingu.

 Stejně jako u předchozích generací se komunikuje v některé formě *duplexu* (FDD = Frequency Division Duplex, TDD = Time Division Duplex), tj. ve frekvenčním spektru se zároveň nacházejí oba směry komunikace. Při použití frekvenčního duplexu jsou různé směry komunikace v různých frekvenčních pásmech spektra, při použití časového duplexu nejsou odděleny frekvencemi, ale PDU se umísťují do časových slotů určených pro daný směr.

FDD je použitelný (doporučovaný) spíše pro symetrický provoz (download a upload mají stejnou propustnost, frekvence jsou staticky rozděleny), a to v pásmech Sub6GHz (tj. na nižších frekvencích). TDD se dá používat na všech frekvencích, ovšem na těch vyšších (mmWave) se používá prakticky výhradně.

 Protože pokrytí novými frekvencemi u nás zatím není zrovna ideální a „starší“ frekvence jsou většinou ještě obsazeny staršími generacemi, můžeme se setkat také se sdílením frekvenčního spektra mezi 5G a 4G. Propustnost a latence nedosahují tak dobrých hodnot jako v případě čistého 5G, protože 5G komunikace se musí spokojit s tím, co „zbude“, a také se přizpůsobit způsobu komunikace podle 4G (aby nedocházelo ke vzájemnému rušení). Navíc

 Protože 5G je určeno pro velmi různorodé zákazníky s odlišnými potřebami (také bezpečnostními), byla zvolena metoda *slicing* (rozdělení komunikačního pásma na plátky, řezy) fungující podobně jako VLAN v lokálních sítích. Jednotlivé slices jsou navzájem izolovány. Například zdravotnické či jiné bezpečnostně relevantní aplikace si o něco takového přímo žádají, stejně tak průmyslová automatizace.



Další informace:

- <https://www.rfwireless-world.com/Terminology/5G-FDD-vs-5G-TDD.html>
- <https://www.gsma.com/spectrum/wp-content/uploads/2021/03/5G-TDD-Synchronisation-QA.pdf>
- https://www.researchgate.net/publication/327787320_4G5G_Spectrum_Sharing_Efficient_5G_Deployment_to_Serve_Enhanced_Mobile_Broadband_and_Internet_of_Things_Applications
- <https://www.zive.cz/clanky/co-je-to-5g-a-proc-technicky-sef-t-mobile-us-varuje-pred-prilisnym-optimismem/sc-3-a-198163/default.aspx>
- <http://www.techtrained.com/what-are-the-multiplexing-techniques-in-5g/>



5.6 Další bezdrátové sítě

Dalších bezdrátových sítí (ať už fixních nebo mobilních) existuje ještě spousta, zde se soustředíme především na sítě doporučené pro Internet věcí a podobné oblasti. Specifikem je především to, že komunikace až tak moc nezáleží na člověku, většinou si „povídají“ zařízení samotná bez lidské asistence.

5.6.1 Zajištění provozu sítě

Pro zařízení Internetu věcí často (ne vždy) platí, že je třeba brát ohled na nízkou energetickou náročnost, často totiž není možné takové zařízení připojit do elektrické zásuvky. Napájení je řešeno buď vyměnitelnou či nabíjecí baterií, případně může být vybaveno solárními panely – tedy tam, kde opravdu slunce svítí. Běžné typy sítí jsou bohužel často poněkud energeticky náročné, zejména Wi-fi nebo klasické mobilní sítě.

Zařízení Internetu věcí mohou samozřejmě používat běžné (výše probírané) typy sítí, a taky některé používají. Buď proto, aby se daly jednoduše připojit do „běžné“ sítě, tedy kvůli kompatibilitě, nebo byly vybrány takové technologie, které mají energeticky úspornou variantu či mód. Například:

- NFC (o tom zde zatím nebyla řeč) nebo Bluetooth, Bluetooth má v novějších verzích speciální variantu BLE (Bluetooth Low Energy) pro snížení spotřeby energie,
- Wi-fi, podle IEEE 802.11ax dokonce lze využít funkci TWT, která značně snižuje spotřebu energie,
- mobilní sítě, zejména LTE, LTE Advanced a 5G, kde opět najdeme funkce pro snížení spotřeby energie a 5G je dokonce pro účely IoT optimalizována.

Samozřejmě lze použít také Ethernet, pokud to je fyzicky možné, přičemž přes PoE lze zařízení také napájet.

 Dále tu máme technologie, které byly optimalizovány pro chytré sítě, Průmysl 4.0, Internet věcí, komunikaci M2M (Machine-to-Machine), V2X (Vehicle to Vehicle/Infrastructure/cokoliv dalšího, tj. vozidlo komunikuje s jiným vozidlem, semaforem, atd.) a další.

- IoT obdobou Ethernetu a Wi-fi (tj. v oblasti LAN) jsou nejznámější ZigBee a Z-Wave,
- v oblasti WAN konkurují běžným mobilním sítím tzv. LPWAN sítě (Low Power WAN), zejména LoRaWAN, Sigfox a NB-IoT.

 **ZigBee**

 **Z-Wave**

 **LoRaWAN**

 **NB-IoT**

5.6.2 Komunikace na logické úrovni

MQTT, XMPP, CoAP,

Síťová a transportní vrstva

 **Rychlý náhled:** V této kapitole se budeme zabývat především tím, co se děje na síťové a transportní vrstvě. Zaměříme se prohloubení vědomostí o těch mechanismech, které nebyly do podrobností probrány na bakalářském stupni studia, a to zejména o překladu adres, skupinách na síťové vrstvě, objevování sousedů, a také se stručně podíváme na další protokoly transportní vrstvy.

Předpokládáme, že čtenář již některé znalosti má, především že ví, jak vypadají IPv4 a IPv6 adresy, jaké jejich typy existují a které adresy jsou vyhrazeny ke speciálním účelům, jak vypadá IPv4 a IPv6 paket, jak se dělá podsíťování a VLSM, u transportní vrstvy pak formát TCP a UDP segmentů a průběh TCP komunikace.

 **Klíčová slova:** IPv4, IPv6, CIDR, Token Bucket, ICMP, NAT, CGN, IP skupiny, IGMP, ARP, NDP, DoD, SEND, TCP, UDP, SCTP, RTCP, DCCP.

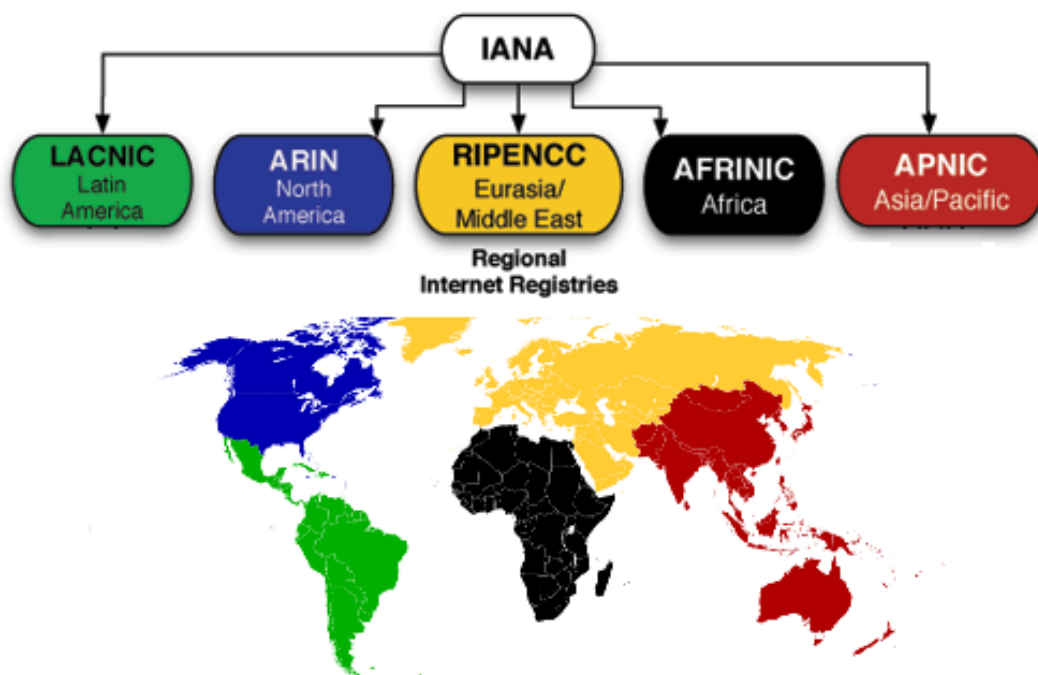
 **Cíle studia:** Cílem této kapitoly je prohloubit znalosti o protokolech síťové a transportní vrstvy a procesech, které na těchto vrstvách probíhají. Půjde o upevnění vědomostí o protokolech IPv4 a IPv6, rozebrání mechanismu NAT, přidělování IP adres a fungování skupin na síťové vrstvě, rozšíření vědomostí o protokolech transportní vrstvy.

6.1 IP adresy

6.1.1 Hierarchie adres a sítí

Hlavním garantem přidělování IP adres (ale také doménových jmen, čísel autonomních systémů, atd.) je ICANN (Internet Corporation for Assigned Network Numbers, <http://www.icann.org>), kdežto organizace IANA (Internet Assigned Numbers Authority, <http://www.iana.org/>) toto přidělování fyzicky provádí.

Celková struktura přidělování adres je hierarchická. IANA přiděluje bloky adres *regionálním registrátorům RIR* (Regional Internet Registry), což jsou RIPE (Evropa a část Asie), ARIN (Severní Amerika), AfriNIC (Afrika), LACNIC (Latinská Amerika), APNIC (Asie a Pacifik). Další patro hierarchie tvoří *lokální registrátoři* (LIR), kteří své bloky získávají od regionálních registrátorů. Od lokálních registrátorů pak své rozsahy adres získávají zákazníci nebo další subjekty, které mohou své rozsahy dále distribuovat.

Obrázek 6.1: Základ struktury přidělování IP adres¹

Na obrázku 6.1 jsou naznačena horní dvě patra této struktury. Níže jsou registrátoři pro jednotlivé země LIR, což mohou být třeba ISP (poskytovatelé Internetu), následují jejich zákazníci a drobní poskytovatelé Internetu.



Další informace:

Seznam LIR pro Evropu je na <https://www.ripe.net/participate/member-support/info/list-of-members/europe>, seznam pro ČR získáte tak, že v mapě klepnete na „Czech“.



6.1.2 CIDR



Beztrždní směrování – CIDR (Classless Inter-Domain Routing) ruší závislost mezi tvarem prvního oktetu a hranicí pro síťovou část adresy (což právě stanovují třídy). Hlavním motivem bylo zpřehlednění a zjednodušení směrování ve vyšších úrovních hierarchie, tedy u poskytovatelů síťového připojení na úrovni RIR a LIR.

Jak to funguje s CIDR: IANA přiděluje základní rozsahy jednotlivým RIR, tedy určitý počet bitů v adrese zleva. Tím je dána určitá základní délka prefixu. Každý RIR svým zákazníkům (tedy jednotlivým LIR) přidělí rozsah adres takový, že LIR „zdědí“ to, co má přiděleno RIR (tedy tuto část mají všichni zákazníci společnou), a k tomu jsou přidány další bity specifické pro dotyčného LIR – provede se podsíťování, tedy posun doprava, každý LIR bude mít (vzhledem k RIR) jednu podsíť. Takže LIR má adresu s delším prefixem. LIR má také své zákazníky, kterým provede totéž – každý zákazník „zdědí“ prefix od LIR plus část specifickou pro daného zákazníka.

¹Zdroj: <http://whitengreen.com/blog-1124-how-to-trace-and-locate-ip-addresses>

**Příklad**

Na jednoduchém příkladu si ukážeme, jak se CIDR projevuje na routerech (poněkud netypicky na soukromých adresách). Předpokládejme, že (pro zjednodušení) máme kromě jiných tyto podsítě, které se případně ještě dělí na další podsítě:

Adresa podsítě	binárně druhý oktet
172.16.0.0/20	...0001 0000....
172.20.0.0/20	...0001 0100....
172.24.0.0/20	...0010 0100....
172.28.0.0/20	...0010 1100....

Pokud máme na některém routeru například první dvě podsítě na stejném portu (cesty k nim se dělí až později), pak můžeme mít ve směrovací tabulce pro obě podsítě jediný záznam, přičemž u nadsíťové adresy v tomto záznamu budeme brát v úvahu pouze ty bity (ve směru zleva), ve kterých jsou shodné. Jak vidíme, prvních 13 bitů (8 + 5) je u prvních dvou podsítí stejných, přičemž právě ve třináctém bitu se liší od následujících dvou sítí, takže bychom mohli souhrnně první dvě podsítě reprezentovat adresou 172.16.0.0/13. Adresa vypadá skoro stejně jako adresa první podsítě, rozdíl je jen v délce prefixu – díky tomu sem spadá i rozsah pro druhou podsít.

Podobně bychom mohli adresou 172.24.0.0/13 reprezentovat souhrnně třetí a čtvrtou síť, kdyby na daném routeru byly za stejným portem.



Tedy CIDR má za úkol především zredukovat směrovací tabulky routerů, zejména routerů na páteřích velkých sítí. Bez tohoto postupu by směrovací tabulky routerů na páteři Internetu byly prakticky nepoužitelné (resp. ani velmi velká paměť by nestačila a procesor routeru by nestíhal).

 *Sumarizace (také agregace) cest, resp. nadsíťování (supernetting)* je jiný název pro výše popsaný postup. Pokud do více (pod)sítí vede cesta přes tentýž port routeru, není důvod mít pro každou zvlášť ve směrovací tabulce samostatný řádek. Všechny budou shrnuty do jednoho řádku, přičemž se zkrátí prefix (posune se hranice síťové části adresy směrem doleva) tak, aby v sobě zahrnoval všechny tyto podsítě. Klidně i před původní hranici tříd A/B/C. Tomuto souhrnnému řádku reprezentujícímu více (pod)sítí se říká *CIDR blok*. Ovšem podmínkou je, aby doopravdy fyzické propojení sítí odpovídalo hierarchii IP adres.

**Poznámka:**

Rozdíl mezi podsíťováním (včetně VLSM) a nadsíťováním (CIDR) je kromě jiného v tom, že u podsíťování posouváme hranici mezi (pod)síťovou a hostitelskou částí směrem doprava, kdežto u nadsíťování ji posouváme doleva. Podsíťování se týká každého prvku hierarchie (firma dostane od ISP adresní rozsah a ten si může libovolně rozčlenit), nadsíťování se týká spíše ISP provozujících hodně vytížené směrovače (potřebují redukovat směrovací tabulky právě pomocí CIDR).



 Právě systém CIDR přišel se zjednodušeným označováním hranice mezi síťovou a hostitelskou částí pomocí lomítka a délky prefixu, tedy v třídním směrování bychom teoreticky měli používat zápis pomocí masky. Zápisu s lomítkem a délkou prefixu se říká *CIDR notace*. Systém adresních rozsahů IPv6 je již přirozeně typu CIDR.

6.2 Překlad adres

6.2.1 NAT na hranici klientské sítě

 Jak už bychom měli vědět, rozlišujeme *veřejné* a *soukromé adresy*. Veřejné adresy jsou celosvětově unikátní a lze je používat bez jakýchkoliv omezení. Soukromé adresy jsou unikátní jen v rámci sítě (nebo sítí) dané organizace a na Internet se nesmí dostat.

Adresy pro privátní sítě najdeme v RFC 1918, používají se pro ně tyto adresní rozsahy:

- 10.x.x.x, tedy 10.0.0.0/8,
- 172.16.x.x až 172.31.x.x, tedy 172.16.0.0/12,
- 192.168.0.x až 192.168.255.x, tedy 192.168.0.0/16.

Takže například pokud je prvním oktetem adresy číslo 10, jde o soukromou adresu.

Jejich hlavním účelem je šetření rozsahů veřejných IP adres. Druhým účelem je mírné zvýšení zabezpečení, protože soukromé adresy je třeba na hranici sítě překládat na veřejné (nebo na soukromé fungující v nadřazené síti), což může být chápáno jako dodatečný ochranný mechanismus.

Soukromé adresy je třeba překládat na hranici sítě provozovatele, ale ve skutečnosti tento překlad můžeme využívat i u jiných typů adres. Překlad v základu spočívá v tom, že v paketu zaměníme jednu IP adresu za jinou.



Poznámka:

Abychom si nepletli pojmy:

- veřejná adresa je globálně platná, může do Internetu,
- soukromá adresa je pouze lokálně platná, nemůže do Internetu.

Dále rozlišujeme staticky a dynamicky přidělenou adresu:

- statická adresa je pevně určena danému zařízení, jiné zařízení ji nemůže získat,
- dynamická adresa je momentálně danému zařízení přidělena, ale zítra ji může mít úplně jiné zařízení.

Adresa zařízení může být soukromá a zároveň dynamická nebo soukromá a zároveň statická. Veřejné adresy bývají statické, ale ne každá statická adresa je veřejná.



 NAT (*Network Address Translation*) je mechanismus překladu IP adres na hranici sítě. V paketu můžeme překládat jen zdrojovou adresu nebo jen cílovou adresu nebo obojí, obvykle se kombinuje překlad zdrojové adresy v jednom směru a překlad cílové adresy v druhém směru.

Rozlišujeme tyto druhy NAT:

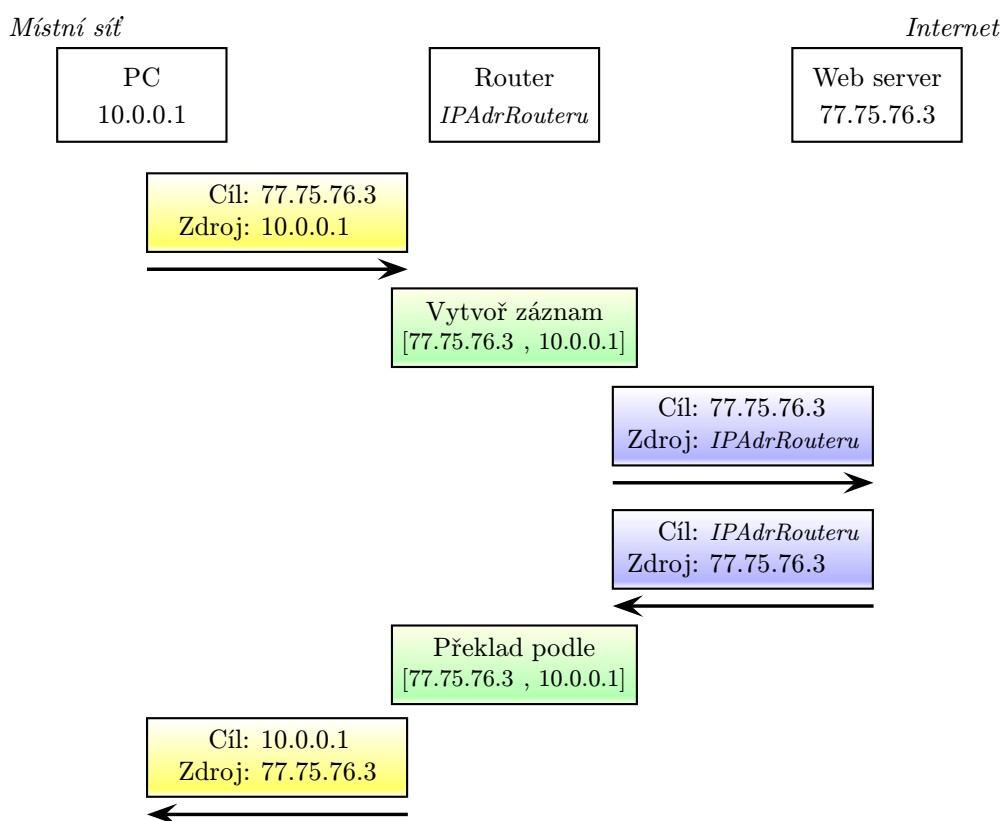
- Statický NAT (bezstavový) – je pevně určeno, kterou adresu z vnitřní sítě překládáme na kterou adresu platnou ve vnější síti a naopak.
- Dynamický NAT (stavový, Masquerade) – dynamicky se určuje dvojice adres pro překlad, vyžaduje uložení záznamu o překladu, který se využije pro zpětný překlad.
- PAT (Port Address Translation, přetížení NAT) – přidává k dynamickému NAT dodatečný mechanismus rozlišení jednotlivých konverzací.

 Při použití *statického NATu* zaměníme v odchozím provozu v paketu zdrojovou IP adresu našeho zařízení za jinou, pod kterou má být toto zařízení viditelné „venku“, a naopak v příchozím provozu

provedeme na (tentokrát cílovou) IP adresu tohoto zařízení zpětný překlad. Statický NAT vyžaduje, abychom měli pro dané zařízení vyhrazeny dvě adresy – jednu pro vnitřní síť a druhou pro vnější síť. Typicky se používá pro skrývání IP adres našich serverů. Server je „venku“ viditelný pod jinou IP adresou než jakou má ve skutečnosti. Dalším důvodem využití statického NATu je dočasná změna v síti třeba z důvodu údržby.

 **Dynamický NAT** (Masquerade, českým patvarem „maškaráda“) je již stavová konfigurace. Používá se v případě, že vnitřní adresy, které překládáme, jsou dynamické, a tedy si je nelze „staticky zapamatovat“. Abychom ušetřili adresní rozsah, máme ve vnitřní síti dynamicky přidělované adresy a směrem ven je překládáme na jedinou adresu (společnou pro všechno dynamické z vnitřní sítě).

NAT server (to bývá router provádějící NAT překlad) si vede tabulku spojení, ve které si pro každé spojení mezi vnitřní a vnější sítí zaznamenává adresu zařízení ve vnitřní síti a adresu zařízení z vnější sítě (typicky některého serveru na Internetu, se kterým je spojení navazováno). Princip je naznačen na obrázku 6.2.



Obrázek 6.2: Zjednodušené schéma dynamického NAT

Jednoznačnost je zajištěna v případě, že každé zařízení z vnitřní sítě komunikuje s jiným serverem. Problém nastává tehdy, když víc zařízení z vnitřní sítě chce komunikovat s tímtež „vnějším“ serverem, pak by záznamy v tabulce spojení přestaly být jednoznačnými. Ve směru od klienta k serveru by se paket poslat dal, ale odpověď od serveru bychom nemohli doručit příslušnému klientovi, protože bychom nevěděli, kterému.

 **PAT** (překlad portů, přetížený NAT) řeší právě tento problém. Do IP paketů se obvykle zapouzdřují TCP nebo UDP segmenty, ve kterých používáme zdrojové a cílové číslo portu. Klientská čísla portů bývají

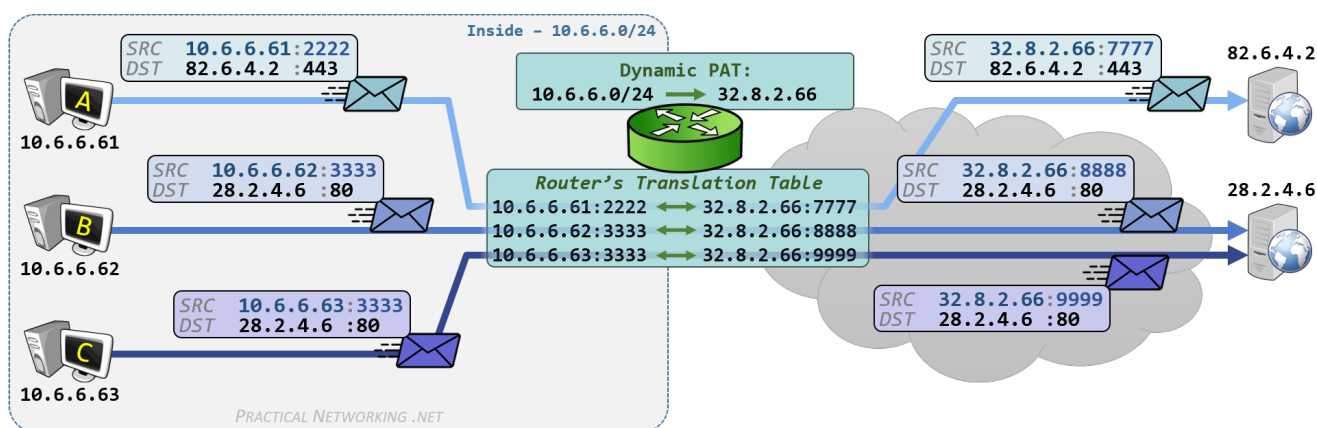
dynamická, tj. je velká pravděpodobnost, že dva různí klienti z naší sítě komunikující s tímto serverem budou mít tato čísla odlišná.

Takže místo abychom do tabulky ukládali jen IP adresy, přidáme i čísla portů (připomeňme si, že například socket je identifikován dvojicí IP adresa + číslo portu). Pokud by přesto nastala situace, že by dva klienti z vnitřní sítě při komunikaci s tímto serverem použili stejné zdrojové číslo portu, na NAT serveru by byla přeložena nejen IP adresa, ale i číslo portu tak, abychom zajistili jednoznačnost.



Příklad

Na obrázku 6.3 je znázorněna situace, kdy tři různá zařízení z vnitřní sítě (se soukromými adresami) komunikují s dvěma servery na internetu. Ve všech odchozích paketech je soukromá adresa odesílatele překládána na adresu 32.8.2.66. V našem případě jde o veřejnou adresu, ale ve skutečnosti to může být (a taky bývá) soukromá adresa z rozsahu určeného pro síť za routerem (jde o síť našeho ISP).



Obrázek 6.3: Překlad portů – PAT²

S překladem čísel portů to obrázek poněkud přehání, porty se v záhlaví paketu mění jen v případě, že v tabulce existuje kolidující záznam se stejným číslem portu. Dynamické porty se generují jako náhodná čísla, pravděpodobnost shody není 100 % (ale ani 0 %).



Další informace:

Mechanismus NAT je obvykle chápán jako berlička pro IPv4, ale i v IPv6 pro něj mohou existovat důvody: <http://www.lupa.cz/clanky/10-duvodu-proc-mit-nat-na-ipv6/>

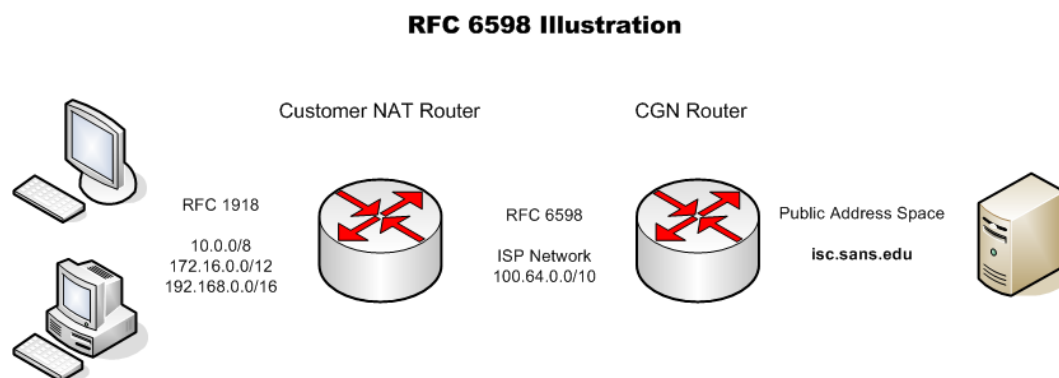


6.2.2 Carrier-Grade NAT se sdíleným adresovým prostorem

RFC 6598 určuje další rozsah soukromých adres, a to „Shared Address Space“, a to 100.64.0.0/10. Tento rozsah adres není určen pro klientské lokální sítě, byl navržen pro poskytovatele internetu (ISP) pro využití v mezilehlých sítích propojujících lokální sítě jejich klientů, tedy na ISP páteři.

²Zdroj: <https://www.practicalnetworking.net/series/nat/dynamic-pat/>

³Zdroj: <https://isc.sans.edu/forums/diary/RFC+6598+Carrier+Grade+NAT/20777/>

Obrázek 6.4: Sdílený adresní prostor podle RFC 6598³

Na obrázku 6.4 vlevo je klientská lokální síť využívající některý z rozsahů podle RFC 1918, například adresy 10.0.0.0/8. Uprostřed je síť poskytovatele internetu, a zařízení na hranici mezi těmito dvěma sítěmi (zde označené Customer NAT router) má síťová rozhraní v obou těchto sítích. Z toho vyplývá, že v žádném případě nesmí nastat kolize adresních rozsahů těchto sítí, jinak by router nedokázal vytvořit použitelnou směrovací tabulku (pokud by přišel paket patřící do průniku adresních rozsahů, router by nedokázal určit, na které rozhraní tento paket poslat).

Jak zajistit, aby ke kolizi nedošlo, tedy aby průnik byl prázdná množina? ISP by v páteřní síti mohl používat veřejné IP adresy, jenže ty jsou vzácné. Při větším počtu klientů (a jejich routerů) by to mohl být vážný problém, navíc plýtvání, protože tady přece ještě nejsme na Internetu. Jenže soukromé rozsahy podle RFC 1918 taky nemůže používat, protože nemůže tušit, jestli se klient ve své vlastní LAN síti někdy nerozhodne náhodou pro tentýž rozsah, jaký si zvolil ISP. Tak co s tím? Nový rozsah, který není ani veřejný, ani určený pro soukromé LAN síť.

Jak na obrázku vidíme, byl opravdu zvolen rozsah podle RFC 6598 určený pro sdílená prostředí (tj. páteře ISP). Ovšem pak je nutné překládat nejen mezi první a druhou sítí, ale taky mezi sítí ISP a další fází cesty. Na zařízení na obrázku označeném CGN router tedy bude taky probíhat NAT, ovšem v masivním měřítku. Tudy poteče provoz všech klientů daného ISP, tedy musíme zde mít dostatečně dimenzovaná zařízení co se týče jak provozu, tak i paměti pro masquerade. Překlad do následující sítě probíhá na veřejné IPv4 adresy nebo na IPv6 adresy (DualStack).

 Tento typ překladu adres se nazývá *Carrier-Grade NAT* (CGN): adresy používané ve sdíleném adresovém prostoru podle RFC 6598 jsou překládány na daný typ veřejných IP adres (IPv4 nebo IPv6).

6.2.3 Koexistence IPv4 a IPv6

V současné době je většina zařízení tzv. Dual Stack, tedy zvládají IPv4 i IPv6 a taky mají funkční oba typy těchto adres, ale ne všechna, a taky ne vždy je podpora obou typů adres na všech úsecích cesty paketu. Jaké jsou tedy možnosti koexistence obou verzí síťových zásobníků TCP/IP?

 **Dual Stack.** Nejjednodušší situace je tedy případ, kdy obě komunikující strany a taky všechny úseky cesty mezi nimi podporují obě verze. Pak má odesílatel na výběr.

 **IP Tunneling.** Standard RFC 4213 (Basic Transition Mechanisms for IPv6 Hosts and Routers) předepisuje způsob tunelování IPv6 provozu do IPv4 tunelu na takovém úseku cesty, kde IPv6 není podporováno. Je to jeden z typů tunelů typu IP-in-IP, tedy IP paket uvnitř jiného IP paketu.

Pojem „tunel“ jednoduše znamená, že původnímu paketu je přidáno nové záhlaví (obvykle jiného protokolu nebo jiné verze téhož protokolu) takové, které bez problémů může projít úsekem cesty, kde by původní paket projít nemohl. V našem případě má původní paket IPv6 záhlaví a na začátku „problematického“ úseku je mu přidáno nové IPv4 záhlaví, přičemž to původní zůstává zachováno, jen je skryto za novým. Na konci IPv4 úseku se provede zpětná operace, zbytek cesty absolvuje původní IPv6 paket.

Dotyčný standard stanovuje, jaké údaje se mají doplnit do přidávaného IPv4 záhlaví. Do pole *Protocol* se doplní číslo 41 znamenající zapouzdřený paket IPv6. Zdrojová IPv4 adresa může být určena ze zdrojové IPv6 adresy, nebo se jednoduše použije podobný princip jako u dynamického NAT (tedy se použije IPv4 adresa vstupního bodu tunelu, kde také tento překlad probíhá). Jako cílová IPv4 adresa se použije adresa druhého konce tunelu.

 **NAT/PT.** Standard RFC 2766 (Network Address Translation – Protocol Translation) popisuje způsob překladi IPv6 záhlaví na IPv4 záhlaví. Zatímco předchozí popsany mechanismus nijak nezasahuje do původního záhlaví a jen přidává nové, tento mechanismus odstraní původní IPv6 záhlaví a vytvoří nové. Standard opět určuje, co dosadit do jednotlivých polí v IPv4 záhlaví, ale tentokrát se nepočítá s tím, že někde jinde na cestě by se rekonstruoval původní IPv6 paket.

Tento mechanismus se už nedoporučuje používat, zejména překládání IPv6 adres na mnohem kratší IPv4 adresy je velmi problematické. Ostatně zmíněný standard byl „Obsoleted by“ standardem RFC 4966 s názvem „Reasons to Move the Network Address Translator – Protocol Translator (NAT-PT) to Historic Status“.



Poznámka:

V případě použití Dual Stack je paket stejný po celou cestu. Druhý mechanismus (IP Tunneling) se používá tehdy, když je třeba překlenout krátký úsek cesty pomocí tunelu, třetí se používá (spíše používal) v případě, že celý zbytek cesty včetně cílového zařízení nebude možné použít původní záhlaví.



6.3 Protokol ICMP a zprávy

6.3.1 Účel protokolu ICMP

Protokol ICMP (Internet Control Message Protocol, tedy protokol pro posílání řídicích zpráv) je také protokolem vrstvy L3. Jeho účelem je zajišťovat posílání krátkých zpráv, obvykle reprezentovaných jedním nebo dvěma čísly (každé takové číslo má určitý význam), výjimečně ještě s dodatečnou datovou informací. Typicky se takto posílají informace o chybách, problémech na cestě nebo jiná upozornění.

 Každá zpráva posílaná pomocí ICMP má své *číslo*. V následující tabulce jsou některé běžné hodnoty pro ICMPv4 uvedeny:

 Samotná zpráva je v některých případech příliš obecná, tedy ji může upřesnit kód zprávy. Například pro zprávu typu 3 (*Destination Unreachable*) existuje několik kódů, které upřesňují, proč je cílové zařízení nedostupné:

- kód 0 – cílová síť je nedostupná (router tuto síť nezná a nedokáže ji směrovat),
- kód 1 – cílová stanice je nedostupná (v zadané síti nelze najít dotyčnou stanici),
- kód 2 – cílový protokol je nedostupný (v poli *Protokol* je neznámá hodnota),

Číslo	Význam
8	<i>Echo Request</i> – žádost o odezvu (posíláme, když příkazem <code>ping</code> zjišťujeme, zda je konkrétní zařízení dostupné)
0	<i>Echo Reply</i> – odpověď na <i>Echo Request</i> (8) ve smyslu „ano, jsem tady“
3	<i>Destination Unreachable</i> – zpráva o nedostupnosti cíle (pokud router obdrží IP paket, který nedokáže doručit, pošle ICMP paket se zprávou 3 zdrojovému zařízení)
11	<i>Time Exceeded</i> – vypršel čas čekání; buď hodnota TTL v IP paketu klesla na 0, vypršel časový limit při čekání na zbývající fragmenty IP paketu apod. (posílá router zdroji IP paketu)
12	<i>Parameter problem</i> – v IP paketu (většinou v jeho záhlaví) je nějaký problém, který nelze oznámit jinou ICMP zprávou
17	<i>Address Mask Request</i> – žádost o informaci o masce (pod)sítě
18	<i>Address Mask Reply</i> – odpověď na <i>Address Mask Request</i> (17)

Tabulka 6.1: Některé zprávy protokolu ICMPv4

- kód 3 – cílový port je nedostupný (tím je myšleno číslo portu uvedené v zapouzdřeném segmentu transportní vrstvy),
- kód 4 – paket je větší než MTU na cestě, ale nelze fragmentovat, protože je nastaven příznak DF (Do not Fragment),
- atd. (pro zprávu 3 je celkem 15 různých kódů).

Pro zprávy typu 11 (*Time Exceeded*) se obvykle používají kódy 0 (hodnota pole TTL je 0) nebo 1 (překročen čas čekání na zbývající fragmenty).

 Takže zatím máme dvojici číslo zprávy + kód. Zpráva ale může být upřesněna ještě jinak, například *jako data* může být přiložena část IP paketu, na který je takto reagováno (přiloží se IP záhlaví a prvních 8 oktetů datové části).

Poznámka:

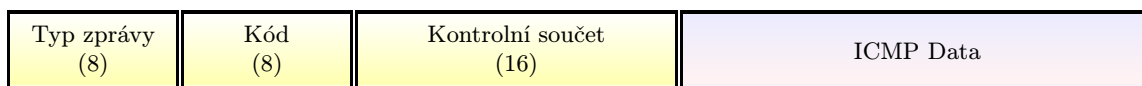
Ve skutečnosti mají datovou část i ty ICMP pakety, které ji vlastně ani nepotřebují – přidávají tam výplň, protože na nižších vrstvách bývá stanovena určitá minimální délka pro zapouzdřená data (payload). S tím se setkáváme například u ICMP zprávy *Echo Request* (8).

 ICMP paket má velmi jednoduchou strukturu – žádné adresy, délka paketu apod., záhlaví má jenom tři pole:

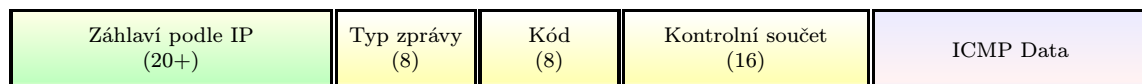
- *Typ zprávy* (ICMP Type, 8 bitů) – určení typu zprávy, například číslo 8 pro *Echo Request*,
- *Kód* (Code, 8 bitů) – upřesňující kód,
- *Kontrolní součet* (Checksum, 16 bitů) – kontrolní součet přes předchozí dvě pole záhlaví.

Formát celého ICMP paketu vidíme na obrázku 6.5.

 ICMP paket se sice při odesílání předává na vrstvu L2, ale ještě předtím se zabalí do IP paketu – samotný ICMP paket nemá ve svém záhlaví ani adresy ani jiné důležité položky. Na obrázku 6.6 vidíme formát paketu do zapouzdření do IP.



Obrázek 6.5: Formát ICMP paketu



Obrázek 6.6: Formát ICMP paketu zapouzdřeného v IPv4 paketu

Některé typy zpráv používají $TTL = 1$, ale většina využívá hodnotu TTL nastavenou v odesílajícím systému (tj. obvykle 128 nebo 64).

Po zapouzdření do IP je paket zpracován na L2 jako každý jiný paket, tj. například je zapouzdřen do paketu typu Ethernet II a poslán po ethernetové síti.

6.3.2 ICMP verze 6

 Pokud používáme protokol IPv6, pak je třeba také používat ICMPv6. Oproti ICMPv4 byly některé zprávy přidány a taky se rozšířily oblasti využití tohoto protokolu (což souvisí) – v protokolu ICMPv6 se shrnuly role tří původních protokolů: ICMPv4, ARP a IGMP.

Jenže zdaleka nešlo jen o přidání typů zpráv pro nové způsoby využití protokolu, mnohé původní typy byly přechíslovány. Už u ICMPv4 platilo, že některé zprávy jsou chybové a jiné informační, ICMPv6 toto rozlišení přenesl i do číslování typů zpráv – pole pro typ zprávy je taky 8bitové, ale první (nejlevější, nejvíce významný) bit je nastaven na 0 u chybových zpráv a na 1 u informačních zpráv.

Číslo	Význam
128	<i>Echo Request</i> – žádost o odezvu, původně číslo 8
129	<i>Echo Reply</i> – odpověď na <i>Echo Request</i> (128), původně číslo 0
1	<i>Destination Unreachable</i> – zpráva o nedostupnosti cíle, původně číslo 3
3	<i>Time Exceeded</i> – vypršel čas čekání, původně číslo 11
4	<i>Parameter problem</i> – v IP paketu je nějaký problém, který nelze oznámit jinou ICMP zprávou, původně číslo 12

Tabulka 6.2: Některé zprávy protokolu ICMPv6, které byly i u ICMPv4

 V tabulce 6.2 je několik typů zpráv, které najdeme v obou verzích. Všimněte si, že tyto zprávy mají v ICMPv6 jiná čísla než v ICMPv4, a že při novém číslování zachovávají pravidlo, kdy nejvíce významný bit z 8 bitů prvního pole je nastaven na 0 u chybových zpráv (tj. číslo typu je v rozmezí 0–127, poslední tři řádky tabulky) a na 1 u informačních zpráv (rozmezí 128–255, první dva řádky tabulky). Také v kódech pro některé konkrétní typy zpráv jsou změny (ale třeba kódy pro zprávu Time Exceeded jsou stejné, jen číslo typu se změnilo).

Nově přidané typy zpráv souvisejí obvykle buď s mapováním IP adres na MAC adresy (což u starší verze byla práce protokolu ARP) nebo se správou skupin (původně v protokolu IGMP), a taky pár typů zpráv pro běžné nebo mobilní sítě. Nejdůležitější z nich jsou v tabulce 6.3.

Číslo	Význam
2	<i>Packet Too Big</i> – v případě, že router dostane k přeposlání příliš velký paket (větší než hodnota MTU na portu pro odeslání), paket zahodí a původnímu odesílateli pošle tuto ICMP zprávu
134	<i>Router Advertisement</i> (vizitka routeru) – routery v pravidelných intervalech posílají tuto ICMP zprávu s informacemi o sobě a síti, například adrese sítě a délce prefixu
133	<i>Router Solicitation</i> (žádost o vizitku routeru) – stanice si touto zprávou může na routeru vynutit vyslání zprávy Router Advertisement (134) mimo interval
135	<i>Neighbor Solicitation</i> (žádost o oznámení souseda) – máme IP adresu některého zařízení v sousedství, touto zprávou se ptáme svého okolí na jeho MAC adresu (tj. „Kdo má tuto IP adresu?“, soused, který pozná svou IP adresu, odpoví následující zprávou:
136	<i>Neighbor Advertisement</i> (oznámení souseda) – v reakci na předchozí zprávu (pokud poznám svou IP adresu) oznámím svou MAC adresu
130	<i>Multicast Listener Query</i> (dotaz na členy skupiny) – posílá router při ověřování, kdo z jeho sítě patří do konkrétní skupiny a tedy odebírá pakety adresované na danou multicast adresu
143	<i>Multicast Listener Report</i> (oznámení členství ve skupině) – posílá zařízení, když se přihlašuje do skupiny (informuje router, že chce dostávat příslušné pakety) nebo jako odpověď na předchozí zprávu

Tabulka 6.3: Některé zprávy protokolu ICMPv6, které nejsou v ICMPv4

 Zajímavá je například zpráva *Packet Too Big* (2), která v předchozí verzi nebyla, třebaže by se v podstatě hodila i tam. Pokud router dostane k přeposlání paket větší než je hodnota MTU na následující cestě, chová se u různých verzí IP odlišně.

- Je to paket protokolu IPv4: buď se ho pokusí fragmentovat, nebo (když to nejde, příznak DF = 1) paket zahodí a jeho původci pošle ICMPv4 zprávu *Destination Unreachable* (3).
- Je to paket protokolu IPv6: v každém případě ho zahodí a jeho původci pošle ICMPv6 zprávu *Packet Too Big* (2).

**Další informace:**

Seznam všech typů ICMP zpráv a k nim příslušných kódů najdete na

- IPv4: <https://www.iana.org/assignments/icmp-parameters/icmp-parameters.xhtml>
- IPv6: <https://www.iana.org/assignments/icmpv6-parameters/icmpv6-parameters.xhtml>



6.4 Správa skupin

Jak už víme, také v adresním prostoru IPv4 a IPv6 existují skupinové adresy, a tedy musí existovat protokol na jejich správu. Skupina v prostoru IP je především dána konkrétní skupinovou IP adresou.



Správa skupin obnáší:

- evidování portů s členy dané skupiny na aktivních síťových prvcích (například routerech),
- možnost ohlásit příslušnost ke konkrétní skupině (zajistit si členství ve skupině),
- ověřování trvání členství ve skupině.

Na aktivních síťových prvech vrstvy L3 je třeba evidovat, za kterými porty jsou členové patřící do konkrétní skupiny (nikoliv samotné členy, to je zbytečné), protože na tyto prvky je třeba rozesílat pakety, které mají danou skupinovou adresu jako cílovou. Router (nebo obdobné zařízení vrstvy L3) se může pravidelně na portech dotazovat, zda se za nimi nachází nějaké zařízení patřící do dané skupiny, a pokud na daném portu žádné zařízení nereaguje, port ze skupiny odstraní (tj. nebude na něj zasílat pakety s danou skupinovou adresou jako cílovou).

Pokud síť prochází multicast paket, na jednotlivých routerech, přes které jde, se podle cílové (multicastové) adresy ověří, na které porty má být tento paket replikován. Pole TTL se na každém routeru dekrementuje, zachází se s ním stejně jako v jakémkoliv jiném paketu.

Ať už jde o IPv4 nebo IPv6, vždy máme vyhrazen určitý rozsah pro skupinové adresy, a dále existují některé rezervované skupinové adresy – typicky adresa pro skupinu všech zařízení podporujících protokol IP dané verze, adresa pro skupinu všech routerů, atd. Multicasty se také využívají pro distribuci multimediálních dat (v komunikaci s více cíli – multipoint, např. IPTV nebo Video-on-Demand), zpráv, bezpečnostních záplat, atd.

 V mechanismu IPv4 je za správu skupin odpovědný *protokol IGMP* (Internet Group Membership Protocol). Struktura IGMP paketu je vcelku podobná ICMP, taky se podobným způsobem používá, včetně zapouzdřování do IP paketu. Tak jako ICMP má zprávy (ICMP Echo Request apod.), IGMP má zprávy týkající se skupin, například *Membership Query* (17) a *Membership Report* (22).

Skupinové adresy mají vyhrazen adresní rozsah 224.0.0.0 až 239.255.255.255, přičemž místní skupiny mají adresy 224.0.0.x (v obalovém IP paketu je TTL při odesílání nastaveno na 1). S některými vyhrazenými skupinovými adresami jsme se seznámili v předchozím textu.



Další informace:

Seznam registrovaných multicast adres je dostupný na

- <https://www.iana.org/assignments/multicast-addresses/multicast-addresses.xhtml>
- <https://www.iana.org/assignments/ipv6-multicast-addresses/ipv6-multicast-addresses.xhtml>



 V mechanismu IPv6 roli správy skupin přebírá *protokol ICMP*, přičemž pro účely správy skupiny jsou do něj přidány nové typy zpráv (nahrazující to, co původně bylo v IGMP):

- *Multicast Listener Query* (130), kterou posílá router na downlink porty pro ověřování, zda na těchto portech existují zařízení patřící do konkrétní skupiny,
- *Multicast Listener Report* (143), kterou posílá zařízení, když chce informovat router (nebo routery na cestě), že chce odebírat pakety pro konkrétní skupinu (přihlášení do skupiny).

U IPv6 jsou pro skupiny rezervované adresy začínající prefixem `ff00::/8`.

 Ve správě skupin se ve skutečnosti používají i další protokoly, například obdoby směrovacích protokolů pro distribuci obsahu multicastových tabulek mezi routery.



Další informace:

- <http://www.samuraj-cz.com/clanek/tcpip-skupinove-vysilani-ip-multicast-a-cisco/>
- https://www.ipv6.cz/cs/skupinove_adresy_multicast
- <http://www.lupa.cz/clanky/ipv6-myty-a-skutecnost-dil-vii-podpora-multicast-a-anycast-provozu/>



6.5 Objevování sousedů

6.5.1 Tabulky sousedů

Představme si tuto situaci: známe svou IP adresu, masku (či prefix) a IP adresu brány. Potřebujeme poslat paket do cizí sítě, a tedy přes bránu. Jenže „kterým směrem“ je brána? Přes které síťové rozhraní je dostupná a jaká je její MAC adresa? Vlastně podobný problém musíme řešit i tehdy, když chceme poslat IP paket někomu do vnitřní sítě. Třeba znám IP adresu dotyčného zařízení, ale jakou tedy má MAC adresu?

 Mapování IP adres na MAC adresy provádí v případě IPv4 *protokol ARP* (Address Resolution Protocol) a v případě IPv6 *protokol NDP* (Neighbor Discovery Protocol). Jejich hlavní funkce:

- vést tabulku sousedů,
- pomocí ARP/NDP dotazů tuto tabulku doplňovat.

 V obou případech si zařízení vede *tabulku sousedů* s informacemi o nejbližším síťovém okolí (sousedech). V tabulce sousedů máme pro každého „sousedu“ tyto informace:

- IP adresa (IPv4 nebo IPv6),
- MAC adresa,
- typ (statická/trvalá nebo dynamická/zjištěná).

První dva údaje nám mapují IP adresu souseda na jeho MAC adresu, třetí údaj určuje, jak se dotyčný údaj dostal do tabulky (zda byl staticky vložen nebo zjištěn dynamicky dotazováním). Pro každé síťové rozhraní s vlastní IP adresou se vede jedna tabulka (včetně virtuálních rozhraní, pokud máme nainstalován virtualizační software).



Poznámka:

Koncová zařízení v síti mají v této tabulce především jeden důležitý záznam – adresu brány, na kterou směřují většinu svého provozu.



 V případě Windows vyšších verzí máme ARP tabulku i NDP tabulku poněkud přepřádané – najdeme tam kromě brány taky multicast adresy a v případě ARP taky broadcast adresy včetně mapování na MAC adresy. Obsah tabulek zobrazíme takto:

- `arp -a` pro ARP tabulku v IPv4,
- `netsh interface ipv6 show neighbors` pro NDP tabulku v IPv6
- v PowerShellu: `get-NetNeighbor` (nefunguje ve Windows 7)

(příkaz `netsh` se dá použít i pro zobrazení ARP tabulky v IPv4, pokud místo `ipv6` napíšete `ipv4`).

Tyto příkazy se také dají použít k ovlivňování tabulky sousedů (přidávání nových záznamů nebo jejich odstraňování), ve většině případů však tyto úlohy zvládají dotyčné protokoly dynamicky.

 V Linuxu a jiných UNIXových systémech se pro zobrazení ARP/NDP tabulky používá příkaz `ip neigh show`

(tedy můžeme napsat `ip neighbor show` nebo `ip n s` nebo něco mezi tím, nejlepší je volit kompromis – nepsat příliš a zároveň psát tolik, aby to bylo srozumitelné). Na výstupu máme obvykle pro každé síťové

rozhraní právě jeden údaj – záznam pro bránu, nejsme zahlcováni dalšími informacemi. Také tento příkaz (`ip neigh` s některým podpříkazem) je možné použít pro ovlivňování obsahu tabulky.

Také v UNIXových systémech existuje příkaz `arp` (pro zobrazení tabulky stačí zadat bez parametrů), ten můžeme použít tehdy, když nám jde opravdu jen o IPv4 adresy.



Poznámka:

Které IPv6 adresy mají být takto překládány? Předně unicast adresy, ať už globální nebo ULA či LLA, patřící do místní sítě (včetně brány), ale také anycast adresy (které ve skutečnosti patří také do rozsahu unicast adres). Zatímco „normální“ unicast adresy se vždy jednoznačně překládají na jedinou MAC adresu, v případě anycast adres to je složitější – v síti můžeme mít i několik zařízení odpovídajících na jednu anycast adresu, tj. jednu anycast adresu ve skutečnosti lze přeložit na několik různých MAC adres.

Mechanismus NDP tedy předně musí zajistit, aby byla vybrána vhodná MAC adresa, a v případě, že tato vybraná přestane vyhovovat (například zařízení přestane být dostupné), musí být zajištěno přemapování na jinou MAC adresu.



6.5.2 K protokolům

Úkolem protokolů ARP a NDP není jen vést tabulku sousedů, ale také v případě potřeby se dynamicky poptávat na to, co by mělo v této tabulce být. Existují tedy pakety posílané dotyčným protokolem.

 *ARP pakety* (tedy pro IPv4) se zapouzdřují přímo do rámců Ethernet II nebo rámců bezdrátové sítě Wi-fi. Rozlišujeme ARP dotaz a ARP odpověď.

Pokud zařízení potřebuje zjistit MAC adresu k určité IP adrese, vyšle ARP paket s dotazem „Who has xxx? Tell yyy.“ (Kdo má MAC adresu xxx? Sdělte to na MAC adresu yyy). Protože MAC adresáta neznáme (teprve ji zjišťujeme), dotaz se odesílá jako univerzální broadcast (tj. na adresu `FF-FF-FF-FF-FF-FF`). Reaguje pouze to zařízení, které pozná svou MAC adresu, zpět odešle odpověď, tentokrát unicast (protože z dotazu zná MAC adresu tazatele).

Každý ARP paket je stejně dlouhý, ať už jde v kterémkoliv směru. V dotazu i odpovědi jsou všechna pole (vlastní IP a MAC adresa, dotazovaná IP a MAC adresa), ale v dotazu je to, co odesílatel neví, vyplněno nulami.

 *ND pakety* (mechanismus pro IPv6) nemají vlastní specifikaci, posílají se jako ICMPv6 pakety se speciálními typy zpráv a zapouzdřují se do IPv6 paketů (jako každá ICMP zpráva). NDP je komplexnější protokol než ARP, plní více úloh. Z ICMPv6 zpráv se sousedy souvisejí především

- *Neighbor Solicitation* (ICMP zpráva číslo 135) – žádost o oznámení souseda, významem odpovídá ARP dotazu,
- *Neighbor Advertisement* (ICMP zpráva číslo 136) – odpověď na předchozí žádost.

Na rozdíl od ARP se formát a velikost dotazu a odpovědi liší, v dotazu najdeme opravdu jen žádanou adresu. V dotazu je jako adresa cíle použita multicast adresa odvozená od skupinové adresy pro směrovače. Odpověď je poslána buď jako unicast (pokud byla v dotazu IP adresa tazatele) nebo jako multicast na všechny uzly v síti zvládající protokol IPv6 (tj. na `ff02::1`).

NDP posílá zprávu *Neighbor Solicitation* na skupinovou adresu s prefixem `ff02:0:0:0:0:1:ff00::/104`, které se říká *Solicited-node Multicast Address*. Jenže to není celá adresa, ještě se musí na konec doplnit

posledních 24 bitů (tj. 3 oktety) z té adresy, na kterou se tážeme. Takže je vlastně dopočítaná z dotazované adresy, čímž se omezuje počet členů dané multicastové skupiny na (maximálně) jednoho. Ostatní se touto zprávou nemusejí vůbec zabývat.



Příklad

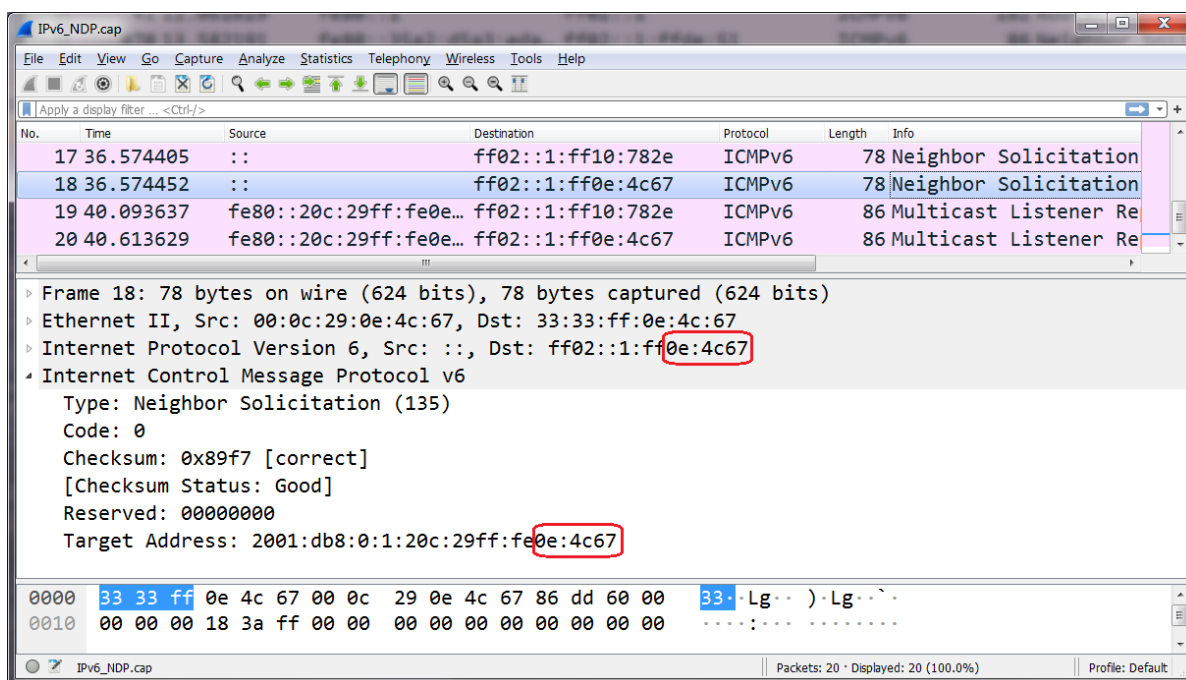
Mechanismus objevování sousedů podle RFC 4861 si ukážeme na příkladu zjištění MAC adresy souseda (v případě, že známe IPv6 adresu) – je to obdoba ARP dotazu. Tento příklad je převzat z webu <https://www.ipv6.cz/>. Postup:

- posílám datagram na skupinovou IPv6 adresu s prefixem `ff02:0:0:0:0:1:ff00:0/104`
- určení celé skupinové IPv6 adresy:
 - znám IPv6 adresu souseda, např. `2001:db8:1:1:22a:fff:fe32:5ed1`
 - z ní vezmu posledních 24 bitů (3 oktety) a dodám do skupinové IPv6 adresy z prvního bodu: `ff02::1:ff32:5ed1`
- na tuto adresu pošlu ICMPv6 zprávu *Neighbor Solicitation* – výzvu sousedovi
- soused odpoví zprávou *Neighbor Advertisement* obsahující jeho MAC adresu



Příklad

Na obrázku 6.7 vidíme paket s dotazem Neighbor Solicitation (screenshot z pcap souboru IPv6_NDP ze serveru packetlife.net). Červeně je vyznačena cílová adresa IP paketu a níže dotazovaná adresa, kterou potřebujeme přeložit na MAC adresu. Jak vidíme, poslední tři oktety (tj. 24 bitů, 6 hexadecimálních číslic) se shoduje, přesně podle výše uvedeného popisu.



Obrázek 6.7: Paket se zprávou Neighbor Solicitation



 Mechanismus NDP slouží i k jiným účelům – komunikaci s routery při zjišťování konfigurace sítě (adresa sítě, délka prefixu apod.), zjišťování duplicitních adres a dalším. Pro každý účel existují zvlášť čísla ICMPv6 zpráv.

Mechanismus zjišťování duplicitních adres se nazývá *Duplicate Address Detection* (DAD). Když si zařízení dopočítá autokonfigurací novou IPv6 globálně platnou adresu (GUA), pak se doporučuje (zejména pokud je použit mechanismus Privacy Extensions s generováním náhodného čísla) zkontrolovat, jestli náhodou v síti není jiné zařízení s toutéž adresou. Zařízení pošle do sítě ND dotaz na svou novou adresu, a pokud se zpět někdo ozve, znamená to kolizi adres a je třeba adresu generovat znovu.

Poznámka:

Otázkou je, na které vrstvě RM ISO/OSI, resp. síťového modelu TCP/IP, protokoly ARP a NDP vlastně pracují. V případě ARP není v standardu o nějaké vrstvě ani slovo. ARP pakety se zapouzdřují do rámců vrstvy L2, takže spíše patří na vrstvu L2 nebo na rozhraní mezi vrstvami L2 a L3, ale rozhodně ne na vrstvu L3.

Protokol NDP využívá ICMP zprávy, takže jednoznačně patří na vrstvu L3.



6.5.3 Bezpečnost

Objevování sousedů má jedno závažné úskalí: pokud obdržíme NS paket (paket s informací o mapování MAC a IP adres souseda), předpokládá se, že údajům v tomto paketu budeme věřit. Jenže tento paket může být podvržený a důsledkem používání této IP adresy je odesílání dat na jiný počítač než předpokládáme, což zvláště ve firemním prostředí může způsobit odcizení dat.

 Pro tento problém existuje několik řešení, z nichž je zajímavý například *protokol SEND* (Secure Neighbour Discovery – RFC 3971), který umožňuje digitálně podepisovat oznámení protokolu NDP. Používá se pro zajištění „objevovací“ a keep-alive komunikace se sousedy.

Možnosti zajištění bezpečnosti v SEND:

- *kryptograficky generované adresy* (CGA) – RFC 3972: koncept soukromého a veřejného klíče, ICMPv6 zprávy jsou digitálně podepsány,
- *certifikační cesty*: ochrana proti falešným směrovačům – řetězce navazujících certifikátů dokazující, že určitá *důvěryhodná* certifikační autorita schválila toto zařízení jako směrovač poskytující dané informace (certifikační autority tvoří hierarchii).

Další informace:

O objevování sousedů najdeme informace například na

- <http://www.abclinuxu.cz/clanky/architektura-ipv6-konfigurace-adres-a-objevovani-sousedu>
- http://www-public.it-sudparis.eu/~lauren_m/articles/M-CGA-Tony-Cheneau-SAR-SSI-2011.pdf



6.6 Jak funguje router

Routery implementují funkcionalitu vrstvy L3, a to softwarově, což znamená, že potřebují velký výpočetní výkon – zpracování IP záhlaví (plus případné související činnosti) je totiž výpočetně mnohem náročnější než zpracování záhlaví ethernetové rámce. V IP záhlaví (kterékoli verze) totiž máme mnohem více polí, tato pole jsou poměrně dlouhá a v některých je třeba jít na úroveň bitů (což znamená použití bitových operací). Routery také obvykle plní různé další funkce, včetně směrování, filtrování provozu, překladu adres, atd.

Proto průměrný router potřebuje dostatečně výkonný procesor (úměrně provozu, který má reálně zvládat) a dostatek paměti (pakety jsou totiž nejdříve uloženy, pak zpracovány a pak teprve odeslány). Typicky mívají routery mnohem méně portů než switche – každý port navíc znamená navýšení množství paketů ke zpracování a tedy navýšení potřeby výkonu.

Protokol IP poskytuje službu typu best-effort, tj. negarantuje doručení. Takže IP paket se může cestou ztratit nebo dokonce zahozen (z různých důvodů: je poškozený, nebo se nevejde do MTU a nelze ho fragmentovat, nelze ho doručit – viz ICMP, router je zahlcen a nestihá přijímat/odesílat, atd.

Z předchozích odstavců vyplývá, že se opravdu může stát, že provoz je nad možností routeru a některé pakety bude nutné zahodit, třebaže by v jiné situaci mohly být doručeny. Jak tedy řešit problém zahlcení routeru?

 Na směrovačích se může používat jednoduchá metoda řízení propustnosti sítě nazvaná *Token Bucket*. Pracuje na principu virtuálního koše (bucket), do kterého jsou konstantní rychlostí neustále doplňovány tokeny („povolenky“). Pokud přeteče, tokeny jsou zahazovány. Jeden token představuje určité množství odesílaných dat (třeba jeden oktet).

Pokud má být zpracován (směrován, přeposlán) paket, nejdříve je zjištěna jeho velikost. Velikost paketu je porovnávána s množstvím odpovídajících tokenů v Token Bucketu. Pokud dostačuje, odstraní se z Bucketu potřebné množství tokenů a paket může být zpracován. Jestliže však nedostačuje, paket bude zahozen.

6.7 Protokoly transportní vrstvy

Na transportní vrstvě pracují i jiné protokoly než jen TCP a UDP.

 *TCP* (Transmission Control Protocol) poskytuje potvrzovanou službu se spojením, garantuje řazení, tedy doručení více úseků dat ve správném pořadí. Používá se pouze pro spojení typu point-to-point, komunikuje vždy jen jedno koncové zařízení s jiným koncovým zařízením.

Výhodou je zabezpečení ve smyslu řízení datového toku a kontroly ztracených segmentů, možnost segmentace dat a mechanismus řazení v cíli, a také možnost lepšího monitoringu datového toku na firewallech (stavová inspekce paketů – lze poznat, které datové jednotky „patří k sobě“).

Nevýhodou je poměrně vysoká režie navazování, udržování a ukončování spojení, při problémech s udržením spojení dokonce komunikace zcela padne. Dá se použít pouze pro unicast komunikaci.

 *UDP* (User Datagram Protocol) poskytuje nepotvrzovanou službu bez spojení, tedy datagramovou službu. Může se používat i pro spojení typu point-to-multipoint. Místo pojmu segment se v případě protokolu UDP často používá spíše pojem datagram.

Výhodou je pružnost a škálovatelnost celého mechanismu, výrazně nižší režie, rychlost. Tam, kde TCP neustále padá nebo vůbec nejde navázat, může UDP projít. UDP může být použit i pro multicast a broadcast komunikaci.

Nevýhodou je, že UDP nemá mechanismus pro členění posílaných dat do segmentů ani potvrzování, tedy chybí možnost kontroly ztracených dat. Taky není možné určovat parametry spojení a řídit datový tok, protože se žádné spojení nenavazuje.

 Z dalších protokolů najdeme na transportní vrstvě například *SCTP* (Stream Control Transmission Protocol), který je podobný TCP, ale dovoluje navázat více nezávislých paralelních spojení (streamů), což snižuje pravděpodobnost ztráty PDU po cestě – každý stream může jít jinou cestou, a pokud má zařízení víc IP adres, lze mezi nimi přecházet. Tento protokol je implementován v jádru mnoha UNIXových systémů (FreeBSD, Linux, atd.), pro Windows existují implementace třetích stran.

 Transportní protokol *RTCP* slouží k rezervaci zdrojů převážně v komunikaci vyžadující upřednostňování při přidělování zdrojů (řízení kvality služby), setkáváme se s ním v některých VoIP technologiích. Těmito technologiemi se budeme zabývat v jedné z dalších kapitol.

 Protokol *DCCP* poskytuje službu se spojením (jako TCP) nebo volitelně bez spojení (jako UDP), ale negarantuje doručení ve správném pořadí ani nepotvrzuje (jako UDP), má jednodušší záhlaví než TCP. Rozlišuje se několik typů DCCP datagramů (DCCP-request pro navázání spojení, DCCP-response jako odpověď, DCCP-data pro běžná data, DCCP-ack a DCCP-DataAck pro potvrzování, dále typy pro uzavření spojení, reset a synchronizaci). Pro určení typu datagramu nepoužívá jednotlivé bity, ale v záhlaví je pole pro určení typu datagramu.

Na rozdíl od TCP tento protokol neřeší ztracené PDU (což znamená, že přenos tímto mechanismem není zdržován), ale dokáže segmentovat, tedy rozdělit posílaná data na více částí. Stejně jako TCP a UDP, první dvě pole jsou zdrojový a cílový port.

DCCP sice negarantuje doručení ani řazení, ale přesto v záhlaví najdeme pole Sequence Number, a v některých typech i Acknowledge Number. Používají se trochu jinak než u TCP, účelem je především zjistit, jestli se cestou neztrácejí datagramy. Zatímco TCP používá čísla ve stylu byte-based (Sequence Number určuje, na které pozici v počtu bytů je právě přenášený blok dat v celém sekvencovaném streamu), čísla v DCCP se v sousedních datagramech zvyšují o 1 a aby nedošlo k přetečení datového typu, ukládají se modulo 2^{48} .

A to je hlavní vlastnost protokolu DCCP – mechanismus řízení zahlcení. Pokud se cestou ztrácejí datagramy, znamená to, že na některých routerech je plný buffer, důsledkem je zahazování provozu, který se „nevejde“. Potom je třeba adaptivně změnit trasu a tím se vyhnout přetíženým linkám.

Protokol DCCP používají například internetová rádia, některé technologie pro internetovou telefonii, on-line videa, ale také on-line hry.



Souhrn kapitoly

V této kapitole jsme se zabývali některými mechanismy probíhajícími na síťové a transportní vrstvě. Zaměřili jsme se zejména na práci s IPv4 a IPv6 adresami včetně CIDR, překlad adres pro lokální sítě i sdílená prostředí, připomenuli jsme si protokol ICMP, navázali jsme správou skupin v prostředí IPv4 a IPv6. Následovalo téma objevování sousedů opět v oblastech IPv4 a IPv6, připomenuli jsme si funkcionality routeru a rozhlédli jsme se mezi protokoly na transportní vrstvě.

Většina kapitoly je typická tím, že se zdůrazňují rozdíly mezi prostředím IPv4 a IPv6, zejména u protokolu ICMP, ve správě IP skupin, získávání adres, mapování IP adres na MAC adresy.



Decentralizované a distribuované systémy

 **Rychlý náhled:** V této kapitole se po probrání pojmů budeme podrobněji zabývat některými síťovými technologiemi, které jsou v principu decentralizované nebo distribuované. Jedná se zejména o směrování, mechanismus DNS, některé druhy řízení kvality služby (QoS), VoIP a videotelefonii.

 **Klíčová slova:** Směrování, autonomní systém (AS), ASN, směrovací algoritmus, algoritmus vektoru vzdáleností (distance-vector), algoritmus stavu spoje (link-state), Dijkstrův algoritmus, směrovací protokol, RIP, OSPF, EIGRP, BGP, DNS, DNS cache, tabulka hostitelů, resolver, doména, zóna, zónový soubor, zone transfer, rekurzivní a iterativní dotaz, DNSSEC, SPF, DKIM, QoS, IntServ, DiffServ, VoIP, videotelefonie, videokonference.

 **Cíle studia:** Cílem v této kapitole je seznámit se s pojmy decentralizovaný a distribuovaný systém, prohloubit si znalosti o nejpoužívanějších směrovacích protokolech a mechanismu DNS, pochopit princip řízení kvality služeb a prioritizace některých typů dat a jejich využití ve VoIP a videotelefonii.

7.1 Typy systémů

 **Centralizovaný systém** je systém s jedinou centrální řídicí jednotkou. Komunikace probíhá obvykle ve formě klient-server, kde server je právě ta centrální řídicí jednotka.

V oblasti počítačových sítí můžeme k centralizovaným architekturám zařadit například

- *protokol RADIUS*, kde centralizovanou jednotkou je RADIUS server, komunikující se svými klienty (například access pointy – přístupovými body bezdrátové sítě),
- *protokol CMIP* (Common Management Information Protocol), což je protokol z ISO/OSI určený pro správu sítě (obdobu SNMP).

 **Decentralizovaný systém** má více řídicích jednotek (serverových/správních uzlů), víceméně rovnocenných (s podobnými funkcemi). Účelem (a výhodou oproti centralizovaným systémům) je zvýšení robustnosti sítě (tutéž funkci poskytuje více „centrálních“ – serverových – uzlů, klientské uzly mohou komunikovat s kterýmkoliv uzlem poskytujícím žádanou službu). Je možné snadno vyřešit situaci, kdy jeden ze serverových uzlů přestane službu poskytovat (například je odpojen od sítě).

Nevýhodou je náročnější implementace, protože serverové uzly musejí být synchronizovány (především jejich databáze).

V oblasti počítačových sítí se s decentralizovanou architekturou setkáme například

- mnohé *směrovací protokoly* jsou decentralizované,
- *protokol SIP* (Session Initiation Protocol), který se používá například u VoIP.

Samotný Internet byl původně decentralizovaný a do značné míry to platí i dnes. Ale plná decentralizovanost by znamenala zachování běžného chodu systému (třeba i s určitým zpomalením nebo neposkytováním některých služeb, které mohou „počkat“, nejsou časově kritické) i po odpojení naprosté většiny serverových uzlů, což už dnes Internet nesplňuje.

 **Distribuovaný systém** je takový systém, jehož funkce jsou distribuovány (rozděleny) na různé uzly sítě se zachováním několika důležitých vlastností.

Především jde o vlastnost *transparentnost*. Například přístupová transparentnost znamená, že k lokálním i vzdáleným prostředkům se přistupuje stejným způsobem, zde přes protokoly. Migrační transparentnost znamená možnost přesouvání poskytovaných služeb mezi různými uzly sítě bez výraznějšího vlivu na výkonnost sítě.

Další důležitou vlastností distribuovaného systému je *flexibilita*. Tato vlastnost znamená přizpůsobivost systému změnám prostředí (včetně poruch a výpadků částí systému). To vylučuje jakoukoliv centralizaci rozhodování, každý uzel sítě musí být ve své činnosti co nejvíc samostatný a služby ostatních uzlů vyžaduje se zachováním transparentnosti.

Flexibilní systém může být jakkoliv rozšiřován (téměř – jsou technické hranice, které se jen velmi nesnadno překračují) nebo naopak omezován, funkce odstraňovaného uzlu může plnit jiný uzel.

S distribuovaností se setkáváme u některých směrovacích protokolů a dále například v systému DNS (obecně u doménových služeb).



Poznámka:

V této kapitole se budeme zabývat různými technologiemi využívajícími zejména decentralizovaný a distribuovaný přístup.



7.2 Bridging, switching, routing

Nejdřív se krátce podíváme na pár pojmů, které souvisejí s distribuovaným přenosem a zpracováním dat v počítačových sítích.

 **Bridging** zahrnuje technologie související s vrstvou L2, podle standardu tedy mosty (bridge). Zahrnuje především to, co najdeme v standardu IEEE 802.1, včetně STP.

Mosty (obecně jakákoliv zařízení na L2) mohou pracovat v jednom z těchto režimů – co se týče práce s informacemi o cestě rámce:

- *Source-route bridging* – odesílající uzel napevno určí seznam mostů (obecně mezilehlých zařízení L2), přes které má rámec projít (buď se použije pro vytvoření obdoby okruhu, nebo se uloží do záhlaví rámce).
- *Transparent bridging* – odesílatel pouze určí fyzickou adresu cíle, mosty si vedou a dynamicky udržují a dolňují přepínací tabulku, pomocí které stanovují svůj úsek cesty vedoucí k danému cíli.

Je zřejmé, že v Ethernetu se používá transparent bridging. Source-route bridging se používal například u sítě Token-Ring, s jeho obdobou se setekáváme ve WAN sítích (tam jsme obvykle taky na vrstvě L2). Při source-route bridging je tedy nutné dodávat informace pro správné určování adres a portů, kdežto při transparent bridging musí v síti existovat mechanismus transportu fyzických adres na mosty.

Jak víme, u Ethernetu je tento problém řešen velice jednoduše – pokud je na portu detekován rámec, jehož zdrojovou adresu neznáme (což pro samotné přepnutí rámce není nutné), poznamenejme si do tabulky tuto zdrojovou adresu a označení portu, ze kterého rámec přišel, protože někde za tímto portem toto zařízení velice pravděpodobně bude připojeno.

 **Switching** se nevztahuje k některému standardu, ale spíše ke konkrétním technologiím, které najdeme na přepínačích, vlastně je to jakási nástavba bridgingu. Rozlišujeme tyto typy přepínání:

- *Store-and-Forward* (ulož a pošli) – přichází rámec je nejdřív celý přijat a uložen do vyrovnávací paměti přepínače, až po přijetí celého rámce je přečteno záhlaví a určen výstupní port pro odeslání rámce. Dokáže odchyťovat chybné rámce, ale pomalý, vhodný pro síť s vyšší chybovostí nebo switche s výkonnějším hardwarem.
- *Cut-Through* (průběžné zpracování, také on-the-fly) – přichází rámec se začíná odesílat průběžně ještě během svého příjmu, hned po načtení informací ze záhlaví potřebných k určení výstupního portu. Velmi rychlé přepínání, menší potřeba cache, ale nedokáže odchyťovat chybné rámce.
- *Fragment Free* (s omezením malých rámců) – něco mezi prvními dvěma metodami. Zpracovávání rámce začne po načtení jeho prvních 64 oktetů (tj. záhlaví a alespoň část dat). Účelem je odhalení alespoň jednoho typu chybných rámců: nepovolených příliš krátkých rámců.

Co je lepší? To je otázka. Záleží na vytíženosti sítě a výkonu switchů. Na zařízeních určených pro korporátní sféru je obvykle možné zvolený mód nastavit. Některé funkce v síti vyžadují určitý konkrétní mód, například pokud potřebujeme kvalitu služby (QoS), mělo by být na switchích nastaveno store-and-forward.

Příklad

Například na některých switchích Cisco je jako výchozí nastaven mód cut-through. Pokud chceme nastavit store-and-forward, zadáme:

```
switch# switching-mode store-forward
```

Pokud bychom chtěli nastavit zpět cut-through, předřadíme před tento příkaz no:

```
switch# no switching-mode store-forward
```

Jiné switche od Cisca mají napevno nastaven mód store-and-forward a nelze to změnit. Pokud v příslušném módu zadáme otazník, získáme seznam podporovaných příkazů. Pokud tam je příkaz `switching-mode`, napíšeme ho a místo parametru dáme otazník. Tak zjistíme, který mód lze nastavit.



Obecně platí, že mosty a přepínače odesílají unicast provoz (se známou cílovou adresou) pouze na jeden port, kdežto broadcast a unicast s neznámou cílovou adresou odesílají na všechny porty kromě toho, ze kterého rámec přišel.

 **Routing** se týká činnosti routerů a dalších zařízení pracujících na vrstvě L3, včetně switchů s funkcionalitou vyšších vrstev. Víme už, jak vypadají směrovací tabulky; v rámci routingu je třeba zajistit distribuci aktuálního obsahu těchto tabulek a tyto tabulky používat při směrování paketů.

Účelem je, aby směrovací tabulky byly vždy aktuální, a proces uvádění tabulek do aktuálního (konzistentního) stavu se nazývá *konvergence sítě*. Pokud se v síti používá dynamické směrování (tj. výměna

směrovacích informací pomocí protokolů), je proces konvergence sítě jedním z nejdůležitějších parametrů daného protokolu – konvergence by měla být rychlá a zároveň by neměla moc vytěžovat síť.

Routingem se budeme zabývat podrobněji hned v další sekci.

7.3 Směrování

7.3.1 Jak směrujeme

Směrování (routing) probíhá na vrstvě L3 a vlastně znamená proces určení cesty do konkrétních IP sítí. Tato informace je dále využívána k určení cesty pro konkrétní IP paket.



Definice (Směrovací tabulka)

Na routeru a všech dalších zařízeních vrstvy L3 je vedena minimálně jedna *směrovací tabulka* (routing table), ve které máme pro každý záznam především tyto informace:

- adresa (pod)sítě a maska nebo délka prefixu, může být agregovaná,
- určení síťového rozhraní pro odeslání a/nebo souseda, přes kterého vede cesta k cíli,
- metrika (kvalita určené cesty), případně další (například zdroj směrovací informace).

Součástí směrovací tabulky také bývá informace o bráně (na routeru „gateway of the last resort“), tedy určení směru, na který se má posílat vše, pro co nemáme určenou cestu jinde ve směrovací tabulce.



Tuto tabulku je třeba nějak naplnit a dále aktualizovat podle momentální situace (některá cesta přestane být platná, nebo naopak jiná cesta je zpřístupněna, mění se metriky či přidáváme novou podsíť). To se buď provádí ručně (staticky) nebo to můžeme nechat na *směrovacích protokolech* (dynamická údržba).



Z toho vyplývá, že směrování může být

- *statické* (static routing) – správce sítě ručně vkládá záznamy do směrovacích tabulek,
- *dynamické* (dynamic routing) – na směrovačích v síti běží směrovací protokoly, které si udržují přehled o topologii sítě (jak je co propojeno) a aktualizují směrovací tabulky.

V praxi se obojí kombinuje, tedy některé (blízké či nějakým způsobem citlivé) cesty vkládáme ručně a zbylé necháme doplňovat a aktualizovat dynamicky.

Údaje ze statického směrování jsou považovány za důvěryhodnější, proto jsou obvykle upřednostňovány před údaji získanými dynamicky pomocí protokolů.



Metrika (kvalita cesty) je číslo, které bylo vypočteno podle určitých kritérií. Každý směrovací protokol používá trochu jiná kritéria a taky trochu jiným způsobem z nich metriku počítá. Jako kritéria se například mohou použít:

- počet routerů na cestě,
- propustnost cesty (například zda jde o Fast Ethernet nebo Gigabit Ethernet),
- latence (zpoždění v ms),
- spolehlivost cesty (nízká chybovost, málo zahozených paketů apod.),
- vytíženost cesty,
- maximální hodnota MTU po cestě (tj. maximální velikost paketu, který lze poslat),
- důvěryhodnost zdroje, ze kterého byla získána informace o cestě, atd.

Obvykle platí: čím menší číslo metriky, tím lepší cesta. U přímo připojených sítí a staticky určených záznamů metriku nepotřebujeme, automaticky jsou považovány za nejdůvěryhodnější a nejlepší.

 Kdy a jak se pracuje se směrovací tabulkou: pokud router přijme paket s cílovou IP adresou, potřebuje zjistit, na který port má tento paket přeposlat. Prochází směrovací tabulku a zjišťuje, kterému záznamu v tabulce tato adresa odpovídá, tedy zjišťuje, zda pro daný záznam adresa patří do podsítě v záznamu uvedené.

 Pokud v tabulce existují dvě cesty k témuž cíli (tj. ve dvou záznamech je adresa podsítě, do které se dá zařadit adresa z paketu), obvykle se upřednostní ta, která má delší prefix (tj. je přesnější, „most specific“). Například pokud máme údaje

10.160.0.0/12

10.160.0.0/17

a adresa konkrétního zařízení, pro kterou hledáme údaj v tabulce, odpovídá oběma, druhý údaj bude upřednostněn. Jde o dva různé záznamy, protože se liší v délce prefixu, třebaže adresa v obou případech vypadá jako stejná. Druhá adresa zřejmě bude podsítí sítě určené první adresou (nebo hlouběji v hierarchii).

Pokud ve směrovací tabulce není nalezena žádná cesta, je buď paket předán *bráně* (gateway), protože paket zřejmě patří do jiné sítě, nebo zahozen (když žádná brána není dostupná).



Příklad

Směrovací tabulka na routeru může vypadat například takto (pro IPv4):

```
R2#show ip route
Codes: L - local, C - connected, S - static, R - RIP, M - mobile, B - BGP
       D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
       N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
       E1 - OSPF external type 1, E2 - OSPF external type 2, E - EGP
       i - IS-IS, L1 - IS-IS level-1, L2 - IS-IS level-2, ia - IS-IS inter area
       * - candidate default, U - per-user static route, o - ODR
       P - periodic downloaded static route
```

Gateway of last resort is 0.0.0.0 to network 0.0.0.0

```
O    10.0.0.0/8 [110/2] via 192.168.4.2, 00:00:19, GigabitEthernet0/0/0
    100.0.0.0/8 is variably subnetted, 2 subnets, 2 masks
C    100.64.0.0/24 is directly connected, Serial0/2/0
L    100.64.0.55/32 is directly connected, Serial0/2/0
O    172.31.0.0/16 [110/2] via 192.168.8.2, 00:00:19, GigabitEthernet0/0/1
    192.168.4.0/24 is variably subnetted, 2 subnets, 2 masks
C    192.168.4.0/24 is directly connected, GigabitEthernet0/0/0
L    192.168.4.1/32 is directly connected, GigabitEthernet0/0/0
    192.168.8.0/24 is variably subnetted, 2 subnets, 2 masks
C    192.168.8.0/24 is directly connected, GigabitEthernet0/0/1
L    192.168.8.1/32 is directly connected, GigabitEthernet0/0/1
S*   0.0.0.0/0 is directly connected, Serial0/2/0
```

Přímo připojené sítě jsou na začátku řádku označeny písmenem „C“, sítě zjištěné směrovacím protokolem OSPF písmenem „O“. U všech dynamicky zjištěných sítí máme jak IP adresu souseda, přes kterého vede do daného cíle cesta, tak i příslušné síťové rozhraní vedoucí k sousedovi.

U dynamicky zjištěných položek je hned za adresou dvojice čísel v hranatých závorkách. Druhé číslo (zde vždy 2) je právě metrika. K prvnímu číslu se vrátíme později.

Směrovací tabulka na tomtéž zařízení pro IPv6:

```
R2#show ipv6 route
IPv6 Routing Table - 9 entries
Codes: C - Connected, L - Local, S - Static, R - RIP, B - BGP
       U - Per-user Static route, M - MIPv6
       I1 - ISIS L1, I2 - ISIS L2, IA - ISIS interarea, IS - ISIS summary
       ND - ND Default, NDp - ND Prefix, DCE - Destination, NDr - Redirect
       O - OSPF intra, OI - OSPF inter, OE1 - OSPF ext 1, OE2 - OSPF ext 2
       ON1 - OSPF NSSA ext 1, ON2 - OSPF NSSA ext 2
       D - EIGRP, EX - EIGRP external
C  2001:DB8:CAFE:4::/64 [0/0]
    via GigabitEthernet0/0/0, directly connected
L  2001:DB8:CAFE:4::1/128 [0/0]
    via GigabitEthernet0/0/0, receive
C  2001:DB8:CAFE:8::/64 [0/0]
    via GigabitEthernet0/0/1, directly connected
L  2001:DB8:CAFE:8::1/128 [0/0]
    via GigabitEthernet0/0/1, receive
O  2001:DB8:CAFE:10::/64 [110/2]
    via FE80::1, GigabitEthernet0/0/0
O  2001:DB8:CAFE:172::/64 [110/2]
    via FE80::209:7CFF:FEA0:E601, GigabitEthernet0/0/1
C  2001:DB8:FEED:1::/64 [0/0]
    via Serial0/2/0, directly connected
L  2001:DB8:FEED:1::55/128 [0/0]
    via Serial0/2/0, receive
L  FF00::/8 [0/0]
    via Null0, receive
```



 Rozlišujeme L2 protokoly *směrovatelné* (routed) a *směrovací* (routing). Směrovatelné protokoly (například IP) pracují na síťové vrstvě, ale směrování neovlivňují, to je záležitostí směrovacích protokolů (pracujících na L3 nebo jiné vrstvě).

 *Jednoprotokolový* router dokáže pracovat jen s jedním směrovacím protokolem, *multiprotokolový* router dokáže ve směrovací tabulce kombinovat údaje od více směrovacích protokolů.

V multiprotokolovém směrovači je třeba navíc rozhodnout mezi cestami nalezenými různými protokoly, jejichž metriky jsou vytvářeny podle různých kritérií a tedy tak jak jsou, nebývají navzájem porovnatelné. Z toho důvodu byla pro každý protokol stanovena hodnotící konstanta nazvaná *administrativní vzdálenost* (administrative distance, AD) umožňující porovnat více cest k těmž cílům od různých protokolů. Toto číslo reálně znamená důvěryhodnost zdroje (jak moc důvěřujeme, že tento zdroj dodává „dobré“ cesty). Ve výpisu v příkladu na straně 160 je AD právě první číslo v hranaté závorce za adresou cíle.

Administrativní vzdálenosti jsou v tabulce 7.1. Všimněte si rozdílů u protokolů, které mohou pracovat jako vnitřní i vnější (například BGP). Na routerech se ve skutečnosti dá AD konfigurovat.

Pokud do daného cíle dokáže směrovat více protokolů, bude vybrána cesta toho směrovacího protokolu, který má menší číslo administrativní vzdálenosti. Jak vidíme, připojené rozhraní je nejdůvěryhodnější, protože router si je do směrovací tabulky řadí sám (sám sobě důvěřuje nejvíce) a cesta tímto směrem je nejkratší. Následují staticky přidané záznamy. Vnitřní EIGRP má AD=90, OSPF má AD=110, takže

Typ cesty	Administrativní vzdálenost
Připojené rozhraní	0
Statické směrování	1
EIGRP summary route (souhrnné cesty pro skupinu sítí)	5
External BGP	20
Internal EIGRP	90
IGRP	100
OSPF	110
IS-IS	115
RIP	120
EGP	140
ODR (On Demand Routing)	160
External EIGRP	170
Internal BGP	200
jiný (neznámý)	255

Tabulka 7.1: Administrativní vzdálenosti směrovacích protokolů

u EIGRP se předpokládá mírně lepší schopnost určování cest (ale na druhou stranu, EIGRP většinou umí jen zařízení od Cisco, takže se dá použít jen v homogenní síti).

 Cesta zjištěná v rámci běžného směrování je *interní*, cesta zjištěná vně směrovacího procesu (typicky z jiného autonomního systému) je *externí*. Většina protokolů upřednostňuje interní cesty, ale například u BGP je to naopak a jiné jejich spolehlivost nerozlišují (například OSPF).

7.3.2 Autonomní systém

 *Autonomní systém* je skupina směrovačů patřících pod správu téže organizace. Každý autonomní systém je identifikován 16bitovým nebo 32bitovým číslem *ASN*.

- *Vnitřní směrovací protokoly* (interior routing protocols) dokážou směrovat jen uvnitř vlastního autonomního systému (cestu ven neznají, všechno „cizí“ posílají prostě na jedinou bránu), nerozlišují jiné autonomní systémy (rozlišují jen „vlastní“ a „cizí“). Ve svých datových strukturách nemají pole pro evidenci autonomního systému.
- *Vnější směrovací protokoly* (exterior routing protocols) se zabývají cestami vedoucími mezi různými autonomními systémy (dokážou pracovat s číslem autonomního systému, ve směrovací tabulce máme v každém záznamu i číslo ASN).



Poznámka:

Musí mít každá organizace svůj autonomní systém a své číslo ASN? Ne. Menším firmám stačí být součástí autonomního systému svého ISP, a jeho číslo ASN obvykle ani neznají. Pokud firma používá pouze vnitřní směrovací protokoly, není důvod mít vlastní ASN (hlavně za to platit).



 Číslo ASN může být dlouhé 16 bitů nebo 32 bitů, tedy 2 B nebo 4 B. K rozšíření datového prostoru na dvojnásobek došlo před lety z obavy, že rozsah 16 bitů v budoucnu nebude stačit. Pro zápis 32bitových čísel ASN existují různé konvence:

- ASPlain: 32bitové ASN se jednoduše se v desítkové soustavě zapíše 32bitové číslo, například 232576.
- ASDot: každá polovina 32bitového ASN se v desítkové soustavě zapíše zvlášť, oddělí se tečkou, například 3.35968 (to je totéž ASN jako předchozí). Pokud se ASN vejde do 16 bitů (a před tečkou by byla 0), zapíše se ve formátu ASPlain.
- ASDot+: podobně jako ASDot, ale i menší čísla ASN se zapisují s tečkou (a před tečkou je nula). Účelem tohoto formátu je jednotný zápis kratších i delších ASN (s tečkou) kvůli snadnějšímu vyhledávání pomocí regulárních výrazů.

BGP routery (které s ASN typicky pracují) obvykle „umí“ oba zápisy, ale jeden z nich je nastaven jako defaultní. Pokud chceme používat jiný zápis, bývá nutné to v konfiguraci změnit.

ASPlain	ASDot	ASDot+
10	10	0.10
64482	64482	0.64482
68521	1.2985	1.2985
859520	12.7552	12.7552

Tabulka 7.2: Příklady zápisu čísel ASN

 Správou čísel ASN je pověřena organizace IANA, stejně jako správou IP adres, čísel portů a dalších parametrů sítí, a způsob jejich využívání je stanoven v řadě dokumentů RFC. IANA určila využití určitých rozsahů čísel ASN tak, jak ukazuje tabulka 7.3.

Rozsah hodnot ASN	Význam
0	rezervováno
1–64.495	přidělováno organizací IANA, veřejná čísla ASN
64.496–64.511	rezervováno pro dokumentaci, testování
64.512–65.534	soukromý rozsah ASN
65.535	rezervováno

Tabulka 7.3: Rozsahy čísel ASN[26]

Z tabulky vyplývá, že existují veřejná a soukromá čísla ASN. Důvod je podobný jako u IPv4, soukromá čísla ASN existují kvůli tomu, aby se zbytečně neplýtvalo těmi veřejnými. U soukromého čísla je ovšem možné, že totéž používá více autonomních systémů, tedy existují určitá omezení, aby se předešlo chybným updatům a případně smyčkám v síti peerů. Soukromá ASN se doporučují například systémům, které nejsou průchozí (jsou připojeny jen k jednomu jinému autonomnímu systému, tzv. stub AS), třebaže použití pro systém propojený s několika jinými systémy se nevylučuje. Důvodem je, že v takovém případě určitě nehrozí vytvoření smyčky.

7.3.3 Směrovací algoritmy

 Každý směrovací protokol se řídí nějakým *směrovacím algoritmem*. Obvykle se jedná o jeden z následujících dvou algoritmů (případně modifikaci):

- algoritmus vektoru vzdáleností,
- algoritmus stavu spoje.

 **Algoritmus vektoru vzdáleností** (Distance-Vector Algorithm) určuje metriku cesty především podle vzdálenosti (v jednodušším případě podle počtu routerů na cestě). Pro každou síť evidovanou ve směrovací tabulce máme vektor, který určuje vzdálenost a směr k dané síti, tj. buď přes které síťové rozhraní nebo přes kterého souseda vede cesta. Pozice cíle je určena pouze vektorem, zařízení nemají přehled o topologii sítě routerů.

Router pracující podle tohoto algoritmu má nejdřív v tabulce cesty do připojených sítí, tj. cesty k přímým sousedům, a od nich postupně dostává informace o jiných routerech a jejich podsítích. Pokud od svého souseda dostane update s informací o (pod)síti s konkrétní metrikou (podle toho souseda), přidá si tuto (pod)síť do směrovací tabulky:

- adresu (pod)sítě zjistil od souseda,
- rozhraním bude adresa či port dotyčného souseda (protože přes něj je ta síť dostupná),
- jako metriku použije o něco vyšší číslo než jaké mu poslal soused, v jednodušším případě o 1 větší než jaká platí pro souseda (tj. k cestě od souseda k síti je třeba započítat i úsek od souseda ke mně).

Pro tento algoritmus je typické, že updaty (informace o aktuálním stavu směrování) se sousedům posílají v pravidelných intervalech.

Jak vidíme, může vzniknout podobný problém jako v síti switchů, kdy jsme museli pro odstranění smyček použít protokol STP. Smyčky mohou vzniknout i zde (máme síť routerů používajících daných algoritmus/protokol), přičemž důsledky by mohly být podobné. Představme si například situaci, kdy tatáž síť je dostupná přes dvě různé cesty. Router dostává informace o síti střídavě ze dvou směrů, což nesvědčí stabilitě směrovacích tabulek. Stabilita fungování algoritmu se vylepšuje různými způsoby, například:

 Při použití tohoto algoritmu může být definován *maximální počet přeskoků* (maximum hop count), který určuje maximální velikost sítě (v počtu routerů na cestě) – obvykle 15 nebo 255. Pokud router dostane informaci o síti s metrikou vyšší než tato hranice (například u RIP s hodnotou 16), pak takovou síť považuje za nedostupnou.

 Dále je určen *hold-down timer*, většinou nastavený na trojnásobek intervalu zasílání updatů sousedům. Pokud router obdrží informaci o nedostupnosti některé sítě, pak po dobu určenou hold-down timerem ignoruje zprávy o cestě do této sítě, považuje je za potenciálně chybné.

 *Poison reverse* (otrávení cesty zpět) je oznámení směrovače o své nedostupnosti. Tento směrovač odešle svým sousedům aktualizaci cesty k sobě s hodnotou metriky „maximum hop count“ a sousedi si k tomu ještě přičtou 1, tj. například u protokolu RIP to bude hodnota 16 nebo u jiných protokolů hodnota 256. Metrika větší než maximální stanovená označuje „nedosažitelnou“ cestu a tedy na port k odpojovanému směrovači nebudou zasílány žádné pakety.

 Další možnost snížení zatížení sítě a rizika zacyklení je postup *split horizon* (rozložený horizont). Směrovač neodesílá jinému směrovači celou svou tabulku, ale jen ty záznamy, které nevedou k tomuto směrovači (není třeba informovat jiný směrovač o tom, o čem původně informoval on). Takto se také snižuje nebezpečí zasílání chybných informací.

 Protokoly využívající algoritmus vektorů vzdáleností:

- RIP (Routing Information Protocol) varianty pro IP, IPX, atd.,
- IGRP (Interior Gateway Protocol) pro IP,
- EIGRP, nicméně tento algoritmus obohacuje o některé prvky typické spíše pro algoritmus stavu spoje, jeho metrika je považována za velmi dobrou.

 **Algoritmus stavu spoje** (Link-state Algorithm, algoritmus nejkratší cesty) vyžaduje budování *topologické databáze sítě* (mapy sítě). Router s protokolem podle tohoto algoritmu si vede tři tabulky:

- tabulku sousedů, kterou si vytvoří jako první hned po zapnutí,
- tabulku reprezentující topologickou databázi sítě, tu si po zapnutí vytvoří postupně podle aktualizací informací ostatních routerů,
- směrovací tabulku, kterou si vytvoří podle topologické databáze.

Každý router si neustále hlídá kvalitu (stav) spojů vedoucích k sousedům (to se týká první zmíněné tabulky). Pokud zjistí změnu (některý spoj/soused se stane nedostupným nebo změní některý parametr), informuje o tom své okolí. Routery, které dostanou tuto informaci, si změnu zanesou do topologické databáze (druhé tabulky) a podle ní propočítají novou cestu k cílům (aktualizují směrovací tabulku).

Pro protokoly podle algoritmu stavu spoje je typické, že mají přehled o topologii sítě, méně zahrnují síť (posílají se jen změny ke konkrétním spojům, a většinou jen při změnách topologie) a doba konvergence je výrazně kratší (aktualizace se jen přeposílají, změnu si propočítává každý router sám, nemusejí na sebe navzájem čekat).

 Samotné propočítávání tras a zjišťování nejkratší (neoptimálnější) cesty k dané (pod)síti se provádí podle topologické databáze *Dijkstrovým algoritmem* (algoritmem nejkratší cesty, shortest path first, SPF) nebo jeho modifikací. Jedná se o jeden z nejznámějších grafových algoritmů, jehož autorem je holandský informatik Edsger Dijkstra.

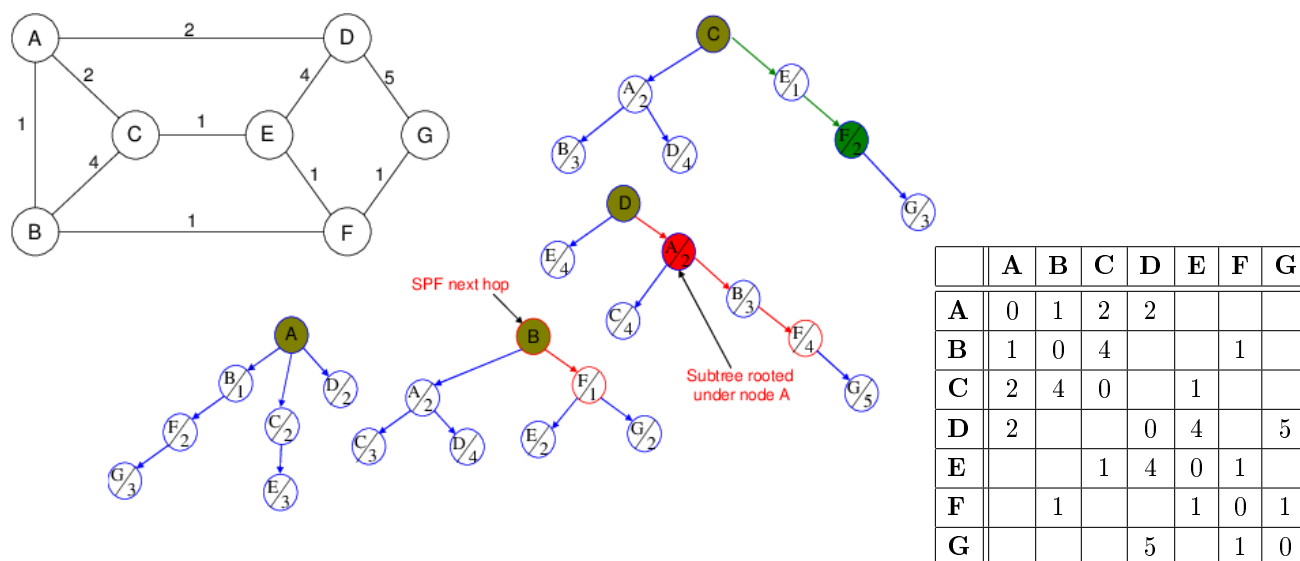
Příklad

Podíváme se, jak funguje Dijkstrův algoritmus. Každý spoj mezi routery je ohodnocen číslem, přičemž menší číslo znamená kvalitnější spoj. Každý router si utváří topologickou databázi obsahující všechny routery v dané doméně a ke každé dvojici routerů informaci, zda mezi nimi vede spoj, a pokud ano, jaké je jeho ohodnocení. Topologickou databázi si můžeme představit jako tabulku, kde řádky i sloupce jsou označeny jednotlivými routery, v buňce pro určitou dvojici routerů je informace o spoji mezi nimi (nebo o tom, že tam spoj není).

Na obrázku 7.1 na straně 166 je graf se sedmi routery (označenými A až G) s ohodnocením spojů mezi nimi. Na obrázku vpravo je tabulka s odpovídající topologickou databází.

Tato tabulka by měla být stejná na všech routerech, každý router má tedy představu o topologii sítě. Podle topologické tabulky se vypočítává cena cesty a tedy nejvhodnější cesta ke všem prvkům sítě, což je základ směrovací tabulky.

Postup stojí na přetvoření původního grafu (ve kterém máme redundantní cesty) do formy stromu – v kořeni stromu je ten router, na kterém probíhá tento výpočet (tj. každý router bude mít jiný strom, sebe dosadí na vrchol), z původních spojů některé odstraní a některé nechá. Zůstanou jen takové spoje, které jsou součástí optimálních cest.

Obrázek 7.1: Ukázka použití Dijkstrova algoritmu a odpovídající topologická databáze¹

Na obrázku 7.1 vidíme stromy pro routery A, B, C, D. Například router A nechal k routeru F pouze cestu přes souseda B, kdežto cesty přes další sousedy odstranil. K routeru G nechal cestu $A \rightarrow B \rightarrow F \rightarrow G$, jejíž cena je $1 + 1 + 1 = 3$, nepoužil například cestu $A \rightarrow D \rightarrow G$, která sice vede přes méně mezilehlých uzlů, ale její cena je $2 + 5 = 7$, tedy horší.

Algoritmus SPF je jeden z nejpoužívanějších grafových algoritmů. Nejznámějším protokolem stavu spoje je OSPF, tento algoritmus používá také protokol IS-IS.

Poznámka:

Typickou vlastností protokolů implementujících algoritmus stavu spoje je, že mají přehled o své síti nebo její podstatné části (autonomní systém může být rozdělen na relativně samostatné *oblasti*, které jsou navzájem propojeny hraničními routery, a pak stačí, aby měl protokol přehled o oblasti, ve které se nachází). Další typickou vlastností je rychlá konvergence a menší náchylnost k zacyklení.

7.3.4 Směrovací protokol RIP

RIP (Routing Information Protocol) je jedním z nejstarších směrovacích protokolů, byl vyvinut společností Xerox už roku 1981. Jedná se o vnitřní směrovací protokol komunikující na portu 520 (využívá protokol UDP). Implementuje algoritmus vektoru vzdáleností, jako metrika se používá počet směrovačů na cestě k cíli.

RIPv1 (verze 1, RFC 1058) využívá třídy adres (nepodporuje beztržní směrování), to znamená, že součástí směrovacích informací není maska podsítě ani délka prefixu. To je další nevýhoda, která RIPv1 prakticky vylučuje z větších sítí.

¹Zdroj: https://www.researchgate.net/figure/Example-topology-and-routing-trees-of-node-A-and-its-neighbours_fig1_225143830

Nejvyšší akceptovaná metrika je 15, číslo 16 již označuje nedostupnou síť, a tedy v mechanismu Poison Reverse se právě číslo 16 používá pro oznámení nedostupnosti. V intervalech 30 s je všesměrově (broadcast) vysílána směrovací informace (celá směrovací tabulka), což ve velkých sítích může značně snižovat propustnost sítě.



Příklad

Adresy podsítí

10.10.22.0/24

10.20.10.0/24

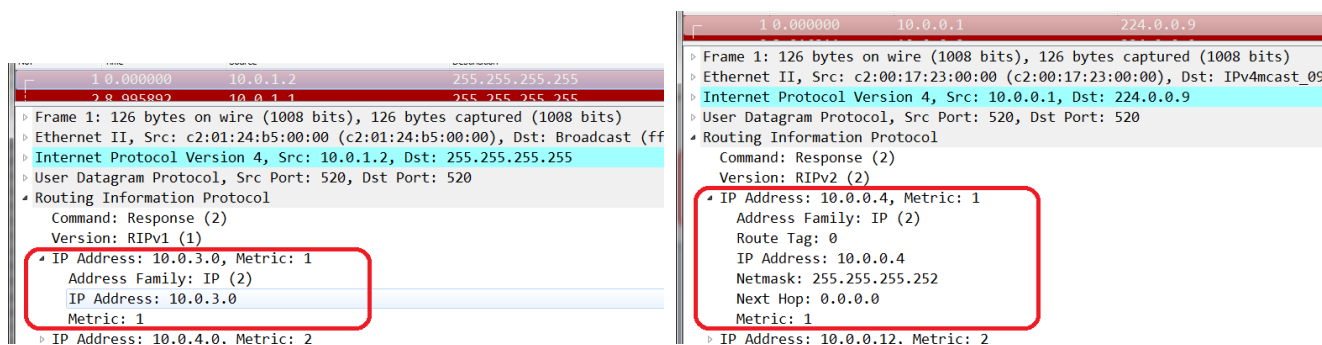
by byly považovány za stejnou síť, protože RIP směrovač by obě adresy považoval za adresu třídy A, kde je síť určena prvním oktetem. Pokud k těmto sítím vede cesta přes různé směrovače, nebude směrování probíhat správně (jeden údaj bude přepsán druhým, jedna ze sítí nebude dostupná).



RIPv2 (STD-56) je aktualizací protokolu RIP z poloviny 90. let 20. století. Hlavním úkolem bylo zohlednění využívání prefixů a beztrždního směrování (VLSM a CIDR). Rozdíly oproti RIPv1:

- lze používat beztrždní směrování (v aktualizacích paketech posílá záznamy o sítích včetně masky),
- aktualizace směrovacích informací se neposílají jako broadcast zprávy, ale jako multicast zprávy na adresu 224.0.0.9, v případě známých sousedů může být použit unicast,
- podpora ověřování mezi dvěma směrovači, ale na bídné úrovni (buď nešifrovaně nebo šifrovaně s MD5).

Srovnání posílaných informací vidíme na obrázku 7.2 (vlevo RIPv1 posílající jen adresu a metriku, vpravo RIPv2 posílající i masku a další informace).



Obrázek 7.2: Srovnání informace o síti v updatu RIPv1 a RIPv2

U obou verzí RIP je metrika považována za nepřiliš kvalitní, tedy použitelnost je pouze v malých sítích, kde to moc nevadí. Omezení na 15 skoků přetrvává.



RIPv2 může nebo nemusí používat *automatickou sumarizaci cest*. Sumarizace zde funguje tak, že pokud jsou ve směrovací tabulce dvě (pod)sítě se stejnou „trždní“ částí adresy, sloučí se do jednoho řádku směrovací tabulky. Slučování se provádí na hranici podle třídy adresy.



Příklad

Podsítě 10.1.0.0/16 a 10.2.0.0/16 mají adresy třídy A, tj. zajímá nás první byte adresy – v obou je to 10. Tyto dvě podsítě by byly při zapnuté autosumarizaci sloučeny do jednoho řádku směrovací tabulky.

Pokud do obou vede cesta přes totéž síťové rozhraní, je to v pořádku, v opačném případě máme problém. Takže rozhodnutí zapnout/vypnout autosumarizaci je velmi důležité.



 *RIPng* (v podstatě třetí verze, next gen, RFC 2080) přidává podporu IPv6, tedy v updatech posílá IPv6 adresy. Další vlastnosti má podobné jako RIPv2.

7.3.5 Směrovací protokoly IGRP a EIGRP

 **IGRP** (Interior Gateway Routing Protocol) je letitý proprietární protokol společnosti Cisco. Implementuje algoritmus vektoru vzdáleností, ale má mnohem lepší vlastnosti než RIP. Součástí konfigurace protokolu je číslo autonomního systému. Bohužel nepodporuje beztrždní směrování (počítá pouze s třídami, žádná maska ani délka prefixu se nepoužívá).

Směrovací informace se posílají každých 90 s (ve výchozím nastavení), a to broadcastem. Metrika je vícekritériální, využívá se kombinace kritérií

- zpoždění v síti (prodleva) za předpokladu nezatížené sítě v ms,
- propustnost v b/s, tj. šířka pásma,
- zatížení sítě, hodnota 255 znamená 100% zatížení,
- spolehlivost cesty, hodnota 255 znamená 100% spolehlivost.

První dvě kritéria jsou povinná (dají se odvodit z topologie sítě), další dvě kritéria jsou nepovinná (lze je odvodit z momentálního provozu). Další parametry trasy – počet směrovačů na cestě a MTU – se projevují při kompletaci výsledné metriky.

Zvolená kritéria se zjišťují pro každý spoj na cestě zvlášť a pak se do výsledku sečtou (například u zpoždění v síti) nebo se vybere nejhorší hodnota přes všechny spoje (například šířka pásma, to je vlastně úzké hrdlo celé cesty).

 Při výpočtu metriky se postupuje podle zvoleného *typu služby* (obdoba QoS), hodnoty jednotlivých kritérií jsou násobeny konstantami specifickými pro daný typ služby (například může být upřednostněna propustnost sítě).

 **EIGRP** (Enhanced IGRP) je vylepšení protokolu IGRP. Algoritmus směrování podle tohoto protokolu již není čistě algoritmem vektoru vzdáleností, objevují se některé prvky algoritmu stavu spoje, ale přesto je tento protokol řazen do skupiny využívající algoritmus vektoru vzdáleností. Metrika je (stejně jako u IGRP) vypočítávána z propustnosti a zpoždění v síti, je možné přidat kritéria zatížení sítě a spolehlivosti, ale obvykle se to nedělá.

Zatímco IGRP je classful, EIGRP je classless, tedy používá beztrždní směrování. Součástí informace o síti je i maska sítě. EIGRP v novější verzi podporuje také IPv6.

Podobně jako algoritmy stavu spoje, také EIGRP si udržuje tři tabulky: tabulku sousedů, topologickou tabulku a směrovací tabulku. Je zajímavé, že tyto tabulky jsou spravovány na aplikační vrstvě, tedy EIGRP je vlastně aplikační protokol, třebaže ovlivňuje fungování směrovací vrstvy. Tabulky jsou aktualizovány při změnách v topologii, a to multicast zprávami. Používá se multicast adresa 224.0.0.10, resp. pro IPv6 to je FF02::A.

 RIP používá jednoduchý algoritmus počtu routerů na cestě, OSPF si vytváří směrovací tabulku pomocí Dijkstrova algoritmu, EIGRP používá algoritmus *DUAL* (Diffusing Update Algorithm). DUAL je poměrně

složitý algoritmus, jehož cílem je nejen vypočítávat co nejoptimálnější cesty k cílům bez smyček, ale také udržuje údaje o některých dalších méně optimálních cestách. Na rozdíl od jiných algoritmů dokáže vyvažovat zátěž na více cest k témuž cíli (takže umí dodat do směrovací tabulky pro jeden cíl i více cest), a to i s různými metrikami (pak je posíláno více paketů optimálnější cestou a méně paketů cestou s horší metrikou). Navíc se uchovávají předvypočtené záložní cesty (backup routes), aby při změně topologie nebylo nutné znovu spouštět DUAL (je dost náročný na procesor).

Existuje více druhů zpráv, které se při používání EIGRP posílají (Hello, Update, Query, Reply, Ack), některé potvrzované, jiné nepotvrzované.

 Na transportní vrstvě používá Cisco svůj vlastní speciální protokol Cisco RTP (Reliable Transport Protocol – neplést si se „standardním“ RTP Real-time Transport Protocol). Cisco RTP je „něco mezi“ TCP a UDP v tom smyslu, že dokáže fungovat potvrzovaně i nepotvrzovaně, podle toho, co aplikačního přenáší. Například Hello pakety přenáší nepotvrzovaně, kdežto Update pakety přenáší potvrzovaně. Protože to je proprietární protokol, v provozu to vypadá, že Cisco RTP a EIGRP mají svá záhlaví smíchaná dohromady a jsou zapouzdřena v IP paketu.

 EIGRP používá především dva timery: *hello timer* určuje, jak často si mají sousedi navzájem posílat hello pakety, *hold timer* určuje, jak dlouho má router čekat, na hello paket od souseda (po uplynutí hold bez odezvy od souseda je soused prohlášen za nedostupného).

Na jednom směrovači může běžet více instancí protokolu (E)IGRP i další směrovací protokoly, protože součástí konfigurace je číslo směrovacího procesu (ve formě ASN).

Protokol EIGRP byl původně proprietární (společnost Cisco ho používala pouze na svých zařízeních), roku 2013 byl však uvolněn a dnes se jedná o volně použitelný standard RFC 7868, který může být implementován i na zařízeních jiných výrobců (nicméně nikdo se do toho moc nehrne, implementace EIGRP je přes nesporné výhody poměrně složitá).



Poznámka:

S protokolem EIGRP se i nadále setkáváme v praxi. Je považován za optimálnější algoritmus než OSPF, ale na druhou stranu není většinou výrobců podporován a je výpočetně náročný.



7.3.6 Směrovací protokol OSPF

 OSPF (Open Shortest Path First) používá algoritmus stavu spoje, což znamená rychlou konvergenci sítě při každé změně topologie. Jde o otevřený protokol, a proto se s ním setkáme na zařízeních od mnoha různých výrobců. Sice se typicky používá jako vnitřní směrovací protokol, ale je možné ho použít i jako vnější (mezi autonomními systémy). Jde o beztržní směrování, proto k informaci o adrese zadáváme i délku prefixu. OSPF verze 2 (RFC 1247 a RFC 2328) podporuje IPv4 adresy, kdežto OSPF verze 3 (RFC 2740) pracuje s IPv6 adresami.

Směrovací informace se zasílají okamžitě po změně topologie nebo minimálně každých 30 minut (to je pořád výrazně méně často než u předchozích protokolů). Ovšem neposílá se směrovací tabulka, posílá se pouze topologická informace.

Podobně jako EIGRP, i OSPF posílá několik druhů paketů (Hello, Database Description se stručným popisem topologie, Link-state Request s požadavkem na podrobnější info o spoji, Link-state Update s podrobnější informací o spoji, Link-state Ack s potvrzením), stejně jako u EIGRP tyto pakety zapouzdřuje do IP paketů (tj. pracuje na vrstvě L3, podobně jako třeba ICMP).

Každý směrovač si udržuje topologickou databázi sítě nebo své oblasti v síti. Topologická databáze má formu orientovaného grafu s tím, že jedna cesta může mít v každém směru přiřazenu jinou metriku. Pro výpočet cest z topologické databáze se používá Dijkstrův algoritmus (SPF), který byl vysvětlen výše.

 Metrika je označována pojmem *cena*. Cena je kombinací více kritérií – propustností sítě, náklady na spoje, atd., podle konfigurace, většinou se volí pouze výpočet z kritéria propustnosti sítě. Pro výpočet ceny spoje se používá referenční číslo. Jako výchozí může být na starších switchích nastavena hodnota 100 000 000, což je 100 Mb/s. Referenční číslo se vydělí propustností spoje, takže například právě 100 Mb/s má cenu 1, rozhraní s propustností 1.2 Mb/s má cenu 84, apod. Pokud vyjde číslo menší než 1, výsledek bude 1.

To je dost nepříjemné omezení, protože routery obvykle propojujeme spíše rychlejšími cestami, není výjimkou použití 10Gb spojů (jenže všechny spoje od rychlosti 100 Mb/s včetně výše by měly cenu 1). Naštěstí se koeficient pro přepočtení dá změnit.

Příklad

Pokud je rozhraní rychlejší než 100 Mb/s, standardně má také cenu 1 (zaokrouhlujeme nahoru). Této *nesrovnalosti* se předchází změnou vzorce v konfiguraci (změna referenční šířky pásma). Například posloupnost příkazů

```
R1(config)# router ospf 1
```

```
R1(config-router)# auto-cost reference-bandwidth 1000
```

Číslo v prvním příkazu je číslo procesu (na routeru teoreticky může běžet několik instancí OSPF, v realu se to nedoporučuje). Takže v OSPF s číslem procesu 1 měníme referenční šířku pásma na 1000 Mb/s, tedy ve vzorci bychom použili číslo 1 000 000 000 (to je třeba provést ve všech routerech v doméně, abychom si nevyrobili nestabilní síť s nepředvídatelným chováním). Tímto odlišíme fast ethernetový spoj od gigabitového (ale 10Gb spoj bude mít stejnou cenu jako gigabitový). Pokud budeme chtít odlišit i 10Gb spoje, ještě jednu nulu přidáme:

```
R1(config-router)# auto-cost reference-bandwidth 10000
```

Po změně referenčního čísla na 10 Gb/s (což je 10 000 Mb/s, 10 000 000 000 b/s) se tedy cena spoje vypočte takto:

Propustnost spoje	Vzorec	Cena spoje
100 Mb/s	$\frac{10\,000\,000\,000}{100\,000\,000} = 100$	100
1 Gb/s	$\frac{10\,000\,000\,000}{1\,000\,000\,000} = 10$	10
2.5 Gb/s	$\frac{10\,000\,000\,000}{2\,500\,000\,000} = 4$	4
5 Gb/s	$\frac{10\,000\,000\,000}{5\,000\,000\,000} = 2$	2
10 Gb/s	$\frac{10\,000\,000\,000}{10\,000\,000\,000} = 1$	1
40 Gb/s	$\frac{10\,000\,000\,000}{40\,000\,000\,000} = 0.25$	1

Aktuální propustnost spoje za konkrétním portem zjistíme takto:

```
R1# sh int g0/0/1 | include BW
```

Cenu pro spoj za tímto portem najdeme v tomto výpisu:

```
R1# sh ip ospf interface g0/0/1
```

```
GigabitEthernet0/0/1 is up, line protocol is up
Internet address is 192.168.8.1/24, Area 0
Process ID 1, Router ID 2.2.2.2, Network Type BROADCAST, Cost: 1
Transmit Delay is 1 sec, State BDR, Priority 1
Designated Router (ID) 3.3.3.3, Interface address 192.168.8.2
Backup Designated Router (ID) 2.2.2.2, Interface address 192.168.8.1
Timer intervals configured, Hello 10, Dead 40, Wait 40, Retransmit 5
Hello due in 00:00:06
Index 2/2, flood queue length 0
Next 0x0(0)/0x0(0)
Last flood scan length is 1, maximum is 1
Last flood scan time is 0 msec, maximum is 0 msec
Neighbor Count is 1, Adjacent neighbor count is 1
Adjacent with neighbor 3.3.3.3 (Designated Router)
Suppress hello for 0 neighbor(s)
```



Cena celé cesty je součtem cen pro jednotlivé úseky cesty. Takže když SPF najde různé cesty k těmž cíli, sežte ceny spojů na každé cestě, porovná výsledky

 OSPF podporuje *hierarchické směrování*. To znamená, že směrovače (a jejich sítě) lze v rámci směrovací domény rozdělit do *oblastí* (area). Rozlišujeme *vnitřní routery* (uvnitř některé oblasti), *hraniční routery oblasti* (ABR, Area Border Router) a *hraniční routery autonomního systému* (zprostředkovávají komunikaci mezi různými autonomními systémy, včetně těch s jinými směrovacími protokoly). Vnitřní router (uvnitř oblasti) si udržuje pouze topologickou bázi své oblasti, aktualizace při změnách topologie se omezuje jen na danou oblast.

Páteří síť mezi routery je označována jako *oblast 0*, routery v páteří síti jsou *páteří routery*. Páteří síť může být i WAN. V rozsáhlejších sítích by vždy měla být páteří síť a všechny „nenulové“ oblasti by měly sousedit pouze s oblastí 0.

Součástí směrovací tabulky jsou samozřejmě také prefixy adres sítí. Podobně jako mnohé další protokoly, také OSPF používá *sumarizaci směrovacích informací* – CIDR. Sumarizace je zapnutá zejména na hraničních směrovačích, ať se zbytečně mezi oblastmi neposílá příliš mnoho updatů topologie.

 V sítích typu broadcast (tj. takových, kde mohou být i více než dva routery, například typu Ethernet, a žádné okruhy) musí existovat jeden *pověřený směrovač* (Designated Router), který řídí veškerou aktualizaci směrovacích informací v síti. Pro případ, že by pověřený směrovač selhal, by měl existovat *záložní pověřený směrovač*. Pověřený směrovač provádí výpočet cest a tím snímá část zátěže z ostatních směrovačů v oblasti. Výběr pověřeného směrovače sice lze nechat na protokolu OSPF, ale lepší je provést to ručně, protože automatický výběr nedopadne vždy zrovna ideálně.

 OSPF podporuje také autentizaci. Ve verzi 2 umí ověřovat pomocí MD5 nebo SHA kontrolních součtů, ve verzi 3 byl autentizační mechanismus přepracován – používá se autentizace a šifrování pomocí IPSec.

Z dalších důležitých vlastností protokolu OSPF:

- podpora *rozložení provozu* k cíli mezi cesty se stejnou cenou (ale všechny pakety se stejnou cílovou IP adresou jdou stejnou cestou),
- podpora směrování podle typu služeb (ToS) IP.

V komerční sféře je OSPF momentálně nejrozšířenějším směrovacím protokolem.

7.3.7 Směrovací protokol IS-IS

 IS-IS (Intermediate System to Intermediate System) byl původně součástí architektury ISO/OSI, ale po úpravě byl zařazen do TCP/IP. Je určen především pro směrování bez navazování spojení. Bývá oblíbený zejména u velkých ISP.

IS-IS používá *algoritmus stavu spoje* a princip jeho činnosti je hodně podobný protokolu OSPF, taktéž vypočítává cesty pomocí Dijkstrova algoritmu (SPF). Rozlišují se směrovače uvnitř běžné oblasti (první úroveň – L1), páteřní směrovače (úroveň L2) a hraniční směrovače (L1/2), pozor, neplést si s L1, L2, atd. u vrstev modelů. Páteřní oblast 0 není nutné vytvořit, záleží, jestli to pro síť routerů má smysl. Sousediti mohou různé oblasti.

Zatímco OSPF posílá směrovací informace v IP paketech, IS-IS pracuje částečně na vrstvě L3 a částečně na vrstvě L2, využívá rámce linkové vrstvy. Ve skutečnosti vůbec nepočítá se síťovým zásobníkem TCP/IP, má vlastní formát L3 PDU a tyto své jednotky zapouzdřuje například do ethernetových rámců.

7.3.8 Směrovací protokol BGP

Pojem EGP má ve skutečnosti dva významy – označují se tak obecně vnější (externí) směrovací protokoly, ale také jeden konkrétní externí protokol (ten nejstarší).

 *EGP* (Exterior Gateway Protocol) byl jednoduchý vnější směrovací protokol pro výměnu informací o dostupnosti či nedostupnosti sítí (předchozí popsané protokoly byly vnitřní). Svého času byl jediným externím směrovacím protokolem a s dalšími se nepočítalo, proto tento protokol dostal název podle své kategorie protokolů.

EGP nepoužíval žádnou metriku, a proto byla vyloučena existence více paralelních cest k téže síti; není použitelný, pokud se na cestách vyskytuje smyčka (vyžaduje stromovou strukturu směrovačů). V současnosti se už s EGP nesetkáváme.

 *BGP* (Border Gateway Protocol) je také vnější (exterior) směrovací protokol, ale již složitější, funkčnější a použitelnější. Propojuje různé autonomní systémy, slouží tedy ke směrování mezi autonomními systémy, nikoliv uvnitř (jakéhokoliv) autonomního systému. Ve skutečnosti se dá použít pro interní (iBGP, v rámci jednoho autonomního systému) i externí (eBGP, mezi různými autonomními systémy) směrování. Podporuje beztržní směrování, VLSM, CIDR s agregací cest a v novější verzi samozřejmě IPv6.

Protokol BGP je hlavním směrovacím protokolem na Internetu a WAN sítích (existují i jeho upravené varianty pro konkrétní typy sítí), jeho využití je poměrně časté.

 **Kdy se BGP používá.** Protokol BGP má smysl používat tehdy, pokud ve svých směrovacích tabulkách potřebujeme mít přímo zaneseny cíle nacházející se na Internetu a nestačí pouze „default route“, takže tento protokol najdeme zejména u ISP: potřebují mít zachyceny cesty k různým páteřním sítím (jsou napojeni na různé přístupové body PoP a případně na body IXP pro peering). Také někteří velcí klienti těchto ISP potřebují BGP, zejména pokud jsou napojeni na několik ISP. Způsoby použití BGP můžeme rozdělit do těchto případů:

- Single-homed: firma má pouze jediného ISP, jediné připojení k Internetu.
- Dual-homed: firma má sice jediného ISP, ale z její sítě k němu vede několik spojů.
- Single-multihomed: firma má několik ISP, ke každému vede jeden spoj.
- Dual-multihomed: firma má několik ISP, ke každému může vést více než jeden spoj.

 *Single-homed*: Jedná se rozhodně o stub network (tj. síť, která není průchozí). V tomto případě většinou nemá smysl používat BGP, s jedinou výjimkou: pokud ISP využívá pro řízení připojení klientů protokol BGP (místo statického směrovacího záznamu), pak firma musí použít BGP pro ohlašování své sítě směrem k ISP, a ISP přes BGP sděluje routeru firmy default route (tedy adresu souseda, přes kterého vede cesta na Internet).

 *Dual-homed*: Jednodušší případ by znamenal, že firma sice má ke svému ISP dva spoje, ale defaultně využívá jen jeden z nich, druhý je záložní a funguje jen v případě, že ten první selže. Tento případ se dá vyřešit pomocí statických cest, přičemž ta defaultní by měla nastavenou lepší metriku, případně menší číslo AD (administrative distance). Konfigurace se liší podle toho, jestli každý ze spojů vedoucích k ISP je na jiném routeru nebo jestli vedou z téhož routeru.

Pokud však chceme zároveň využívat oba spoje (nebo více než dva) a určitým způsobem vyvažovat zátěž, nebo pokud chceme určovat spoj podle konkrétního cíle (tedy pro různé cíle vybírat různé spoje), pak přichází v úvahu použít BGP. Přes BGP se ohlašujeme svému ISP a on nám může posílat

- buď jen *default route* (stejně jako v předchozím případě),
- nebo *full updates*, tedy celé BGP směrovací tabulky, ovšem routery musejí být dostatečně dimenzované (procesor, paměť), aby si s nimi dokázaly poradit,
- nebo *partial updates*, tedy default route a některé vybrané cílové sítě, kde default route není vyhovující.

 *Single-multihomed*: Také v tomto případě si můžeme vystačit s default routes nebo můžeme využít BGP, je to vlastně podobné prvním dvěma případům. Z pohledu vnitřní sítě je víceméně jedno, jestli máme jednoho nebo dva ISP, důležitější je počet spojů vedoucích „ven“. Rozdíl v konfiguraci bude samozřejmě na hraničním routeru (resp. routerech, opět záleží, jestli všechny tyto spoje vedou z jediného routeru nebo každý z jiného).

 *Dual-multihomed*: Máme tedy několik ISP, ke každému z nich vede z naší sítě několik spojů. Cílem je co nejvyšší míra redundance a máme takto ošetřeny různé situace – selže spoj k ISP, pak najede záložní spoj; selže ISP (například tehdy, když nezvládne kybernetický útok), pak použijeme jiného ISP. Nebo také můžeme dělit síťový provoz mezi několik spojů či ISP. V každém případě je vhodné, aby vše fungovalo co nejvíce automaticky, protože administrativní zásahy prováděné člověkem obvykle zbytečně dlouho trvají. Dual-multihomed situace bude proto velmi pravděpodobně řešena pomocí BGP, třebaže existují i jiné možnosti včetně statického směrování.

 **Jak se BGP používá.** Routery využívající BGP udržují „sousedské vztahy“ (peering) – ke svým sousedům (peers) pravidelně vysílají zprávy *KeepAlive*, jejichž účel je ujišťování o vlastní funkčnosti. Své sousedy však nezjišťují automaticky, sousedství musí být ručně nakonfigurováno, v konfiguraci zadáváme pro každého souseda jeho IP adresu a číslo ASN. Při první výměně směrovacích informací router informuje o svém *prefixu* a čísle ASN a obdrží celou směrovací tabulku, při změnách pak jen části směrovací tabulky, kterých se změna týká.

BGP rozlišuje dva typy sousedů: internal a external. Internal (iBGP) peer je soused ve stejném autonomním systému, tedy v konfiguraci u něj máme stejné ASN jako router, na kterém konfiguruje, kdežto external (eBGP) peer je soused nacházející se v jiném autonomním systému a jeho ASN bude jiné.

Při oznámení svého prefixu směrem k eBGP sousedovi směrovač doplní k prefixu své ASN. Během předávání mezi směrovači každý směrovač také přidá své ASN, tedy celková délka identifikačního řetězce

cesty, `AS_SEQUENCE`, se neustále prodlužuje. Vyšší prioritu má kratší cesta, tedy cesta s menším počtem ASN, vedoucí přes méně BGP směrovačů.

Protože ke každé cestě eviduje sekvenci čísel ASN, pro daný cíl je vybrána většinou cesta, která má tuto sekvenci kratší. Metrika se ve skutečnosti počítá z více parametrů, ale tento je nejdůležitější.

 Zatímco jiné směrovací protokoly z důvodu větší pružnosti využívají nespolehlivé přenosy (UDP) nebo částečně spolehlivé (Cisco RTP), pokud tedy vůbec využívají nějaký transportní protokol, BGP využívá spolehlivý přenos (zapouzdřuje do TCP, číslo portu 179). TCP spojení pro BGP říkáme *BGP session*.

Důvodem využití spojované komunikace je charakter WAN sítí, ve kterých se používá – tam se jinak než spojovaně nekomunikuje. Dalším důvodem je to, že sousedé BGP často nebývají v téže síti (na rozdíl od vnitřních směrovacích protokolů) a TCP spojení je pro cestu vedoucí přes několik mezilehlých routerů vhodnější.

Po prvním navázání TCP spojení s peerem (BGP sousedem) se přenáší celá BGP směrovací tabulka, kdežto v následujících BGP sessions už jen updaty – změny oproti předchozímu poslanému stavu.

 Jak vidíme, protokol BGP nelze přesně zařadit mezi algoritmy vektoru vzdáleností ani algoritmy stavu spoje. Využívá se *algoritmus vektoru cest*, který je podobný algoritmu vektoru vzdáleností, ale poněkud složitější a místo routerů na cestě se zabývá čísly ASN. Pokud k některému autonomnímu systému vede více cest, většinou vybere tu, která vede přes méně „průchozích“ autonomních systémů.



Poznámka:

Na protokol BGP je třeba si dát pozor. Protože se při jeho použití vlastně napojujeme do směrovací domény celého internetu, musíme dát velký pozor, jaké informace do internetu přes BGP pouštíme. Nejde jen o to, co nemá být vidět „z bezpečnostních důvodů“, ale při neopatrnosti můžeme ovlivnit fungování celého internetu. To už bylo průsvihů – legendární je například problém s chybnou konfigurací BGP směrovačů v Pákistánu, který způsobil kromě jiného několikahodinový výpadek serverů YouTube.



Další informace:

- https://www.cisco.com/c/en/us/td/docs/ios-xml/ios/iproute_bgp/configuration/xr-16/irg-xr-16-book/configuring-a-basic-bgp-network.html
- BGP Fundamentals. Online. <https://www.ciscopress.com/articles/article.asp?p=2756480>
- <https://www.noction.com/blog/bgp-hijacking>
- <https://inl.info.ucl.ac.be/blogs/08-02-25-bgp-misconfigurations-strike-back-pakistan-and-affect-youtube.html>
- <https://www.networkworld.com/article/2272520/six-worst-internet-routing-attacks.html>
- <https://www.9tut.com/basic-bgp-configuration>



7.3.9 Softwarový router

Routery bývají obvykle hardwarové, ale pro tyto účely lze použít i starší počítač (směrování v menší síti není moc výpočetně náročné) se speciálním softwarem. Pokud si vystačíme se statickým směrováním, pak se dá použít to, co je v každé linuxové distribuci. Pro dynamické směrování se může hodit něco navíc.

Quagga je open-source software pro UNIXové systémy (včetně Linuxu, FreeBSD, NetBSD, Solarisu) podporující kromě jiného protokoly OSPF, RIP, BGP. Quagga vznikla jako fork staršího projektu Zebra a hlavní daemon Quaggy se dodnes jmenuje **zebra**. Ovšem tento projekt už není aktivní.

*FRRouting*² (Free Range Routing) vznikl pro změnu jako fork projektu Quagga. V UNIXových systémech se stal velmi populárním. Přidává do UNIXového jádra podporu protokolů BGP, OSPF, RIP, IS-IS, EIGRP (v alfa verzi), VRRP, NHRP a dalších. FRRouting míří nejen do sítí ISP, ale také do cloudu včetně SaaS, web 2.0 apod.

*BIRD*³ (rekurzivní zkratka BIRD Internet Routing Daemon) je původně český projekt, za kterým stojí CZ.NIC. Ze směrovacích protokolů podporuje BGP, RIP, OSPF a další, dokáže pracovat s více směrovacími tabulkami (což je ostatně typické pro každý Linux, podívejte se do přílohy B v tomto dokumentu). Dá se nainstalovat do Linuxu, ale také dalších UNIXových systémů (včetně FreeBSD).



Další informace:

- <https://teklager.se/en/best-free-linux-router-firewall-software/>
- <https://devconnected.com/how-to-configure-linux-as-a-static-router/>
- <https://www.root.cz/serialy/linux-jako-internetova-gateway/>



7.4 Mechanismus DNS

 *Jmenná služba* je obecným principem překladu jmenných názvů (textových řetězců) na číselné adresy. Typickým představitelem jmenných služeb je DNS, ale existují i další jmenné služby, které jsou nicméně postupně odbourávány a nahrazovány službou DNS – například WINS (Windows Internet Name Service, slouží k překladu názvů NetBIOS na IP adresy), Microsoft doporučuje přechod na DNS.

 Mechanismus DNS je typickým zástupcem distribuovaných systémů v oblasti počítačových sítí. Datábase jmenných názvů, se kterými DNS pracuje, je distribuovaná na mnoha síťových zařízeních v síti, přičemž každé zařízení si eviduje především adresy ze své vlastní sítě – *princip lokality*. Distribuovanost zde funguje především proto, že systém DNS adres je hierarchický.

Dále budeme předpokládat, že čtenář ví, co je to DNS a jak vypadá doménová adresa (jmenná adresa serveru). Soustředíme se na základní definice, funkčnost celého systému a formát PDU.

7.4.1 Domény a jmenné adresy

Nebudeme zde opakovat, co je to doména a doménové jméno, jaké musí splňovat podmínky, co je to TLD, SLD doména, FQDN apod., předpokládáme, že každý čtenář to zná.

 Existuje možnost používání ne-ascii znaků v názvech domén (*IDN* – International Domain Names), ale překlad (nahrazení těchto znaků probíhá v klientské aplikaci, nevyskytují se v zónových souborech). Z toho vyplývá, že IDN musí být v klientské aplikaci implementován, včetně lokalizace pro danou znakovou sadu. Pokud tomu tak není, aplikace nedokáže s adresou *vůbec* pracovat (tj. ani zobrazit žádnou stránku z příslušného serveru).

²<https://frrouting.org/>

³<https://bird.network.cz/>

IDN se týká pouze názvů domén. Zbytek adresy také může obsahovat ne-ascii znaky (pozor, často nefunguje, není zatím plně dopracováno), ale jde o IRI (International Resource Indicator).

7.4.2 Zóny a DNS servery

 Jednu nebo více domén sdružujeme do společné *zóny* (typicky to bývá část podstromu domén). Pro každou zónu je stanovena konkrétní *autorita*, tedy zodpovědná organizace. Členění na zóny (coby skupiny domén) souvisí s odpovědností a technickou správou, zóny se obvykle nepřekrývají, až na jejich hranice (mezi zónami musí existovat určitá komunikace). Například celý podstrom `slu` tvoří zónu ve správě Slezské univerzity, kdežto nadřizená doména `cz` patří do zóny spravované organizací CZ.NIC.

 Správu zóny zajišťujeme na *jmenném serveru* – *DNS serveru*, který plní tyto úlohy:

- vést tabulku hostitelů (host table), je uložena v *zónovém souboru* s informacemi o zóně,
- odpovídat podle této tabulky na DNS dotazy, například překládat jmennou adresu na IP adresu.

V zóně musíme mít jeden *primární DNS server* a dále tam můžeme mít pomocné servery – *sekundární DNS servery* a *caching-only servery*. Ve vztahu k zónovému souboru:

- Primární DNS server vede a garantuje zónový soubor s tabulkou hostitelů, změny v tabulce provádíme vždy právě na tomto severu a záznamy na tomto serveru jsou vždy důvěryhodné, *autoritativní*.
- Sekundární DNS servery si v pravidelných intervalech kopírují zónový soubor z primárního DNS serveru (synchronizují, tomu se říká *zone transfer*), jejich účelem je rozložení zátěže, aby primární server nebyl příliš zatěžován. Jejich odpověď je také *autoritativní*.
- Caching-only servery neobsahují zónový soubor. Když takový server obdrží dotaz na adresu, „zeptá se“ primárního nebo sekundárního serveru, ale odpověď si na určitou krátkou dobu uloží do cache a pokud je v krátké době tázán na tutéž adresu, použije údaj z cache (takže se nemusí dotazovat dále). Odpověď caching-only serveru není autoritativní, ale většinou dostačuje.

 Nejznámější DNS servery jsou BIND, CoreDNS (je doporučený pro Kubernetes), DNSmasq (odlehčený DNS server pro malé sítě), Knot DNS (projekt CZ.NIC), NSD (je také testovacím serverem pro DNSSEC) a další, v serverových Windows máme také roli DNS Server.

 *DNS cache* používají nejen klienti (hlavně ve Windows, kdežto v Linuxu to není povinná součást), ale také DNS servery. Používá se pro uložení DNS záznamů, které dané zařízení nemá jinak uloženy (například v zónovém souboru nebo v souboru `.../etc/hosts`. Také DNS servery ji potřebují, protože jsou na ně směrovány i dotazy na záznamy z jiných zón.

Záznam se dostává do cache obvykle tehdy, když zařízení dostalo odpověď od některého DNS serveru (ta odpověď se uloží do cache). V cache ovšem nezůstane napořád, ale po určité době se odstraní. Tato doba je dána hodnotou TTL daného záznamu.

 *DNS resolver* (dotazovač) je komponenta, která zajišťuje dotazování (vytváří dotazy podle zadání od procesů, odesílá, přijímá odpovědi a předává procesům). Tuto komponentu najdeme v každém systému, který je připojen k síti používající jmenné názvy, tedy i v obyčejném počítači (včetně počítačů s Windows nebo mobilních zařízení) a v DNS serverech.

Nejdůležitější informace, kterou DNS resolver potřebuje, je adresa nejbližšího nadřizeného DNS serveru. Právě tomuto serveru odesílá dotazy, které pak mohou být buď přímo zodpovězeny nebo předávány.

Jak víme, v koncových zařízeních máme soubor `/etc/hosts`, resp. `... \system32\drivers\etc\hosts`, což je seznam dvojic IP adresa + jmenná adresa, které tam lze „staticky“ uložit a urychlit překlad, nebo naopak určitý překlad zablokovat (nefunguje, pokud blokový uživatel ví, jak to obejít).

7.4.3 DNS dotazy

DNS server, jak bylo výše uvedeno, si v zónovém záznamu vede informace o své doméně, obsažených subdoménách, a také informaci o nadřazeném serveru (tj. DNS serveru nadřazené domény) a kořenových DNS serverech.

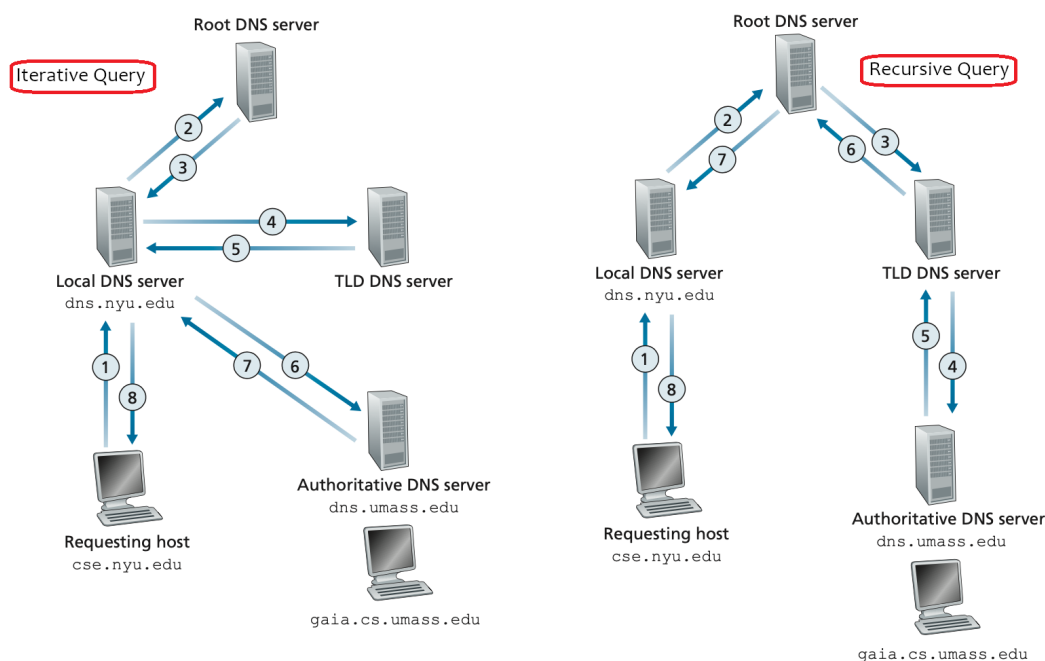
 Rozlišujeme *dopředné vyhledávání* (nalezení IP adresy podle doménového jména) a *zpětné vyhledávání* (reverzní, nalezení doménového jména podle IP adresy). Reverzní vyhledávání se provádí přes doménu `in-addr.arpa` pro IPv4 adresy, `ip6.arpa` pro IPv6 adresy. Zde je hlavním problémem to, že jak jmenné, tak IP adresy jsou hierarchické, ale každá „v jiném směru“ – doména TLD ve jmenné adrese je zcela vpravo, kdežto adresu sítě (hlavní část) v IP adrese najdeme zcela vlevo.

Dotazování v dopředném vyhledávání (máme jmennou adresu a chceme IP adresu) probíhá distribuovaně – DNS server často nedokáže na dotaz odpovědět okamžitě podle svých záznamů, proto se obrací nebo odkazuje na některý z kořenových serverů zadané domény první úrovně.

 Rozlišujeme dva typy dotazů:

- rekurzivní dotazy – tazatel dostane již hotovou odpověď,
- iterativní dotazy – tazatel postupně získává odkazy na místa, kde může dostat odpověď.

Tazatel (ve skutečnosti resolver) svůj dotaz (jaká je IP adresa zařízení se jménem `www.abcd.cz`?) posílá tzv. *lokálnímu DNS serveru* (nejbližšímu, který zná). Pokud tento server již zná odpověď, poskytne ji. Pokud ne, další postup se liší podle používaného typu dotazování:



Obrázek 7.3: Srovnání iterativního a rekurzivního dotazování⁴

⁴Podle https://gaia.cs.umass.edu/kurose_ross/interactive/dns_query.php

 **Rekurzivní dotaz.** Pokud DNS server nezná odpověď, obrátí se na kořenový server nebo DNS server té domény z adresy, kterou ještě zná (to může být TLD). Dotazovaný DNS server se chová stejně – pokud zná cíl, odpoví adresou, pokud cíl nezná, dotazuje se dále (obvykle ve svých subdoménách). Takto rekurzivně je dotaz posílán po stromě domén směrem dolů. Když se konečně dostane k DNS serveru, který zná odpověď, tento server odešle odpověď přímo lokálnímu DNS serveru, na kterém dotazování začalo, a ten ji odešle tazateli.

Takže postup je následující:

- Pokud je dotazovaný název z podřízené zóny (pravá část adresy je stejná), najde jmenný server ve směru zprava doleva první doménu v FQDN, která není jeho. Má odkazy na všechny domény na hranici s podřízenou zónou (například v zónovém souboru pro `cz` je odkaz na DNS server domény `slu`), tedy předá dotaz DNS serveru této podřízené domény. V podřízené doméně je dotaz vyřízen nebo opět předán níže. Rekurzivně směrem dolů je nalezena odpověď, která je pak stejnou cestou distribuována zpět tazateli.
- Pokud dotazovaný název není z vlastní ani podřízené zóny, předá DNS server dotaz DNS serveru z nadřízené zóny. Ten dotaz buď vyřídí nebo předá některému podřízenému nebo nadřízenému uzlu. Po vyřízení je odpověď opět distribuována stejnou cestou k tazateli.

Příklad

Pokud se některé zařízení z domény `fpf.slu.cz` dotáže na IP adresu serveru `www.fpf.slu.cz`, DNS server domény `fpf.slu.cz` hned odpoví, protože tento webový server patří do jeho zóny.

Pokud se některé zařízení z domény `fpf.slu.cz` dotáže na IP adresu serveru `www.cuni.cz`, nejdřív se dotazem zabývá DNS server domény `fpf.slu.cz`. Adresa je z jiného podstromu, tedy dotaz přepošle do domény `slu.cz`. To je podobný případ, proto je dotaz přeposlán DNS serveru domény `cz`. Tento server již pozná, že dotazovaná adresa patří do podřízené domény `cuni.cz` (protože pravá část adresy – `cz` – je shodná), v zónovém souboru najde kontakt na DNS server z podřízené domény, který již najde odpověď – IP adresu webového serveru `www.cuni.cz`, protože ten patří do jeho zóny. Odpověď půjde tazateli zpět stejnou cestou.



 **Iterativní dotaz.** Zde je aktivita především na straně tazatele. Pokud dotazovaný DNS server (kterýkoliv, nejdřív lokální) nezná odpověď, nedotazuje se dále, ale tazateli odešle místo odpovědi seznam DNS serverů, které by mohly znát odpověď („nevím, zeptej se serveru/ů xxxx“, doporučí buď kořenový server, některý TLD, anebo servery ze svých subdomén). Tazatel si vybere některý z doporučených DNS serverů podle toho, jak odpovídají doménové adrese, a táže se dále. Každý dotazovaný server opět pošle buď přímo odpověď, a nebo doporučení s názvy dalších DNS serverů.

Příklad

Pokud potřebujeme zjistit IP adresu počítače `www.firma.cz`, obrátíme se opět na lokální DNS server. Ten nám doporučí, abychom dotaz odeslali na některý z kořenových serverů domény `cz`, což uděláme. Dotazovaný kořenový server buď adresu dohledá ve svých záznamech, anebo vrátí pouze adresy name serverů, které mohou znát odpověď.

Takto distribuovaně probíhá vyhodnocování dotazu (postupně podle vnořování domén různých úrovní), poslední tázaný name server vrátí požadovanou IP adresu.



 V reálném světě se rekurzivní a iterativní způsob dotazování kombinuje. Rekurzivní způsob se používá v prvních fázích cesty, kdy se táže jednoduchý resolver, v dalších fázích probíhá iterativní dotazování (když se táže plnohodnotný DNS server).

 Ve skutečnosti dotazy na jinou TLD spíše směřují na kořenové servery (tj. servery, na kterých jsou informace o doménách první úrovně). Kořenový server opravdu není jen jeden, kopie této maximální kopie se nacházejí na mnoha zařízeních po celém světě. Informace o nich najdeme například na <https://root-servers.org/>.

7.4.4 Tabulka hostitelů

 Klíčovým prvkem pro DNS překlad je *tabulka hostitelů* (host table) v zónovém souboru. Můžeme si ji představit jako tabulku záznamů, kterým říkáme *DNS záznam*, anglicky RR – *Resource Record*. Každý RR obsahuje tyto údaje (ve sloupcích):

- typ záznamu,
- jmenná adresa nebo obecně to, co je v dotazu,
- IP adresa nebo obecně to, co má být v odpovědi,
- TTL (doba platnosti záznamu pro případ, že záznam je uložen do cache),
- třída (class) – většinou „IN“ jako Internet,
- případně další údaje k záznamu.

Údaj TTL je životnost záznamu v sekundách pro případ, že je tento záznam stažen do cache jiného zařízení (ať už koncového zařízení nebo caching-only serveru).

 Nejpoužívanější typy RR jsou:

- typ A – pro překlad jmenné adresy na IPv4 adresu,
 - typ AAAA – pro překlad jmenné adresy na IPv6 adresu (čtyři „A“) jsou tam proto, že IPv6 adresy jsou čtyřikrát delší než IPv4 adresy),
 - typ CNAME (Canonical Name) – pro překlad aliasu na kanonické jméno,
 - typ PTR – pro reverzní (obrácený) překlad, kdy známe IP adresu, ale potřebujeme jmennou adresu,
 - typ NS (Name Server) – v tomto záznamu se dozvíme IP adresu autoritativního DNS serveru pro některou doménu (okolní domény nebo takové, se kterými se v naší doméně hodně komunikuje),
 - typ MX (Mail eXchanger) – adresa e-mail serveru pro danou doménu,
 - typ SOA (Start of Authority) – administrativní informace o doméně,
 - typ TXT – záznam typu „děvče pro všechno“, používá se například pro ukládání veřejných klíčů, pomocí kterých si kdokoliv může ověřit, jestli e-mail poslaný z naší domény opravdu prošel naším mail serverem nebo byl podvržen,
 - několik typů záznamů je určeno pro uložení veřejných klíčů pro asymetrickou kryptografii nebo mají jiný „bezpečnostní“ účel,
- atd.

Většinou potřebujeme právě záznamy typu A nebo AAAA, kdy známe jmennou adresu a potřebujeme buď IPv4 adresu nebo IPv6 adresu.

Příklad

Při zpětném (reverzním) vyhledávání se používají záznamy typu PTR, které jsou „schovány“ pod top-level doménou `in-addr.arpa`. Ukážeme si, jak takový záznam vypadá.

Účelem je vyřešit to, že IP adresy jsou hierarchické v opačném směru než jmenné adresy. Zatímco u jmenné adresy se nejvyšší doména zapisuje zcela vpravo (na konci řetězce vnořených domén) a podřízené domény jsou směrem vlevo, u IP adresy se síť nejvyšší úrovně naopak vlevo (na začátku adresy) a podsítě jdou směrem vpravo. DNS je stavěný pro domény, tedy formát IP adresy v `in-addr.arpa` se tomu musí přizpůsobit.

Budeme pracovat v Linuxu, protože Windows před uživatelem některé aspekty poněkud „skrývají“. Nejdřív provedeme běžný dotaz programem `nslookup` na adresu Googlu.

```
sarka@sarka-VB-pc:~$ nslookup google.com
```

```
...
```

```
Non-authoritative answer:
```

```
Name: google.com
```

```
Address: 142.251.36.132
```

```
Name: google.com
```

```
Address: 2a00:1450:4014:80b::200e
```

Ted' provedeme zpětný dotaz na IPv4 adresu:

```
sarka@sarka-VB-pc:~$ nslookup 142.251.36.132
```

```
132.36.251.142.in-addr.arpa name = prg03s12-in-f4.1e100.net.
```

Všimněte si dvou věcí:

- IP adresa byla přeložena na jinou adresu než `google.com`. Důvodem je, že Google (stejně jako jiné velké platformy) používá CDN síť pro transport obsahu po celém světě, získaná adresa je jeden z blízkých CDN serverů.
- V PTR záznamu je sice hledaná IP adresa, ale její oktety jsou v opačném pořadí a na konci je řetězec `in-addr.arpa`.

Ano, v PTR záznamu je IP adresa opravdu v opačném pořadí, protože hlavní síť potřebujeme vpravo, jak je tomu při zápisu domén.

Příkaz `dig` má podrobnější výstup. Použijeme ho také pro zpětný překlad IP adresy `142.251.36.132`.

```
sarka@sarka-VB-pc:~$ dig -x 142.251.36.132
```

```
; <<>> DiG 9.16.48-Ubuntu <<>> -x 142.251.36.132
```

```
;; global options: +cmd
```

```
;; Got answer:
```

```
;; ->>HEADER<<- opcode: QUERY, status: NOERROR, id: 35852
```

```
;; flags: qr rd ra; QUERY: 1, ANSWER: 1, AUTHORITY: 0, ADDITIONAL: 1
```

```
;; OPT PSEUDOSECTION:
```

```
; EDNS: version: 0, flags:; udp: 65494
```

```
;; QUESTION SECTION:
```

```
;132.36.251.142.in-addr.arpa. IN PTR
```

```
;; ANSWER SECTION:
```

```
132.36.251.142.in-addr.arpa. 85049 IN PTR prg03s12-in-f4.1e100.net.
```

```
;; Query time: 8 msec
```

```
;; SERVER: 127.0.0.53#53(127.0.0.53)
```

```
;; WHEN: Thu May 09 15:02:45 CEST 2024
```

```
;; MSG SIZE rcvd: 94
```

Červenou jsou zvýrazněna umístění dotazovacího řetězce v částech Question a Answer. Jak vidíme, v PTR záznamu jsou oktety IP adresy opravdu v opačném pořadí.

Ted' se zaměříme na IPv6 adresu (tu jsme také získali z prvního příkazu v tomto příkladu).

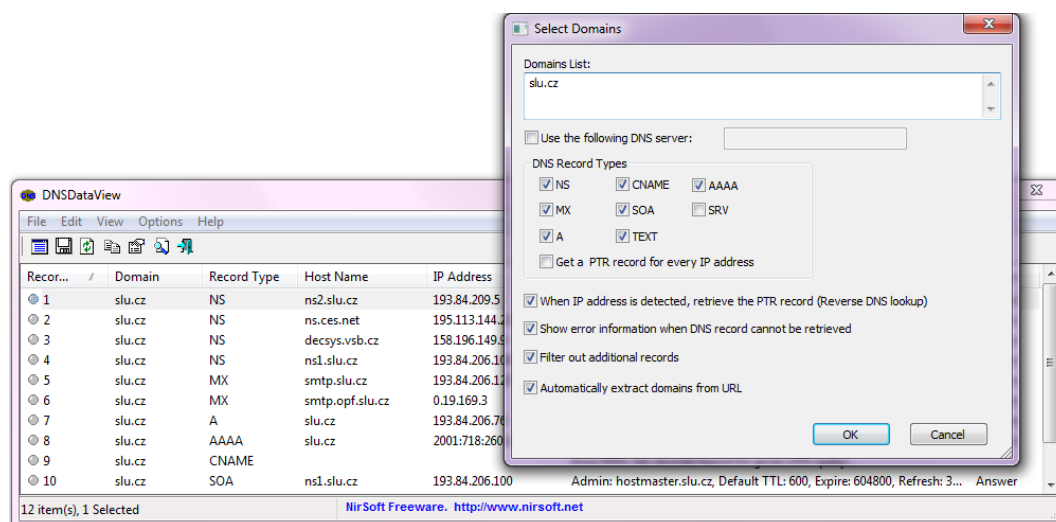
```
sarka@sarka-VB-pc:~$ nslookup 2a00:1450:4014:80e::2004
4.0.0.2.0.0.0.0.0.0.0.0.0.0.e.0.8.0.4.1.0.4.0.5.4.1.0.0.a.2.ip6.arpa name = prg03s12-in-x04.1e100.net.
```

Hlavní doména PTR záznamu pro IPv6 adresu je `ip6.arpa`. Další rozdíl je, že „subdomény“, které v případě IPv4 byly prostě jednotlivé oktety adresy v opačném pořadí, jsou zde sice taky v převráceném pořadí, ale po jednotlivých hexadecimálních číslicích. Navíc jsou zahrnuty i ty číslice 0, které se běžně v adrese nezapisují. V adrese tedy máme před doménou `ip6.arpa` právě 32 jednoznakových subdomén, protože právě tolik hexadecimálních číslic má IPv6 adresa bez použití krácení.



Příklad

Pro přístup k DNS záznamům můžeme použít příkaz `nslookup` nebo v UNIXových systémech lépe příkaz `dig`, ale zajímavou alternativou s grafickým rozhraním je aplikace *DNS Data View*⁵ od společnosti Nirsoft. Okno s výstupem pro doménu `slu.cz` vidíme na obrázku 7.4.



Obrázek 7.4: Program Nirsoft DNS Data View



Další informace:

<https://phoenixnap.com/kb/dns-record-types>



7.4.5 Bezpečnost v DNS

V poslední době se hodně diskutuje o zabezpečení DNS. Překladač jmenné adresy na IP adresu obvykle věříme a mnoho uživatelů vůbec nenapadne, že ve skutečnosti mohou komunikovat s protistranou s jinou

⁵Ke stažení na http://www.nirsoft.net/utls/dns_records_viewer.html

IP adresou, než kterou předpokládají. Důsledkem napadení DNS může být například provedení útoku Man-in-the-Middle, přeměrování na napadený server za účelem šíření škodlivého kódu, odposlechnutí hesla, čísla karty apod.

V rámci zabezpečení DNS je tedy třeba zajistit, aby DNS zprávy s informací o překladu nemohly být podvrženy. Je to v podstatě obdoba mechanismu SEND, který zabezpečuje ARP.

 **DNSSEC** (DNS Security Extensions) je bezpečnostní rozšíření systému DNS umožňující DNS klientovi ověřit původ dat (zda pocházejí od správného DNS serveru).

Při použití DNSSEC se používá asymetrické šifrování následovně:

- provozovatel DNS domény vygeneruje dvojici klíčů (soukromý a veřejný klíč),
- veřejný klíč uloží u nadřazené autority své domény (vytváří se jakási hierarchie důvěry),
- soukromým klíčem šifruje DNS server zprávy, které odesílá,
- příjemce (DNS klient, resolver) je dešifruje veřejným klíčem, který získal od nadřazené autority DNS serveru.

Jedná se vlastně o elektronicky podepsané DNS pakety.

Také správce české TLD domény na svých DNS serverech používá DNSSEC, přičemž jeho nadřazenou autoritou (a místem, kde je příslušný veřejný klíč) je správce celosvětové sítě DNS serverů. Nadřazená autorita ručí za správnost překladu svých bezprostředních podřízených.



Další informace:

Další informace o DNS:

- <http://www.fit.vutbr.cz/study/courses/ISA/public/xkupci00.pdf>
- <https://www.lupa.cz/clanky/neplechy-v-dns-a-jak-se-jim-vyhnout/>
- <https://www.dnssec.cz/>
- <https://www.samuraj-cz.com/clanek/dns-domain-name-system-zamereno-na-microsoft/>
- <https://www.root.cz/clanky/dnssec-cast-prvni-aneb-je-potreba-zacit-od-piky/>
- <https://www.abclinuxu.cz/clanky/site/nastaveni-dns>



7.4.6 Využití DNS při zabezpečení e-mailové komunikace

Naopak mechanismus DNS se používá k obraně před spamem či podvrženými e-maily. V zónovém souboru můžeme mít speciální záznamy typu TXT nesoucí v sobě informaci, která podobné nepravosti umožní odhalit.

 **SPF** (Sender Policy Framework) slouží ke stanovení povolených e-mailových serverů v dané doméně. Příslušný TXT záznam nese informaci

```
v=sfp1 include:...
```

obsahující seznam e-mail serverů v doméně a kdokoliv se tento záznam může vyžádat. Mohou být uvedeny jmenné i IP adresy. V případě, že v doméně lze používat i „cizí“ poskytovatele e-mailových služeb (například veřejné služby typu Gmail), taky by to mělo být v záznamu uvedeno, jinak mohou být při odchodu z naší domény považovány za nedůvěryhodné.

Pokud adresát (mail server, který spravuje schránky uživatelů) zjistí, že příchozí e-mail byl (podle svého záhlaví) odeslán e-mailovým serverem z určité domény, ale v SPF TXT záznamu vyžádaném z dané domény ten server není uveden, půjde určitě o podvržený e-mail.

 **DKIM** (Domain Keys Identified Mail) je ověřovací mechanismus používající kontrolu digitálním podpisem. E-mail server v naší doméně každou odesílanou zprávu (část záhlaví a tělo zprávy) digitálně podepíše svým soukromým klíčem, do záhlaví vloží tento podpis (pole DKIM-Signature). V zónovém souboru naší domény je TXT záznam s obsahem

```
v=dkim1;...
```

V datech následuje typ klíče (většinou rsa) a samotný veřejný klíč. Takže cílový mail server si v záhlaví příchozího e-mailu zjistí doménu odesílatele, od tamního DNS serveru si vyžádá DKIM TXT záznam, z něj vytáhne veřejný klíč a ten použije pro ověření digitálního podpisu. Pokud digitální podpis nesedí, pak e-mail byl ve skutečnosti odeslán z jiné domény, než kterou deklaruje, nebo byl odeslán sice z dané domény, ale nikoliv ze „správného“ mail serveru.

7.5 QoS

QoS (Quality of Service) je mechanismus prioritizace daných typů provozu v síti. Existují metody, které nabízejí garanci určitých parametrů poskytované služby. Garantovat lze například:

- šířku pásma pro přenos, volnou kapacitu sítě,
- minimální ztrátovost datových jednotek (případně téměř nulovou ztrátovost),
- dynamiku ztrátovosti (tedy aby nedocházelo ke ztrátám ve shlcích, to by mohlo mít negativní vliv na přenos multimédií, která občasnou ztrátu datové jednotky tolerují),
- zpoždění, latenci,
- změny – dynamiku – zpoždění (jitter), atd.

Běžné síťové protokoly (e-mail, přenos souborů apod.) obvykle nemají velké nároky na QoS, ale naopak VoIP a videopřenosům mohou vadit vyšší hodnoty zpoždění, a také vyšší hodnoty jitter.

Stav, kdy se QoS neimplementuje, označujeme *Best Effort* (BE).

Díky IETF existují od 90. let 20. století tyto dva mechanismy:

- IntServ (Integrated Services)
- DiffServ (Differentiated Services)

Kromě toho existují i další mechanismy garance kvality služeb. Vlastní mechanismy mají také WAN sítě, včetně ATM a MPLS.

 **IntServ** je starší jednoduchý mechanismus, který funguje takto:

- aplikace, která chce využít QoS, si rezervuje pásmo pro své přenosy,
- toto pásmo je aplikaci plně vyhrazeno po celou dobu trvání rezervace, ať už ho používá nebo ne,
- je simplexní, tedy rezervace platí v jednom směru komunikace.

Jednosměrnost není na závadu, protože se IntServ typicky používá pro multimediální přenosy, které jsou v principu asymetrické, případně není problém navázat dva simplexní spoje pro oba směry.

Rezervace by měla být platná na celé cestě, proto by příslušný protokol měl být podporován na všech uzlech, přes které cesta vede. Pokud je systém rezervací někde na cestě přerušen (třeba proto, že daný protokol tam není podporován), v daném úseku půjdou pakety běžným způsobem.

Nejznámější implementací mechanismu IntServ je *protokol RSVP*, který se používá pro rezervaci komunikačního pásma například jako doplňující protokol protokolu SIP.

Hlavní nevýhodou IntServ je to, že na jednotlivých L3 zařízeních jsou zdroje rezervovány bez ohledu na jiný provoz, toto řešení není škálovatelné.

 **DiffServ** pracuje jiným způsobem. Nerezervuje staticky šířku pásma ani nenutí aplikace k dopřednému oznamování svých potřeb, ale místo toho do záhlaví paketu dodává informaci o prioritě, přičemž tato informace bývá brána v úvahu na síťových prvcích na cestě tohoto paketu. PDU jsou tříděny do kategorií a každá kategorie má přiřazenu určitou prioritu nebo obecněji způsob zacházení při odbavení na síťových prvcích.

 Pracuje s 6 bity označenými zkratkou *DSCP* (Differentiated Service Code Point), což znamená 64 možných tříd přístupu (v praxi se využívá jen několik z nich). Třídy jsou shrnuty v několika kategoriích, z nichž jsou nejdůležitější tyto:

- *Urychlené předávání* (EF, Expedited Forwarding) – garance malého a téměř konstantního zpoždění, latence, šířky pásma; je složité, proto může být poskytnuto jen nízkému počtu přenosů. Typické využití je VoIP.
- *Zajištěné předávání* (AF, Assured Forwarding) – funguje podobně jako CIR a DE u Frame Relay. Používá se u služeb, které upřednostňují volitelnou úroveň QoS. Určitá propustnost je garantovaná, a co je nad ní, má vyšší prioritu při zahazování paketů.
- *Základní služba* (BE, Best Effort) – pro běžné datové přenosy.

Při vstupu do sítě jsou pakety klasifikovány, označeny třídou (pokud je to potřeba). Pokud ne (všechny bity jsou 0), jedná se o kategorii *Best Effort*.

 Klasifikace (zařazení do kategorie) probíhá při vstupu do sítě DiffServ (hovoříme o *DiffServ doméně*), na následujících směrovačích je tato informace pouze využívána. Router si řadí příchozí pakety do několika výstupních front (pro různé kategorie jsou různé fronty), s každou frontou zachází jinak podle jejího typu (přednost, nezahazovat při využití mechanismu Token Bucket, apod.). Jedna fronta je typu Best Effort, tam se řadí běžné pakety bez speciálního určení kategorie.

Tento způsob řízení QoS je jednoduše slučitelný s jinými QoS protokoly. Na hranici DiffServ domény (na hranovém, ingress edge, směrovači) probíhá výše popsaná klasifikace paketu, případně překlad z jiné QoS technologie (ATM, MPLS, atd.).

Konkrétní umístění pole bitů pro DiffServ závisí na tom, na které vrstvě a v kterém protokolu je tato technologie implementována (nemusí to být na síťové vrstvě s IP protokolem). Může to být uvnitř hlavičky datové jednotky (pokud je tam místo), třeba v záhlaví IP paketu, nebo vně (před záhlavím). Také s konkrétními hodnotami jednotlivých bitů se někdy musí zacházet poněkud volně, protože DiffServ hodnoty se často mapují do polí v záhlavích PDU, která mají jiný počet bitů než jak DiffServ určuje.

V IPv4 paketu máme 8bitové pole třídy služby TOS (z toho se první 3 bity využívají podle IEEE 802.1p), v IPv6 paketu jde o pole pro prioritu. V MPLS záhlaví jsou pro QoS pouze tři bity.



Další informace:

- <http://www.kiv.zcu.cz/~ledvina/Prednasky-PSI-2007/qos-text.pdf>
- https://www.cisco.com/en/US/technologies/tk543/tk766/technologies_white_paper09186a00800a3e2f.html
- <https://www.samuraj-cz.com/clanek/cisco-qos-1-uvod-do-quality-of-service-a-diffserv/>



7.6 VoIP a videotelefonie

7.6.1 Princip

 *VoIP* (Voice over IP, „internetová telefonie“, IP telefonie) je technologie pro přenos digitalizovaného hlasu v IP datagramech. Hovor může být veden i mezi více než dvěma účastníky. Koncová zařízení mohou být buď hardwarové VoIP telefony anebo jde o software, VoIP aplikaci.

 Telefonovat můžeme

- v rámci sítě poskytovatele,
- mimo síť poskytovatele, ale pořád v rámci datové sítě,
- na běžné telefonní číslo do veřejné telekomunikační sítě.

 Třetí možnost vyžaduje vytvoření spojení přes *VoIP bránu*, která tvoří rozhraní mezi datovou a veřejnou telekomunikační sítí. Volí se VoIP brána co nejbliž cíli, aby co nejdelší část spojení vedla po datových linkách (resp. aby část spojení po telekomunikačních linkách byla pokud možno „místní hovor“).

V současné době se setkáváme s komplexně řešenými pobočkovými ústřednami pro firmy, které si takto mohou zařídit vlastní VoIP síť. Jejich konkurencí jsou však nabídky mobilních operátorů, které jsou, zvláště pro větší firmy, velmi výhodné.

 Obvykle se vyžaduje splnění těchto požadavků:

- rychlost připojení k Internetu přes 128 kb/s
- maximální zpoždění na spoji k serveru poskytovatele 150 ms
- maximální kolísání zpoždění 30 ms
- ztrátovost datových jednotek pod 2 %, lépe kolem 1 %

Dále si může poskytovatel stanovit požadavky na způsob implementace NAT apod. Tyto požadavky jsou však relativní, záleží na celkovém vytížení spoje (proto je důležité používat QoS) a obvyklém počtu hovorů na spoji vedených. VoIP vyžaduje sice menší šířku pásma než datové služby, ale zato pokud možno konstantní.

 Nejdřív je nutné navázat spojení. Navazování spojení provádí příslušný aplikační protokol, a to buď SIP nebo H.323.

Po navázání spojení se přenášený zvuk nejdřív digitalizuje – provádí se vzorkování na 8 kHz. Přenos se provádí pomocí kodeku (coder/decoder), který provádí kódování signálu a zároveň jeho kompresi.

7.6.2 Protokoly

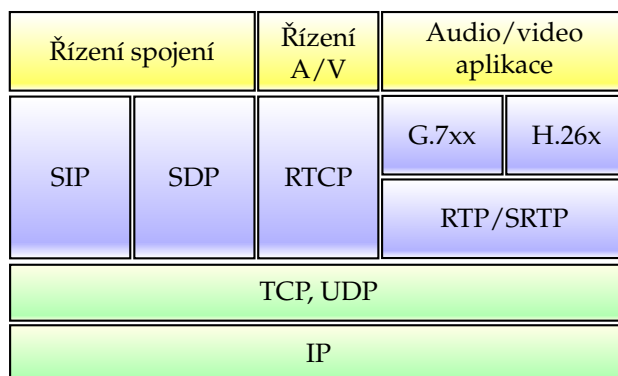
 **SIP** (Session Initiation Protocol) je textově orientovaný protokol vyvinutý přímo pro použití na Internetu (IETF), v roce 1999. Je to protokol aplikační vrstvy pro vytváření, udržování a ukončování interaktivních relací, kromě jiného VoIP. Sám o sobě nezajišťuje QoS, ale dokáže spolupracovat s protokoly, které jsou k tomu určeny. Jeho úkoly:

- lokalizuje příjemce volání,
- ověří vlastnosti použitého zařízení,
- zajistí přenos dat a zabezpečení u jiných protokolů.

SIP je *distribuovaný*, co nejvíc posouvá inteligenci přenosu ke koncovým zařízením.

Další významnou vlastností SIP je *pružnost*. Dokáže spolupracovat s mnoha protokoly a kromě VoIP je použitelný například i pro Instant Messaging a dokáže dobře spolupracovat s mobilními technologiemi. Lze používat jak číselné adresy (telefonní čísla), tak i URI (Universal Resource Identifier) – webovou adresaci.

Na transportní vrstvě je využíván především protokol UDP, což znamená vyšší rychlost přenosu (menší zpoždění).

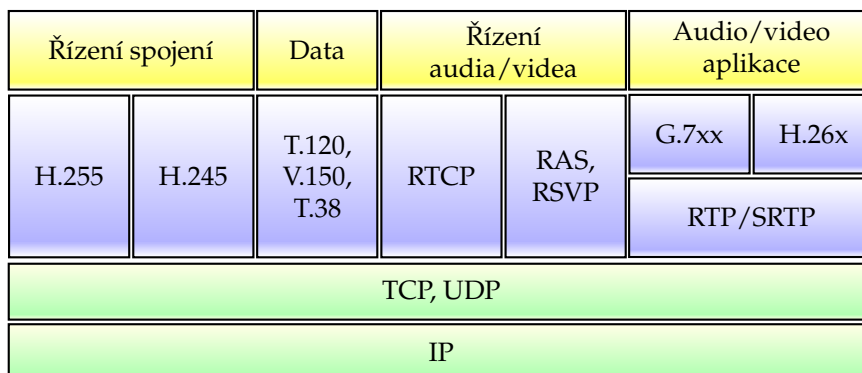


Obrázek 7.5: Protokolový zásobník pro SIP

 **H.323** je sada binárních protokolů pro konverzi signalizace paketového protokolu na signalizaci telefonní sítě. Je starší než SIP (rok 1996) a vychází ze signalizace obvyklé v telekomunikačních sítích. Jde o komplexnější přístup, H.323 řeší kromě navazování a udržování spojení také QoS a další charakteristiky přenosu. Při adresování používá telefonní čísla.

H.323 má větší zpoždění při navazování spojení a navíc jeho spolupráce s jinými protokoly je problematická. Jedná se o protokol do značné míry *centralizovaný*.

Na transportní vrstvě je využíván především protokol TCP, který má větší režii přenosu, to znamená větší zpoždování přenosu.



Obrázek 7.6: Protokolový zásobník pro H.323

Podpůrné protokoly transportní vrstvy:

- RTP (Real-time Transport Protocol) doplňuje činnost protokolu UDP v oblasti rychlé paketizace datového toku v reálném čase, synchronizuje přenos a umožňuje zjistit ztracený paket nebo nesprávné pořadí paketů.

- SRTP (Secure RTP) je bezpečnější varianta protokolu RTP.
- RTCP (Real-time Transport Control Protocol) zprostředkovává zpětnou vazbu pro protokol RTP, umožňuje využívat jeho výše popsané vlastnosti. Využitím RTCP lze regulovat datový tok, například dynamicky měnit rychlost přenosu podle požadavků přijímajícího uzlu.
- SDP (Session Description Protocol) se používá u protokolu SIP pro vyjednávání parametrů multimediální komunikace.
- RSVP (Resource Reservation Protocol) slouží pro rezervaci zdrojů, je to prostředek řízení QoS.
- G.711, G.729, G.723.1, atd. – audio kodeky, nad nimi běží audio aplikace.
- H.261, H.263, atd. – video kodeky, nad nimi běží video aplikace.
- T.120 – konferenční datové přenosy.
- RAS (Registration, Admission, Status) – podpůrný protokol pro vyjednávání a udržení spojení při použití H.323. Je součástí protokolu H.255.0 sloužícího k administraci multimediální komunikace.

Z dalších protokolů se využívá například STUN nebo TURN pro podporu NAT, SCCP u zařízení Cisco, IAX v systému Asterisk, a další.

7.6.3 Videotelefonie a videokonference

 *Videotelefonie* je vzájemný hovor dvou či více účastníků se současným sledováním synchronizovaného videostreamu (obvykle se snímá obraz videokamery). Dnes obvykle nejde ani tak o samostatná zařízení (videotelefony), setkáváme se hlavně se softwarovou implementací (například výše zmíněný Skype nebo FaceTime od Applu, obojí na principu VoIP). Funkce videotelefonie je také obvykle implementována v softwaru pro pobočkové ústředny (např. ve výše zmíněném Asterisku).

Videotelefonie stojí především na protokolu *H.264*. Tento protokol standardizovaný ITU-T je vlastně notoricky známým videokodekem pro kompresi videa, mnohým bude hodně říkat zkratka MPEG-4. Používá se všude tam, kde je potřeba přenášet komprimované digitální video.

 Zatímco videotelefonie vyžaduje vytvoření spoje typu point-to-point, *videokonference* je již náročnější způsob komunikace. Jde vlastně o spoje typu multipoint-to-multipoint, signál je šířen formou multicastu. Videotelefonii můžeme považovat za speciální jednodušší případ videokonferencí. Existují řešení využívající webovou aplikaci, mobilní aplikaci, desktopovou aplikaci, nebo to vše kombinují.

 Nejmenší šířku pásma (do několika desítek kb/s) vyžaduje přenos zvuku. Větší šířku pásma vyžaduje přenos obrazu i v nižší kvalitě. Kromě přenosu zvuku a obrazu existují i další možnosti, jak využít kapacitu multimediálních přenosů, například *sdílená tabule* (účastníci „kreslí“ na tabuli, ta je viditelná a přístupná všem účastníkům videokonference).

Existuje celá řada komplexních nebo jednoduchých řešení pro videokonference či obecně týmovou spolupráci, liší se v mnoha ohledech. Často se setkáváme i s možností vytváření oddělených „místností“ v rámci videokonference, do kterých se rozdělují účastníci, aby mohli pracovat odděleně v menších týmech. V anglických verzích se této možnosti říká „breakout rooms“.

 **MS Teams** je jednoduchý videokonferenční systém založený na cloudu, který je součástí MS Office. Nedá se nijak zvlášť široce konfigurovat, ale základní funkce obsahuje – vytváření týmů, plánování a navazování videokonferenčních hovorů, jejich nahrávání, sdílení plochy, sdílení souborů, funkce tabule. Uživatel je buď v roli vlastníka kanálu (učitele), nebo v roli člena nebo hosta (studenti).

 **BigBlueButton**⁶ je otevřený projekt umožňující provozovat videokonferenční systém ve vlastním privátním cloudu, ve skutečnosti se jedná o upravenou a obohacenou linuxovou distribuci (dá se získat i ve formě virtuálního stroje do VirtualBoxu pro různé operační systémy či VMWare Fusion na MacOS). Vznikl v univerzitním prostředí a je především určen pro on-line výuku. Umí videokonference, sdílení obsahu, sdílení plochy, chat, sdílenou tabuli. Rozlišují se dvě role: moderator (učitel) a viewer (student).

 **Google Zoom a Google Meet**⁷ jsou videokonferenční řešení od Googlu. Používají se při online výuce i při vzdálené týmové spolupráci ve firmách. Meet je vhodnější pro pracovní schůzky, konzultace či školení, kde je občas třeba sdílet i data (i v bezplatné verzi je součástí úložný prostor) a používat plánovací kalendář, Zoom se zase hodí pro větší videokonference.

V obou případech existuje verze zdarma, v placených verzích se přidává možnost nahrávání, delší doba jednoho meetingu, automatické vytváření prezenční listiny, větší úložiště apod.

 **Apache OpenMeetings**⁸ je řešení spíše pro komerční sféru než pro školství, nicméně se také jedná o videokonferenční systém. Vznikl rozšířením Red5 media serveru a nabízí videokonference, schůzky, sdílení plochy, sdílení dokumentů, atd. Je integrován v některých CMS systémech.

 **Cisco Webex Meetings**⁹ je řešení pro komerční sféru, placená verze zvládne videokonferenci s až 200 účastníky. Neplacená verze má určité limity, nicméně také je použitelná, když nemáme vysoké nároky. Webex se dá používat ve formě webové aplikace, ale desktopová aplikace nabízí širší možnosti konfigurace. Celkově se dá říct, že Webex je snad nejjednodušší konfigurovatelné řešení.

 **Jitsi Meet**¹⁰ je jednoduché open-source řešení pro menší týmy, které je postaveno na webovém rozhraní, takže není závislé na operačním systému (ale existuje mobilní aplikace). Nabízí běžné vlastnosti, které od videokonferenčních nástrojů čekáme, včetně nahrávání, sdílení obrazovky či okna, atd.

 **MBone**¹¹ je distribuovaný systém složený z volně dostupných nástrojů. Vyžaduje vhodně multimediálně vybavený počítač připojený do sítě, funguje nad protokolem IP. Podporuje různé operační systémy. Umožňuje organizovat videokonference, přihlásit se do existující videokonference, přenášet zvuk a obraz, sdílet text i multimédia (sdílená tabule).

 **Ekiga** (GnomeMeeting)¹² je svobodný software pro VoIP a videokonference vyvíjený v rámci projektu Gnome. Komunikuje i s jinými klienty podporujícími protokol SIP nebo H.323 (včetně MS NetMeeting). Klient je dostupný pro Linux i Windows.



Souhrn kapitoly

V této kapitole jsme se seznámili s pojmy centralizovaný, decentralizovaný a distribuovaný systém a dále jsme se zabývali především distribuovanými systémy – směrovacími mechanismy, systémem DNS, řízením kvality služby se zřetelem především na DiffServ, VoIP a videotelefonii.



⁶<https://bigbluebutton.org/>

⁷<https://support.google.com/meet>, <https://zoom.us/>

⁸<https://openmeetings.apache.org/>

⁹<https://www.webex.com/>

¹⁰<https://jitsi.org/jitsi-meet/>

¹¹<https://www.savetz.com/mbone/>

¹²<https://www.ekiga.org/>

Management sítě a monitoring

 *Rychlý náhled:* V této kapitole se zaměříme na správu sítě a monitorovací techniky a nástroje. Podíváme se na běžné protokoly, architektury, principy a nástroje – systémy NMS (Network Management Systems), protokoly pro získávání informací o provozu zařízení, jejich evidenci, logování, monitoring sítě a VPN. V tomto tématu budeme pokračovat také v následující kapitole o bezpečnosti.

 *Klíčová slova:* Správa sítě, NMS, agent a správce, sonda, CMIP, SNMP, MIB-II, trap, RMON, Syslog, monitoring sítě, Snort, NetFlow, VPN, IPSec, GRE, L2TP, OpenVPN, SSH, MPLS VPN, IP-in-IP, VPLS, DMVPN.

 *Cíle studia:* Cílem této kapitoly je porozumět principům správy a monitorování sítě. Cíl je pochopit činnost protokolů pro management sítě, NMS a jejich architekturu, a související postupy a software. Důležitou částí správy sítě je monitoring, tedy cílem kapitoly je obeznámit se se základy monitoringu. V této kapitole se zaměřujeme také na různá řešení VPN, tedy studenti se naučí principům VPN a získají přehled o existujících nástrojích.

8.1 Správa sítě

8.1.1 NMS

Správa sítě v sobě zahrnuje tyto činnosti:

- dohled, kontrola (monitorování) – na fyzické vrstvě (stav média apod.), linkové vrstvě (chybovost rámců, atd.) i vyšších, interpretace shromážděných údajů,
- diagnostika a řešení poruch, předcházení poruchám,
- správa jednotlivých prvků sítě, řešení změn topologie (například přidání nového uzlu),
- účtování využívání zdrojů.

V rámci modelu ISO/OSI většina těchto činností probíhá na aplikační vrstvě, a to vždy v některé formě na všech uzlech sítě.

Správa sítě je služba využívající řadu nástrojů, aplikací a zařízení k monitorování, údržbě a zabezpečování sítě, a to v co nejvyšší míře automatizace.

 Rozlišujeme *řídící* a *řízené entity*. Řízené entity odesílají zprávy obsahující hlášení o událostech nebo problémech, řídící entity na ně adekvátně reagují, například

- zpraví operátora mailem nebo jiným způsobem,
- zapíší událost do LOG souboru (logování událostí),
- vypnou nebo restartují systém, odhlásí „problémového“ uživatele, odstřelí „problémový“ proces,
- provedou automatickou úpravu konfigurace systému, přenastaví určené limity.

Softwarové agenty (řízené entity) sbírají informace z koncových zařízení a odesílají řídícím entitám.

 **Network Management System (NMS)** je systém řízení sítě zahrnující řídící entity, databázi a protokoly řízení sítě. Nejznámější protokoly pro NMS:

- SNMP (Simple Network Management Protocol)
- CMIP (Common Management Information Protocol)

8.1.2 ISO model řízení sítě

Řízení sítě lze rozdělit do několika oblastí.

 **Řízení výkonu (performance management)** – účelem je zajistit, aby výkon sítě byl na přijatelné úrovni. Patří zde proměnné veličiny jako je průchodnost sítě, lhůty pro odezvu uživatelů, optimální využívání kabelů a jiných spojových cest a prvků, atd. Zahrnuje:

1. shromáždění souvisejících dat o veličinách
2. analýza dat a stanovení úrovně pro normální provoz
3. stanovení horní (resp. u některých veličin dolní) hranice s tím, že překročení této hranice je indikováno jako problém, který je třeba řešit

Při překročení stanovených hranic je informován NMS.

 **Řízení konfigurace (configuration management)** znamená monitorování konfigurace sítě a připojených systémů (hardware i software) a jejich ovlivňování. Pro každou sledovanou veličinu (verze operačního systému s aktualizacemi, verze protokolu TCP/IP, apod.) je vytvořeno pravidlo, a pokud není splněno, musí být vyvozeny důsledky.

 **Řízení uživatelských účtů (accounting management)** znamená definování regulačních parametrů pro uživatele a skupiny uživatelů. Vhodná regulace minimalizuje problémy v síti, především při využívání zdrojů (hardware, vyhrazené místo v paměti, atd.), a maximalizuje férovost využívání zdrojů v síti vzhledem k ostatním uživatelům.

 **Řízení chyb (fault management)** představuje detekci, evidenci (log soubory), oznamování a, pokud je to možné, také automatické řešení problémů, které v síti mohou vzniknout. o jeden z nejdůležitějších modulů systémů řízení sítě, protože neodchycené chyby mohou způsobit zpomalení komunikace, ztrátu dat nebo obecně neakceptovatelné zhoršení celkového provozu na síti (případně zhroucení sítě).

Modul využívá údaje získávané a uložené v rámci ostatních funkcí řízení sítě, detekuje možné problémy (aby byly co nejdříve zachyceny), navrhuje a testuje řešení.

 **Řízení bezpečnosti (security management)** znamená řízení přístupu ke zdrojům na síti podle stanovených pravidel, zabránění poškození systému (ať už úmyslnému nebo neúmyslnému) a vynesení citlivých informací.

V hierarchické síti jsou stanoveny oblasti autorizované a neautorizované, autorizace může být také víceúrovňová.

8.1.3 Management v ISO/OSI

 Jedná se o centralizovaný systém založený na protokolu *CMIP* (Common Management Information Protocol). Management můžeme rozdělit do tří částí:

- *systémový management* (také síťový, network management) – působí v aplikační vrstvě, slouží jako rozhraní k ostatním částem managementu, tj. provádí úkoly související s více vrstvami,
- *vrstvový management* – pracuje v rámci jedné vrstvy,
- *protokolová operace* – také v rámci jedné vrstvy, ale pouze pro jediný případ komunikace.

Jedná se o objektový model (pracuje s objekty a jejich vlastnostmi – atributy, operacemi, vztahem k jiným objektům, používá ISA hierarchii):

- *manažer* (správce) – programové vybavení na stanici síťového managementu (centrální),
- *agent* – programové vybavení na jednotlivých uzlech sítě, posílá zprávy o (mimořádných) událostech manažerovi,
- *MIB* (Management Information Base) – objektová databáze řízených objektů, v tomto modelu binární.

 **MIB** je tedy objektová databáze řízených objektů. Pro každý objekt je zde

- jméno, třída objektu,
- atributy jako uspořádané dvojice typ–hodnota,
- chování – reakce předpokládaná a skutečná na řídicí operace
- hlášení – o událostech souvisejících s objektem včetně informace, co má být agentem sděleno manažerovi
- výčet operací, které je možno s objektem provádět (u třídy)

 **Protokol CMIP** pracuje na aplikační vrstvě uzlů sítě. Při komunikaci používá službu se spojením (tj. pokud nelze navázat zcela spolehlivě spojení, CMIP nemůže komunikovat s agenty), což může být svým způsobem nevýhoda. Pro reprezentaci řídicích informací se používá jazyk ASN.1 (Abstract Syntax Notation One), který je pro tyto účely běžný (setkáme se s ním i u SNMP).

8.2 Management v TCP/IP

8.2.1 SNMP

 V TCP/IP je správa prováděna pomocí protokolu *SNMP* (Simple Network Management Protocol). Ve skutečnosti může běžet i nad jiným protokolem než IP, ale téměř výhradně se setkáme s použitím nad IP. SNMP (resp. jeho nástavby) je v současné době nepoužívanějším protokolem pro správu sítě a je široce podporován prakticky v každém zařízení, které je možné připojit k síti (také tiskárny, nejrůznější čidla a analyzátoři, přístupové body, atd.).

 Administrativní doménu lze rozdělit na *komunity*, každá je identifikována dvěma řetězci *community string* (read community string pro čtení, write community string pro povolení zápisu). Pro zjišťování informací od řízené entity je nutné znát příslušný read community string, pro změnu dat na řízené entitě je třeba znát write community string.

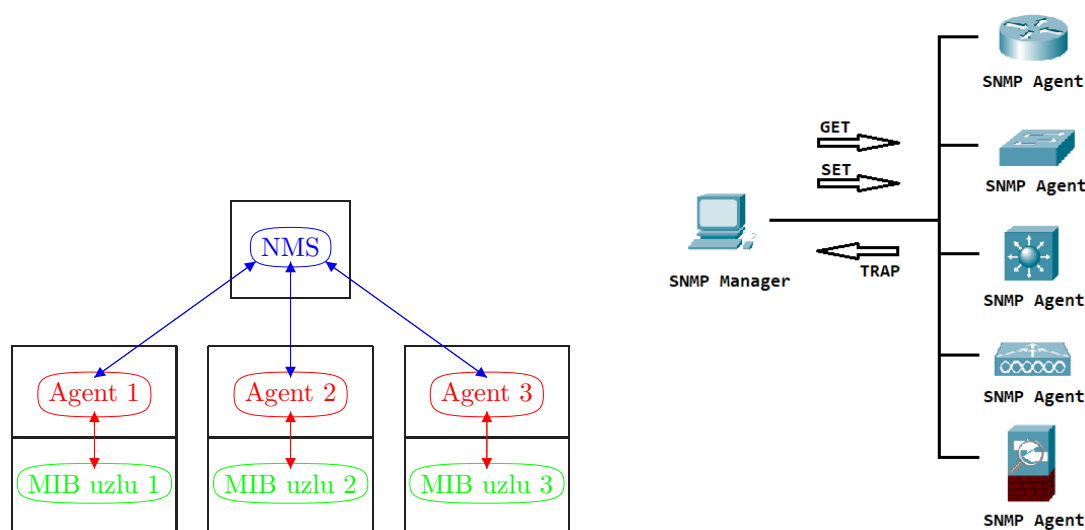
V první verzi existovalo pouze zabezpečení ve formě nutnosti znát tento řetězec (pro danou operaci čtení/zápis), byla to obdoba hesla. Ovšem po síti se přenášel jen jako čistý text.

 Současná verze je SNMPv3, ale přesto je momentálně ještě hodně používaná verze SNMPv2, přesněji její varianta SNMPv2c. Rozdíly:

- SNMPv1 (RFC 1157) autentizuje pouze textovým řetězcem community string,
- SNMPv2 (RFC 1901 až 1908) umožňuje při autentizaci používat digitální podpisy (MD5, SHA),
- SNMPv3 (RFC 2273 till 3415) taky umí šifrovat komunikaci (DES, AES).

SNMPv2 přinesl oproti předchozí verzi například podporu komunikace mezi různými správci (ve verzi SNMPv1 se s více správci ani moc nepočítalo) a vylepšení zabezpečení komunikace. I nadále se sice používají community strings, ale byly přidány nové mechanismy jednoznačné identifikace entit.

 Hlavní vlastností SNMP je jednoduchost. Podobně jako CMIP, i SNMP používá koncept *agent-správce* (manager), ale funkce těchto modulů jsou jinak rozděleny. Každý uzel sítě (spravované zařízení – *managed device*) může obsahovat jakékoliv množství agentů specializovaných na určitou činnost. V síti musí existovat alespoň jeden správce, může jich být více.



Obrázek 8.1: Komunikace v SNMP

 Komunikace mezi agentem a správcem probíhá především přes UDP. Tento protokol však nepodporuje potvrzení doručení (ale na druhou stranu je rychlý a zasílání obvykle funguje), proto SNMPv2 a vyšší obsahuje vlastní mechanismus kontroly doručení.

SNMP podporuje dva typy komunikace:

- *Dotaz–odpověď* – aktivita je na straně správce, správce odesílá dotazy agentům a přijímá jejich reakce. Podle SNMPv2 správce posílá dotazy agentovi na port 161, agent odpovídá z portu 161 na dynamický port správce (tzn. číslo portu se mění). Podle SNMPv3 agent naslouchá na jiném portu – číslo 10 161.
- *Trap* – aktivita je na straně agentů, agenty odesílají trapy (oznámení) správci například při výskytu definované události, překročení zadané hodnoty, změně topologie sítě (například připojení nového uzlu) nebo v pravidelných intervalech. Agent podle SNMPv2 posílá správci trapy ze svého dynamického portu (s různými čísly) na port 162. Podle SNMPv3 správce naslouchá na portu 10 162.

Jak vidíme, mezi verzemi 2 a 3 jsou odlišnosti dokonce i v číslech portů.

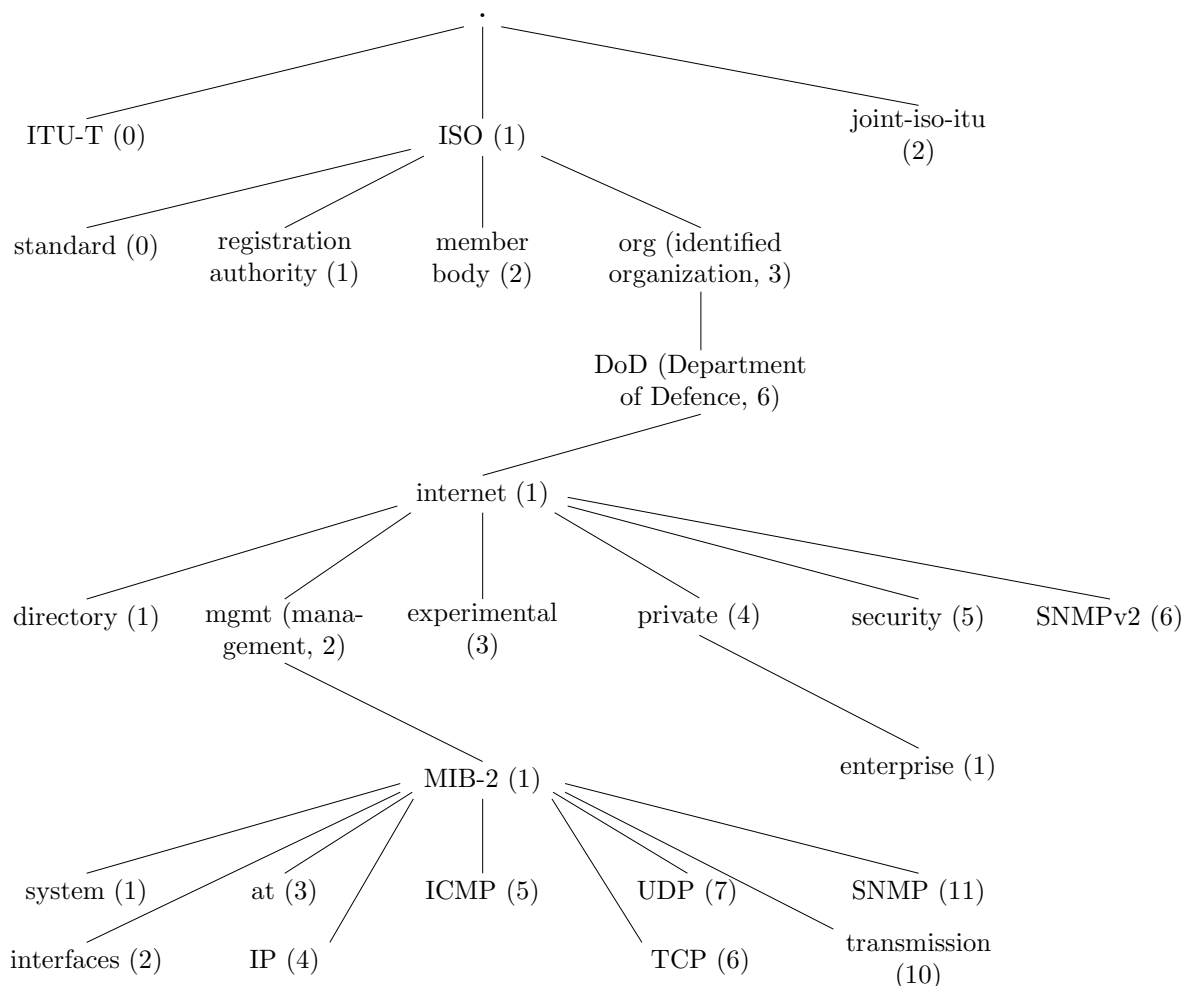
Protokol SNMP je asynchronní, transakčně orientovaný, typu klient/server (komunikace je většinou typu dotaz–odpověď).

8.2.2 MIB-II

 Také se používá *databáze MIB*, ale má obecně jinou strukturu, data jsou uložena v textové formě (obslužný modul MIB je naprogramován v jazyce ASN.1, který je pro tyto účely v TCP/IP typický). Je sice objektová (v tom smyslu, že pracujeme s objekty), ale na rozdíl od OSI MIB je řízená atributy (nepoužívá plně objektový přístup, přistupujeme přes vlastnosti objektů).

Ve skutečnosti se v současných zařízeních setkáme vždy s podporou MIB-II, tedy druhé verze MIB. V následujícím textu budeme pro stručnost používat název MIB, ale vždy půjde o MIB-II.

 MIB je tvořena stromem s jediným kořenem. Kořen obsahuje „prázdnou hodnotu“, jeho účelem je pouze spojit z něho vycházející větve. Všechny uzly kromě kořene mají přidělen textový název a číselnou hodnotu – *OID* (Object ID), OID jednoznačně odlišuje uzly se stejným „rodičem“ – viz obrázek 8.2. Textový název je určen spíše pro „lidskou orientaci“, nemá v databázi žádný jiný význam.



Obrázek 8.2: SNMP MIB (zkrácená)

Uzly (objekty) v MIB, které jsou přímými potomky kořene, jsou nazvány podle organizací, jejichž protokoly využívají podstromy těchto uzlů (například ISO nebo ITU). Některé části (větve) MIB jsou standardizovány organizací IETF, další vydávají výrobci síťových zařízení. Fyzicky je tedy MIB uložena obvykle ve více než jednom (textovém) souboru, aby zpracování zde uložených dat bylo dostatečně pružné a aby byla snadněji rozšiřitelná.



Adresa objektu v MIB je sekvence

- *názvů uzlů* na cestě od kořene stromu – pro lidi, objekty oddělujeme lomítkem nebo tečkou, NEBO
- *čísel OID uzlů* na cestě od kořene stromu – pro kohokoliv/cokoliv jiného, objekty oddělujeme tečkou (včetně prázdného názvu kořene, tedy celý řetězec vlastně začíná tečkou).



Příklad

Například podle obrázku 8.2 má uzel *enterprise* určený pro firemní uzly (zařízení různých výrobců) tuto textovou a OID adresu:

- `.iso.org.dod.internet.private.enterprise`
- `.1.3.6.1.4.1`



Každé síťové zařízení, protokol či hodnota (resp. přiřazený objekt), které lze přes MIB spravovat, má svou unikátní OID adresu, která je stejná v jakémkoliv MIB, tedy všechny OID adresy musí být globálně jedinečné.



Listy stromu MIB jsou *skalární objekty*, jedná se o *proměnné* obsahující konkrétní hodnotu využívanou pro účely řízení sítě (například počet paketů procházejících routerem). Proměnné mohou být uspořádány i do tabulky (*tabulkové objekty*). SNMP poskytuje omezenou možnost pohybovat se v tabulce – po sloupcích. Typy proměnných v MIB:

1. číselné

- běžný *Integer*
- *Counter* – pro čítače, při dosáhnutí maximální hodnoty se vynulují, slouží především pro zjištění rychlosti sledování změn
- *Gauge* (míra) – pokud sledovaná veličina přesáhne danou hranici, hodnota Gauge zůstává na maximální hodnotě a „nepřeteče“
- *Time Ticks* – sleduje čas (v setinách sekundy) od zadané události, například od spuštění zařízení

2. *IP Address* – zde bývá uložena IP adresa

3. *Octet String* – posloupnost oktetů, používá se pro ukládání znakových řetězců

4. *Object Identifier* – OID řetězec některého objektu

5. *tabulka* hodnot výše uvedených typů.

K proměnným se přistupuje poněkud zvláštním způsobem. Pokud se jedná o jednoduchou proměnnou (takovou, která není tabulkou), syntaxe je

`.cesta.název_proměnné.0`

Pokud se jedná o tabulku, pak místo čísla 0 na konci použijeme index v tabulce.



Příklad

K proměnné *sysUpTime* ve větvi *system* vede cesta

`.iso.org.dod.internet.mgmt.mib-2.system.sysUpTime.0` nebo číselně
`.1.3.6.1.2.1.1.3.0`

K hodnotám v tabulkách přistupujeme tak, že místo čísla 0 na konci dosadíme index v tabulce (od 1), například k hodnotám v tabulce ukazujícím stav portů zařízení (živý/mrtvý, v Ethernetu to obvykle znamená je/není něco připojeno) přistupujeme takto:

```
.iso.org.dod.internet.mgmt.mib-2.interfaces.ifTable.ifEntry.ifOperStatus.1
```

```
.iso.org.dod.internet.mgmt.mib-2.interfaces.ifTable.ifEntry.ifOperStatus.2
```

```
.iso.org.dod.internet.mgmt.mib-2.interfaces.ifTable.ifEntry.ifOperStatus.3
```

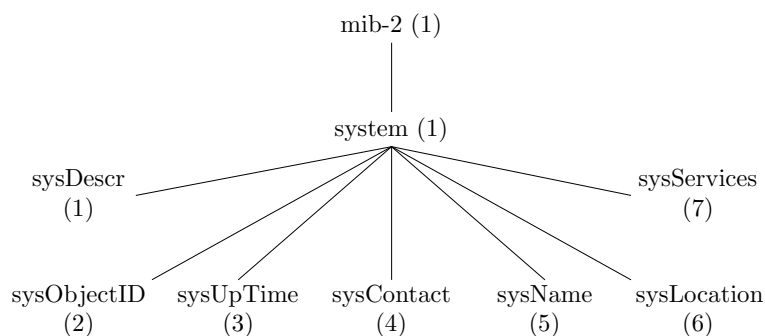
atd., hodnota může být buď `up(1)` nebo `down(2)`. V *interfaces* je jednoduchá číselná proměnná *ifNumber*, ve které máme uložen počet rozhraní v systému (number of interfaces). Je zřejmé, že na pracovní stanici bude `ifNumber.0 = 2` (jedno fyzické rozhraní a jedno loopback), případně nějaký ten výtvar virtualizačního softwaru, na směrovači bude rozhraní více.



Uvádíme zde sice „absolutní“ adresy, ale je možné adresovat také relativně uvnitř daného MIB souboru (skupiny), například ve skupině *interfaces*, *ip* nebo *system*.

 Struktura MIB je celkem logická. Jak bylo výše uvedeno, ve větvi *mib-2* najdeme proměnné, které se netýkají přímo konkrétního výrobce (tj. obecné informace), kdežto ve větvi *private.enterprise* jsou proměnné týkající se výrobků od konkrétních výrobců (záleží, co konkrétně v příslušném síťovém zařízení najdeme). Zařízení, která nejsou pro MIB standardizovaná, jsou obvykle umístěna ve větvi *experimental*.

Obecné informace o zařízení, s jehož MIB pracujeme, najdeme především v podstromu uzlu *system*. Proměnné v něm obsažené vidíme na obrázku 8.3. Je zde popis zařízení (*sysDescr*), adresa OID vedoucí do větve *enterprises* s podrobnějšími informacemi (*sysObjectID*), doba od zapnutí zařízení v setinách sekundy (*sysUpTime*), kontaktní údaje ke správci zařízení (*sysContact*), název zařízení přiřazený adminem (*sysName*), info o fyzickém umístění zařízení (*sysLocation*) a číslo určující vrstvy v ISO/OSI, na kterých zařízení pracuje (*sysServices*).



Obrázek 8.3: Podstrom uzlu *system*



Příklad

Vraťme se k hodnotě *sysServices* určující vrstvy v ISO/OSI, na kterých dané zařízení (to, na kterém je uložena databáze daného agenta) pracuje. Jednotlivé bity této hodnoty jsou nastaveny podle podpory vrstev (bity zprava doleva od nejméně významného bitu).

Například `sysServices.0 = 10` znamená, že zařízení pracuje na L2 a L4 ($2^{2-1} + 2^{4-1}$, tedy je nastaven druhý a čtvrtý bit zprava) a znamená to, že jde zřejmě o switch (L2) a obsahuje funkcionality L4 pro účely správy (transportní vrstva).



 Dalším užitečným uzlem v obecné části MIB je uzel *interfaces*. Zde zjistíme informace o rozhraních zařízení, informace o jednotlivých zařízeních jsou v tabulce *ifTable* (do níž jsme už trochu nahlédli).

Tato tabulka má mnoho sloupců, například *ifSpeed* (nominální rychlost přenosu na rozhraní), *ifOperStatus* (operační stav rozhraní), počet oktetů, přijatých/odeslaných/zahozených unicast paketů, non-unicast paketů (*ifOctets*, *ifInUcastPkts*, ...), atd.

Skutečnou rychlost rozhraní zjistíme pouze tak, že dvakrát za sebou zjistíme počet přijatých oktetů, odečteme a vydělíme délkou intervalu, který uplynul mezi těmito dvěma načteními hodnot.

8.2.3 Používání protokolu SNMP

 Protokol SNMP definuje příkazy (pro agenty a správce):

- **get-request** – správce žádá informace z MIB; uvede název proměnné, jejíž hodnotu požaduje, získá hodnotu a název této proměnné,
- **get-next-request** – účel je stejný (žádost o informace z MIB), správce také uvede název proměnné, ale získá hodnotu proměnné a dále název následující proměnné z databáze, která je potomkem téhož uzlu (tj. pokud žádal o proměnnou ...2.1.1.3, dostane informaci o existenci proměnné ...2.1.1.4, na kterou se může dotazovat následně),
- **set-request** – správce ukládá informace do MIB, je nutný autorizovaný přístup správce,
- **trap** – agent předává správci nevyžádanou informaci o mimořádné události,
- **get-response** – odpověď agenta na předchozí **get-request** od správce, kromě hodnot z MIB obsahuje také původní dotaz, protože SNMP neumožňuje správci párovat dotaz/odpověď (nestavové chování),
- **get-bulk** – podobně jako **get-request**, ale je možné žádat i několik řádků tabulky (od SNMPv2),
- **inform** – komunikace mezi dvěma správci (od SNMPv2).

Příkaz **get-next-request** použitý ve vhodné smyčce umožňuje získat postupně hodnoty stejných atributů přes všechny objekty stejného typu (sloupcová operace).

 Konkrétně můžeme tyto příkazy (ať už v textové nebo grafické podobě) zadávat v některém z nástrojů vytvářejících rozhraní k SNMP. Lze použít například tyto nástroje:

- *net-snmp* je open-source nástroj pracující v textovém režimu, a to pod Linuxem, dalšími UNIXovými systémy a také pod Windows. Je to ve skutečnosti balík programů, který umožňuje využívat plně SNMP v kterékoliv jeho verzi, včetně příjmu trapů.
- *Kiwi* je syslog server s funkcionalitou SNMP, s grafickým rozhraním, existuje free varianta.
- *MIB Browser* je volně šiřitelná aplikace s grafickým rozhraním umožňující pracovat s MIB.
- *GNetWatch* je open-source aplikace s grafickým rozhraním pro real-time monitoring sítě.



Další informace:

- <https://www.samuraj-cz.com/clanek/snmp-simple-network-management-protocol/>
- <https://www.samuraj-cz.com/clanek/zarizeni-v-siti-pod-kontrolou/>
- <https://www.cisco.com/c/en/us/td/docs/ios-xml/ios/snmp/configuration/x6/snmp-x6-book/nm-snmp-cfg-snmp-support.html>
- <https://www.abclinuxu.cz/clanky/system/monitoring-pomoci-nastroju-z-baliku-net-snmp>
- <http://www.snmplink.org/>
- <http://www.net-snmp.org/>, <https://www.solarwinds.com/free-tools/kiwi-free-syslog-server>,
<https://www.ks-soft.net/hostmon.eng/mibbrowser/index.htm>, <https://gnetwatch.sourceforge.net/>



8.3 Syslog

Syslog je mechanismus logování a souvisejících úloh dostupný na každém zařízení, na kterém běží (téměř) jakýkoliv UNIXový systém včetně Linuxu. Jádrem mechanismu je démon¹ *syslogd*.

Démon *syslogd* čte svůj vstup ze zařízení (socketu) `/dev/log`. Tento vstup filtruje (vybírá to, co ho zajímá) podle konfigurace, která se nachází v souboru `/etc/syslog.conf` (to je tedy jeho konfigurační soubor) a pak podle nastavení ukládá hlášení o takto zjištěných událostech do jednoho nebo více LOG souborů.

 Data, která *syslogd* přijímá, se skládají ze tří částí:

1. *kategorie* určuje typ nebo odesílatele události, existují tyto kategorie:

- *auth* – autentizace uživatelů,
- *authpriv* – informace o autentizaci určená pro administrátora,
- *cron* – zprávy od démona *cron* (plánování procesů),
- *daemon* – zprávy od různých démonů,
- *kern* – zprávy od jádra (kernel),
- *lpr* – zprávy od tiskového subsystému,
- *mail* – zprávy týkající se elektronické pošty,
- *mark* – časová razítka (timestamps), které se pravidelně zapisují do logu,
- *news* – diskusní skupiny,
- *security* – totéž co *auth*,
- *syslog* – vlastní zprávy syslogu (i z jiného uzlu v síti),
- *user* – obvykle zprávy od aplikací v uživatelském režimu,
- *local0–local7* – pro tyto kategorie lze definovat vlastní význam,

2. *priorita* určuje důležitost události:

- *emerg* – systém je nepoužitelný nebo vážně ohrožen (emergency),
- *alert* – je nutný okamžitý zásah,
- *crit* – kritická situace,
- *err* – chyba,
- *warning* – varování,
- *notice* – normální, avšak významná, zpráva,
- *info* – informativní zpráva,
- *debug* – ladící zpráva (debugger),

3. vlastní *text* zprávy.

 Tento typ informací tedy *syslog* získá. Jak bylo výše uvedeno, ve svém konfiguračním souboru má určeno, co s takovým záznamem provést. V tomto souboru jsou záznamy ve formě

<code>kategorie.priorita</code>		<code>cíl</code>	<i>tato a vyšší priority</i>
<code>kategorie.=priorita</code>		<code>cíl</code>	<i>pouze tato priorita</i>
<code>kategorie.!priorita</code>		<code>cíl</code>	<i>všechny priority kromě této a vyšších</i>
<code>kategorie.!=priorita</code>		<code>cíl</code>	<i>všechny priority kromě této</i>

¹Pro ty, kteří zapomněli: démon je v UNIXových systémech něco podobného jako služba ve Windows. Jde o systémový proces, který běží na pozadí. Názvy démonů obvykle končí písmenem „d“.

(priorit může být i více, oddělují se středníkem, totéž platí i o dvojicích `kategorie.priorita`). Pokud je před prioritou jen tečka, nastavení platí pro uvedenou a všechny vyšší priority. Pokud chceme nastavení omezit jen na zadanou prioritu, přidáme před ni „=“. Symbol „!“ znamená negaci, tedy pokud je uveden, daná priorita (a všechny vyšší) nebude zahrnuta. Lze kombinovat: „!=“, nebude zahrnuta pouze uvedená priorita. Kategorii můžeme zadat symbolem *. To znamená, že se konfigurace týká jakékoliv kategorie.

Mezi prioritami a cílem musí být vždy *alespoň jednou stisknutý tabulátor*, mezera nestačí.

Cíl určuje, kam konkrétně má být událost oznámena, logována. Může to být buď konkrétní log soubor (výchozí nebo jakýkoliv jiný), nebo výpis na konzolu či přeposlání na jiný počítač v síti. Pokud chceme, aby byla událost oznámena více cílům (například zobrazena určitému uživateli a zároveň uložena do souboru), oddělíme cíle čárkou. Možnosti:

- soubor – výchozí je `/var/log/messages`, ale můžeme si určit názvy souborů například pro různé kategorie nebo priority,
- `@server.firma.cz` nebo `@IPadresa` – událost bude přeposlána,
- `user1, user2, ...` – událost bude oznámena zadanému uživateli (to může být `root`, `admin`, `operator`, apod., podle toho, jaké máme vytvořené uživatele), ale jen tehdy, když je uživatel právě přihlášen,
- `*` – událost bude oznámena všem přihlášeným uživatelům,
- `@loghost` – můžeme použít, pokud v souboru `/etc/hosts` je definován cíl `loghost`.

Lze použít také přesměrování do pojmenované roury, ze které pak může číst jakýkoliv proces, který určíme (a průběžně zpracovávat oznámení o událostech), stačí uvést název souboru a před něj napsat symbol roury:

```
kern.=err          | /var/log/jadro
```



Postup

Takto nějak mohou vypadat záznamy v souboru `/etc/syslog.conf`:

```
mail.*              /var/log/maillog
```

Všechny události z kategorie *mail* jsou uloženy do souboru `/var/log/maillog`

```
security.*;security.!=debug /var/log/secure
```

Všechny události z kategorie *security* kromě události s prioritou *debug* jsou uloženy do souboru `/var/log/secure`

```
cron.*              /var/log/cron
```

Všechny události z kategorie *cron* jsou uloženy do souboru `/var/log/cron` (to znamená události související s plánovaným spouštěním procesů)

```
kern.debug;auth.notice /dev/console
```

Události týkající se ladění jádra a běžné a přesto významné události autentizace jsou vypsány na systémové konzoli

```
authpriv.*          /var/log/secure
```

Všechny „citlivější“ události autentizace jsou uloženy do souboru `/var/log/secure`

```
lpr.info             /var/log/lpd-info
```

Informační zprávy a všechny s vyšší prioritou o událostech z tiskového subsystému se uloží do `/var/log/lpd-info`

*.err admin,/var/log/errors

Všechny chybové a vážnější zprávy jsou okamžitě oznámeny uživateli *admin* (pokud je přihlášen) a zároveň uloženy do souboru */var/log/errors*

*.=crit /var/log/messages,@loghost,root

Kritické události jsou uloženy do souboru */var/log/messages*, poslány na adresu definovanou pod aliasem *loghost* a zároveň je okamžitě informován *root*, pokud je přihlášen

*.emerg *,/var/log/emergency

O mimořádně nebezpečných událostech jsou informováni všichni přihlášení uživatelé a zároveň je přidán záznam do souboru */var/log/emergency*



Pokud chceme, aby *syslog* přijímal zprávy z jiných systémů, musíme spustit démona *syslogd* s parametrem *-d* (platí v Linuxu):

```
/usr/sbin/syslogd -m 0 -r
```

Přepínač *-m* určuje délku intervalu, v jakém se *syslogd* ozývá, tj. do logu zapisuje, že „žije“. Nastavením na 0 toto chování zrušíme. Parametr *-r* znamená „remote“, *syslogd* pak naslouchá na portu 514 a přijímá všechny UDP pakety.

Na FreeBSD je nutné použít jinou syntaxi:

```
/usr/sbin/syslogd
```

(bez parametrů, protože naslouchání na portu 514 je zde výchozí chování). Na FreeBSD je také možné určit uzly v síti, jejichž UDP pakety budou přijímány. Odlišnou syntaxi (i od této) mají systémy OpenBSD, Solaris a jiné, je tedy třeba vždy prostudovat manuálovou stránku:

```
man syslogd
```

Také je možné, že budeme muset povolit naslouchání služby *syslogd* na UDP portu. To se provede v */etc/services* řádkem *syslog 512/UDP*, ale je pravděpodobné, že tam takový záznam už je. Na zařízeních, ze kterých jsou UDP pakety přijímány, je pak nastaveno výše popsáním způsobem zasílání na tento „sběrný“ počítač.

Pokud *syslogd* právě běží a my jsme provedli změnu jeho konfigurace (v souboru */etc/syslog.conf*), musíme *syslogd* restartovat, aby zaregistroval změny ve své konfiguraci.

 Zatím jsme si ukázali, jaké údaje dostává *syslogd* na svůj vstup, podle čeho je filtruje a kam je ukládá. Zbývá nám podívat se, v jakém formátu. Na výstupu je záznam obsahující

- čas, kdy k události došlo,
- kategorii (kernel, mail apod.),
- text zprávy.

Součástí není priorita, protože ta už byla uplatněna při filtrování. Může zde být například záznam:

```
Feb 21 10:00:28 IOL kernel: device eth0 left promiscuous mode
```

Když syslog nastavujeme, hodí se možnost vyzkoušení jeho funkčnosti. K tomu může sloužit nástroj *logger*. Například událost kategorie *daemon* a priority *info* se zadanou zprávou vygenerujeme takto:

```
logger -p daemon.info "testujeme syslog"
```

Ruční správa logů může být celkem náročná. Je třeba hlídat obsah souborů a navíc sledovat jejich délku (včas umazat starší záznamy), stroj naslouchající na UDP portu by měl být dostatečně chráněn (je zde vysoké riziko DoS útoků, tedy firewall je rozhodně na místě). S některými úlohami mohou pomoci

přídavné nástroje. Například rotaci logů (včetně označování či odstraňování starších záznamů) zvládne *logrotate*.

Existují také propracovanější verze *syslogu*, například *syslog-ng* (umí komunikovat i přes TCP a zvládne také regulární výrazy, nejen hvězdičku), *nsyslogd* (komunikuje přes TCP/SSL), *Secure Syslog*, *Modular Syslog* a další. Volně šiřitelný program *NTSyslog* (vlastně služba, kterou můžeme instalovat), umožňuje používat *syslog* také ve Windows (logy z Windows lze posílat na vzdálený systém a tam například vyhodnocovat). Program *swatch* je užitečný, když naopak potřebujeme být o určité události informováni pokud možno okamžitě – detekuje námi definované situace a stanoveným způsobem informuje, a protože je psán v Perlu, je to velmi pružný nástroj využívající regulární výrazy.



Další informace:

- <https://www.root.cz/clanky/linux-jako-internetova-gateway-9/>
- https://www.softpanorama.org/Logs/Syslog/syslog_configuration_examples.shtml
- <https://www.root.cz/clanky/synchronizace-casu/>
- <https://ntsyslog.sourceforge.net>, <https://www.solarwinds.com/kiwi-syslog-server>



8.4 Monitorování sítě

8.4.1 Protokol RMON

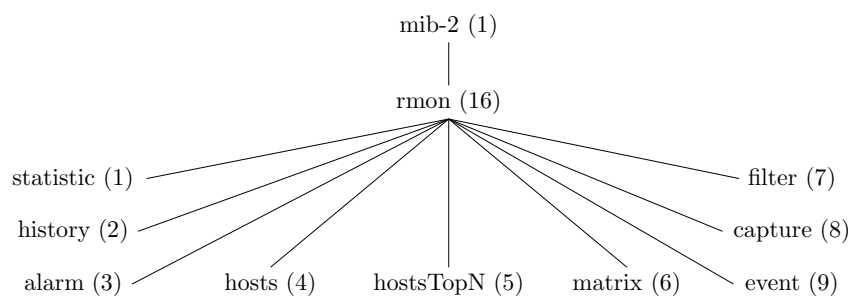
 RMON (Remote Monitoring) je protokol používaný při monitorování sítě. Je úzce spjat s TCP/IP a SNMP, své zprávy posílá přes SNMP. Informace jsou uloženy v MIB ve větvi *rmon* (OID 1.3.6.1.2.1.16). Na rozdíl od SNMP agentů dokáže kromě současného stavu sledovat i historii (vývoj) dané veličiny a na základě vyvozených důsledků reagovat.

 Z hlediska protokolu RMON rozlišujeme dva druhy zařízení v síti:

- RMON-compliant console manager (*správce*) – správce celé LAN
- RMON probe (*sonda*) – zasílá informace správci, konfiguruje se například na switchi

 RMON pracuje s tzv. *skupinami*. Každá skupina v sobě zahrnuje určitý typ informace, která je sledována. Například skupiny pro Ethernet jsou

- *Statistics* – statistika (okamžitý stav) sledovaných zařízení (množství odeslaných paketů, broadcast a multicast paketů, oktetů, chyb v CRC součtech, kolizí apod., také čítače, například pro počty paketů o velikosti z daného intervalu),
- *History* – totéž jako statistika, ale evidováno v historii,
- *Alarms* – pro některé sledované veličiny jsou stanoveny prahové hodnoty, a když je tato prahová hodnota překročena, generuje se příslušná událost,
- *Hosts* – vedou se statistiky typické pro jednotlivé uzly v síti (segmentu sítě) – adresa, odeslané/přijaté pakety, chyby v CRC součtech a zahozené pakety či rámce,
- *HostsTopN* – podobně jako předchozí, uzly jsou seřazeny v tabulkách podle konkrétní sledované veličiny,

Obrázek 8.4: Větev *rmon* v MIB-II

- *Matrix* – sledují se veličiny související s konverzací mezi dvěma uzly v síti, pro každou komunikující dvojici uzlů,
- *Packet Capture* – definují se podmínky pro zachytávání paketů (například velikost vyrovnávací paměti pro zachycené
- *Filters* – umožňují definovat filtry pro zachycení některých paketů nebo generování určité události při průchodu určených paketů, pakety, počet zachycených paketů, kdy vyvolat alarm apod.),
- *Events* – generování a ověřování událostí (eviduje se typ události, její popis a časový údaj posledního vygenerování této události daným zařízením).

Zařízení připojená v síti obvykle podporují většinu těchto skupin, v ideálním případě všechny. Tedy sledujeme ty proměnné v MIB, které nás zajímají.

Hlavním přínosem RMON je přenos velké části zpracování dat od správců na agenty (sondy), čímž se také snižuje množství přenášených dat v síti.

8.4.2 Snort

 Snort je jednoduchý volně šiřitelný systém (open-source licence, ale pro aktualizace je nutné se registrovat, „okamžité“ aktualizace jsou navíc placené), který můžeme zařadit mezi *NIDS* (Network Intrusion Detection System).

 Snort pracuje v jednom ze tří možných režimů:

- *sniffer* („prohlížeč“ provozu) – data o paketech vypisuje na obrazovku nebo někam posílá,
- *logování provozu* – data ukládá na disk a/nebo do databáze,
- *plný NIDS* včetně používání pravidel – analýza hlaviček paketů.

Výstupy dokáže Snort buď ukládat do LOG souboru, nebo zasílat *syslogu*, posílat jako SNMP trap, ukládat do databáze a nebo dokonce podle nich nastavovat konfiguraci některých síťových prvků.

 Snort funguje na principu *detekce signatur* v kombinaci s *pravidly*. Mnoho pravidel je vytvořeno už při instalaci (je jich opravdu hodně, někdy to chce trochu je promazat), a také je možné přidat pravidla vlastní podle konkrétní situace v síti. Využívání pravidel dává systému obecně větší sílu (například kombinace „mírně podezřelých“ akcí je *velmi podezřelá* a narozdíl od systémů založených pouze na signaturách systém generuje méně falešných poplachů.

 Obecný tvar pravidel je

```
action protocol sourceIP sourcePort direction destIP destPort (options)
```

Action (akce) může být například *pass* (předání), *log* (zaznamenání), *alert* (výstraha), *drop* (zahodit paket), atd. Následují *vlastnosti paketu*, podle nich pozná Snort, že má pravidlo použít (protokol, třeba TCP, dále zdrojová a cílová adresa a port, a také směr přenosu paketu zachycený „šipkou“). Poslední částí pravidla jsou *options* – volby, které ve skutečnosti popisují, co se má stát, když Snort zachytí paket odpovídající údajům v pravidlu a také co dalšího má být testováno (například pole TTL paketu nebo ICMP informace). Například:

```
alert tcp any 3389 -> 81.28.192.0/24 3389 (msg: "TCP paket na Net5");
```

To znamená, že má být generována výstraha, pokud je nalezen TCP paket z jakékoliv adresy na portu 3389 mířící do sítě s danou IP adresou (používají se CIDR adresy s prefixem), a to na tomtéž portu. S výstrahou se má vygenerovat zpráva v zadaném znění (zpráva je nahlášena a uložena do logu). Ve skutečnosti bývají pravidla poněkud sofistikovanější a mohou obsahovat také například popis porovnání se signaturou a další volby.

V uvedeném příkladu jsou konkrétní adresy, ale lze využít také proměnné, ve kterých má Snort předdefinovanou adresu vnitřní sítě a některé další adresy, například

```
alert icmp $EXTERNAL_NET any -> $HOME_NET any (msg:"ICMP Large ICMP Packet"; dsize: >800;)
```

 Snort funguje podobně jako předchozí popsané mechanismy včetně SNMP (správce+agenty). Agenty se nazývají *sondy* a nacházejí se v různých oblastech infrastruktury, podle potřeby (umístění stanovíme podle typu informací, které chceme získávat, záleží na tom, jestli je sonda před či za firewallem, v DMZ, v dané kolizní doméně apod.).

 Existují frameworky (rozhraní), které zjednodušují přístup a analýzu dat Snortu, a také se Snortem spolupracují další podobné projekty. Například *Prelude* je framework obsahující vlastní IDS modul, ale může používat i modul pro Snort. Výstupy Snortu jsou typicky směrovány do některého dohledového systému, případně SIEM (budeme probívat v další kapitole).

Další systém, který zjednodušuje používání Snortu, je *ACID* (Analysis Console for Intrusion Databases). ACID je sada skriptů a je zajímavý už proto, že se s jeho použitím počítá v doporučeních pracovních skupin CERT a CSIRT. ACID potřebuje webový server (třeba Apache), funkční PHP a databázi, do které Snort ukládá záznamy (třeba MySQL). ACID pak provádí analýzy nad daty z této databáze.



Další informace:

- <https://www.snort.org/>, <https://www.scribd.com/document/6777057/Snort-Manual>
- https://i.info.cz/r/kd/Intrusion_Detection_with_SNORT.pdf
- <https://www.cs.vsb.cz/grygarek/SPS/projekty0405/IDS/ids.html>
- <https://www.cs.vsb.cz/grygarek/SPS/projekty0506/Snort.pdf>
- <http://www.inguardians.com/research/docs/snortguis.pdf>
- <https://www.andrew.cmu.edu/user/rdanyliw/snort/snortacid.html>



8.4.3 NetFlow

 NetFlow je protokol vyvinutý společností Cisco. S jeho podporou nicméně setkáme nejen na zařízeních od této společnosti, ale také v zařízeních Juniper a některých dalších výrobců (přesněji: tato zařízení

mohou exportovat informace ve stejném formátu jako NetFlow). Nemusí jít jen o routery a switche, ale existují také jednoduchá přenosná zařízení, která plní pouze roli sledování sítě tímto způsobem.

Existuje více verzí NetFlow s odlišnými vlastnostmi. Pokud potřebujeme podporu IPv6, je nutné používat minimálně NetFlow verze 9. NetFlow verze 10 byl standardizován organizací IETF jako RFC 7011 pod názvem IPFIX (IP Flow Information Export).

 NetFlow pracuje na principu *toků*. Tok (flow) je to, co obvykle považujeme za úplnou síťovou konverzaci. V případě přenosů realizovaných spojovými stavovými protokoly (jako je například TCP) do jednoho toku patří nejen odeslání jednoho paketu, ale komunikace v rámci celého spojení (ale pouze jedním směrem, tedy obousměrná komunikace je technicky ve dvou tocích). Do jediného společného toku také patří například sekvence ICMP paketů vyslaných v rámci jediného příkazu `ping`. Každý tok je jednoznačně popsán těmito údaji:

- zdrojová IP adresa
- cílová IP adresa
- zdrojový port
- cílový port
- protokol

Protokoly, které se ve všech těchto parametrech shodují, jsou řazeny do téhož toku. Tady je důvod, proč jsou různé směry komunikace v různých tocích – opačné směry mají v záhlavích přehozené adresy a porty. Komunikace také může být rozdělena do více toků, pokud trvá příliš dlouho (na směrovačích Cisco je to například 30 minut), nebo je delší dobu neaktivní anebo je zaplněna vyrovnávací paměť zařízení (toky není kam ukládat, proto ty starší budou vymazány).

 Místo zjišťování informací (s protokolem NetFlow) se nazývá *Observation Point* (také *NetFlow Exporter*). Toto zařízení zjišťuje z průchozích paketů potřebné informace a odesílá je flow kolektoru. *Flow kolektor* je hardware nebo software, který přijímá data o tocích a zpracovává. Jeden flow kolektor může přijímat data od více observation pointů (takže opět jde o architekturu agenti–správce). Na flow kolektor se informace o toku dostává až po ukončení tohoto toku.

Ke každému toku se kromě určujících parametrů evidují také další informace – počet paketů v toku (na Observation Pointu se během evidence toku tento údaj při každém identifikovaném paketu zvyšuje o 1), doba trvání (před odesláním Flow kolektoru se srovnají časová razítka začátku a konce toku), celkové množství přenesených dat (s každým paketem se tato hodnota postupně zvyšuje).

 Kromě ukládání dat na Flow kolektor ještě potřebujeme aplikaci, která by zjednodušila analýzu těchto dat. Existuje více takových aplikací, mnohé volně šiřitelné.

- *OSU Flow-Tools* jsou volně šiřitelný balík programů pro textový režim vyvíjený na univerzitě v Ohio. V balíku najdeme nástroje na shromažďování dat na flow kolektoru a zpracování údajů o tocích.
- *NFDump* je volně šiřitelný kolektor (distribuovaný pod BSD licenci). Existuje pro něj také grafická nástavba *NFSen*, která je vlastně webovým rozhraním.

 Díky „koncentrovanějším“ informacím, které NetFlow umožňuje získávat, lze odhalit například DoS útok (všechny pakety v rámci tohoto útoku mají společné posuzované parametry, patří do jednoho toku), ale ne například DDoS útok (je z mnoha zdrojů, obvykle z počítačů některého rozsáhlého botnetu) ani skenování uzlů v síti (různé cílové adresy).

Oproti předchozím technologiím má NetFlow výhodu v tom, že administrátora nezahluje informacemi: do Flow kolektoru se dostane jen jediný záznam o kompletním (jednosměrném) toku, i kdyby to byly stovky paketů.

NetFlow je výborným nástrojem také pro drobné poskytovatele internetového připojení, protože takto mohou jednoduše provádět monitoring a správu sítě, mít přehled o vytíženosti sítě, provádět optimalizaci (včetně nastavení QoS) nebo účtování na základě vytížení sítě, a hlavně mohou splnit prováděcí vyhlášku 357/2012 Sb. o uchovávání, předávání a likvidaci provozních a lokalizačních údajů doplňující Zákon o elektronických komunikacích.



Další informace:

- <https://www.flowmon.com/cs/solutions/use-case/netflow-ipfix> (včetně odkazu na video)
- <https://www.cisco.com/c/en/us/products/ios-nx-os-software/ios-netflow/index.html>
- <https://nfdump.sourceforge.net/>, <https://nfsen.sourceforge.net/>



8.5 VPN

 VPN (Virtual Private Network) je zabezpečené spojení procházející nedůvěryhodným prostředím (obvykle Internetem). Jedná se o vytvoření tunelu, komunikačního kanálu mezi dvěma body (to mohou být konkrétní koncová zařízení nebo třeba routery, pak propojujeme sítě).

 Používá se v těchto případech (tj. dělení podle typu ukončujících bodů):

- *Remote Access*: mobilní zaměstnanec potřebuje na svých cestách zabezpečený přístup do firemní (lokální) sítě, aby mohl přistupovat do firemního informačního systému a dalších zabezpečených zdrojů, nebo se jedná o zaměstnance pracujícího doma (Home Office).
- *Site-to-Site* (nebo také *Network-based*): je třeba propojit vzdálené lokální sítě (například pobočky téže firmy).
- Je třeba komunikovat s obchodním partnerem zabezpečeným komunikačním kanálem, ale zároveň ho nechceme pustit přímo do naší lokální sítě. Může se jednat o *site-to-site* nebo *remote access*, podle skutečné konfigurace VPN, tato možnost je vlastně speciálním případem obou předchozích.

Ve všech těchto případech je třeba vybudovat zabezpečený šifrovaný *tunel*, přes který povede komunikace nedůvěryhodným prostředím – zajišťujeme předně vhodné zapouzdření s případným překladem adres (protože pakety půjdou přes cizí síť s jinými adresními rozsahy a případně jinými protokoly), a dále zajišťujeme šifrování (na to je někdy potřeba použít přídatné řešení). Ve třetím případě navíc musíme použít firewall, který bude propouštět pouze komunikaci povolenou pro tento účel.

VPN tunel do firemní sítě ústí většinou buď přímo na routeru s firewallem na hranici lokální sítě, který je viditelný na Internetu, anebo v demilitarizované zóně (tam se často umísťuje VPN koncentrátor, což je vlastně jakási VPN brána), záleží, kam v lokální síti je třeba přes tunel přistupovat.

 VPN řešení by mělo zajistit následující:

- autentizace – je třeba ověřit totožnost obou komunikujících bodů (např. uživatel s notebookem na pracovní cestě a firewall s podporou VPN v LAN firmy), případně se zajišťuje autentizace přenášejících PDU,

- autorizace – stanovení konkrétních přístupových oprávnění,
- zajištění důvěrnosti dat – přenos je vždy šifrován,
- zajištění integrity dat – detekce pozměnění či poškození paketu po cestě.

 **Tunelovací problémy.** Někdy bývá problém se zprovozněním tunelu – všechno se zdá v pořádku, ale v případě, že má být přenesen větší paket, nic nedorazí a do zdroje nepříjde ani ICMP zpráva se zdůvodněním. Důvodem je právě velikost paketu, která je příliš blízko hodnotám MTU na cestě. Nezapomeňte, že příliš velký paket je cestou buď fragmentován nebo zahozen (když se nedá fragmentovat), a přitom VPN obvykle funguje tak, že se původní paket/rámec zapouzdřuje do obalového a nosného paketu – přidávají se další záhlaví, velikost paketu naroste.

Většinou se na síťových zařízeních počítá s tím, že by měly projít rámce zapouzdřující maximálně 1500 B, a s tím taky počítá odesílající koncové zařízení – obvykle už na transportní vrstvě se nastaví maximální velikost segmentu tak, aby po přidání IP záhlaví tato hodnota nebyla překročena (až delší PDU jsou rozděleny do více segmentů). Jenže odesílatel „neví“ o tom, že jeho dílo nabobtná tunelováním.

Dobře, a proč se tedy nefragmentuje? IPv6 se po cestě fragmentovat nesmí, a u IPv4 se často používá mechanismus *Path MTU Discovery*, který je sice užitečný (zjistí MTU po celé cestě a už před odesláním upraví paket tak, aby prošel na všech prvcích), ale automaticky zakazuje fragmentaci po cestě (nastaví příznak DF v záhlaví IP paketu). O problému se často řádným způsobem nedozvíme, protože o zahození nefragmentovatelného paketu je odesílatel informován zprávou ICMP Destination Unreachable, kód 4, a zároveň na některých routerech jsou paušálně zahazovány všechny ICMP zprávy (ano, takové jsou pak následky).

Takže jaké je řešení? Například můžeme na transportní vrstvě určit menší hodnotu pro maximální délku segmentu. Vezmeme obvyklé minimum MTU (1500) a odečteme délku všech záhlaví, která se během zapouzdřování přidávají. Například u GRE se dá použít hodnota 1220.

 **Jak pracují VPN protokoly.** VPN tunel se vytváří zapouzdřením takto:

- uvnitř je PDU přenášeného protokolu, většinou IP (ale může být také IPX, NetBEUI nebo jiný), nebo v některých případech jenom jeho datová část, přičemž záhlaví se použije v třetím kroku,
- následuje obalový protokol, což bývá právě VPN protokol (IPSec, GRE, PPTP, L2TP nebo jiný), zapouzdřená PDU je zašifrována a je přidáno bezpečnostní záhlaví,
- nosný protokol (obvykle IP) v sobě zapouzdří PDU vytvořenou v předchozím kroku, aby bylo možné přenést paket přes „nechráněný“ Internet či jinou veřejnou síť, zde záleží, na které vrstvě se vlastně zapouzdření provádí.

Obecně pojem tunel nezahrnuje šifrování, základem je „obalení“ PDU, přičemž můžeme zpracovat i takové PDU, které by danou cestou samo o sobě nemohlo být posláno (nepodporovaný protokol, soukromé IP adresy, atd.). Když k tunelování přidáme šifrování, získáme VPN.

Dále bude představeno několik běžných VPN řešení. Jedná se o protokoly pracující na různých vrstvách ISO/OSI, další odlišnosti jsou v náročnosti nasazení, (ne)nutnosti mít externího poskytovatele služby, a také v tom, zda je v řešení zahrnuto i šifrování a další bezpečnostní mechanismy.

Některé VPN protokoly jsou typu *multipoint* (je možné sestavit síť tunelů typu mesh, kde každý bod může přímo nebo zprostředkovaně komunikovat s více dalšími body) nebo typu *point-to-point* (každý tunel musí být nakonfigurován zvlášť, všechny tunely mají právě dva konce).

8.5.1 IPSec VPN – na síťové vrstvě

 Protokol IPSec (Internet Protocol Security) umožňuje vytvářet zabezpečené tunely typu point-to-point. Zajišťuje jak vytvoření tunelu, tak i šifrování. Pracuje na síťové vrstvě, díky tomu je transparentní pro aplikační protokoly. Používá se jak pro řešení site-to-site, tak i pro remote access.

IPSec se také často kombinuje s protokoly GRE nebo L2TP pracujícími na vrstvě L2. Jedním z důvodů těchto kombinací je problém IPSec při průchodu přes NAT, který modifikuje IP záhlaví. I pro tento případ však existuje řešení, lze využít mechanismus *NAT-T* (NAT Traversal), který zapouzdří paket do UDP a tím znemožní pozměnění IP záhlaví. Označuje se IPSec over UDP nebo IPSec over NAT-T.

 Šifrování a zapouzdření v IPSec se provádí jedním z těchto dvou způsobů:

- *tunelovací režim* – celý původní IP paket se zašifruje a připojí se záhlaví IPSec (a pak samozřejmě záhlaví nosného protokolu s novou IP adresou),
- *transportní (přenosový) režim* – šifruje se datová část zapouzdřovaného paketu bez záhlaví, pak se přidá IPSec záhlaví, před ním se předstune původní záhlaví zapouzdřených dat.

Tunelovací režim je bezpečnější (celé zapouzdřené IP záhlaví je skryto, šifrováno), transportní režim je pružnější (můžeme používat prvky z původního IP záhlaví, například pro QoS) a pakety jsou kratší (tedy menší dopad na propustnost sítě). Srovnání vidíme na obrázku 8.5.

Původní IP paket:

IP záhlaví1	Tělo paketu
-------------	-------------

- *tunelovací režim:*

IP záhlaví2	IPSec záhlaví	Šifrováno:
		IP záhlaví1 Tělo paketu

- *transportní (přenosový) režim:*

IP záhlaví1	IPSec záhlaví	Šifrováno:
		Tělo paketu

Obrázek 8.5: PDU před a po zašifrování do IPSec tunelu

Tunelovací režim je používán především při spojeních *site-to-site*, zatímco transportní režim je typičtější pro připojení vzdáleného počítače s firemní sítí (*remote acces*).

 Technicky vzato, existují dva typy *IPSec záhlaví*: AH (Authentication Header) a ESP (Encapsulated Security Payload). Zatímco AH je jednodušší, zajišťuje pouze autentizaci a integritu dat, ale ne šifrování, záhlaví ESP je bezpečnější (zajišťuje autentizaci a šifrování). Lze použít obojí najednou, tedy AH/ESP.

Při vytváření jakéhokoli (tedy i IPSec) tunelu je třeba nejdřív dojednat *parametry bezpečného přidružení* (SA – Security Association). IPSec pro tento účel používá *protokol IKE* (Internet Key Exchange), nyní ve verzi 2.

Poznámka:

V protokolu IPv6 je již IPSec nativně implementován (tak to bylo už v původním návrhu IPv6, k verzi IPv4 byl IPSec dodatečně přidán jako „pomocný“ protokol zajišťující šifrování). Proto přechod na IPv6 má jeden bonus – snadnější vytváření IPSec tunelů.





Úkol

Záhlaví AH a ESP mají odlišnou strukturu. Podívejte se na <http://www.unixwiz.net/techtips/iguide-ipsec.html> a srovnajte jejich strukturu a jak se používají (včetně NAT).



Jak se IPSec konfiguruje. IPSec se konfiguruje podobným způsobem jako například firewall – definujeme pravidla a další parametry (v tzv. tabulce politik). Určujeme například, že pakety příchozí z konkrétní adresy (vzdálené pobočky) se mají zpracovat mechanismem IPSec (dešifrovat apod.), pakety odchozí na tutéž adresu naopak zašifrovat, přidat IPSec záhlaví apod., pakety směřující do vnitřní sítě se mají nechat tak jak jsou, atd. Také je obvykle potřeba upravit směrovací tabulku.

V Linuxu záleží na konkrétní distribuci. V distribucích odvozených z Debianu bývá k dispozici (nebo se doinstaluje) balíček *Freeswan*, kde se konfigurace provádí v souboru `/etc/ipsec.conf` a příkazem `ipsec`, v jiných distribucích máme balíček *IPSec Tools* s příkazy `setkey` a `racoon`, konfigurační soubory jsou `/etc/racoon/racoon.conf` a `/etc/racoon/setkey.conf`.

Ve Windows a na síťových zařízeních konfigurovaných přes webové rozhraní se konfigurace provádí většinou „myší“, například ve Windows Server máme k dispozici konzolu v *Start – Programs – Administrative Tools – Local Security Settings*. Kromě toho je ve Windows Server k dispozici program `ipsecpol` pro definování IPSec politik (zásad).

Na klientovi je připojení vcelku jednoduché. VPN je předem nakonfigurována na serveru, na klientovi v podstatě konfiguruje nové připojení k síti: *Vytvořit nové připojení*, zvolíme připojení k firemní síti, VPN, atd. V průvodci potřebujeme IP adresu VPN serveru, se kterým budeme komunikovat. Ve vlastnostech připojení dále nakonfiguruje protokol IPSec, včetně bezpečnostních nastavení.



Další informace:

- Podrobnosti především o Linuxu najdeme v odkazech na konci sekce o VPN, především v odkazu <http://www.ipsec-howto.org/ipsec-howto.pdf>
- Konkrétní postup konfigurace na Linuxu je v <https://www.root.cz/clanky/tuneluji-tunelujes-tunelujeme-ipsec/>, je to jeden z dílů seriálu o tunelování.
- <http://www.unixwiz.net/techtips/iguide-ipsec.html>
- https://www.cisco.com/c/en/us/td/docs/net_mgmt/vpn_solutions_center/2-0/ip_security/provisioning/guide/IPsecPG1.html
- https://www.juniper.net/documentation/en_US/junos/topics/topic-map/security-ipsec-vpn-overview.html



8.5.2 GRE tunely



Protokol *GRE* (Generic Routing Encapsulation) pracuje na síťové vrstvě a vytváří point-to-point tunely, typicky site-to-site. Je standardizován IETF jako RFC 2784.

Obvyklý postup je takový, že původní paket je opatřen GRE záhlavím (velice jednoduchým, pár polí) a následně je přidáno záhlaví nosného protokolu, obvykle IP (v něm je v poli Protocol/NextHeader číslo 47 určující protokol GRE). Všimněte si, že v postupu není ani slovo o šifrování, to totiž neumí.

Jeho výhodou je univerzálnost, dokáže zapouzdřit i jiné typy PDU síťové vrstvy ISO/OSI než jen IP. Dokonce může zapouzdřit i protokoly jiných vrstev, včetně rámců z vrstvy L2. Dokáže přes tunel

transportovat i multicast a broadcast vysílání. Další výhodou je, že tento protokol je podporován prakticky všude (v Linuxu, ve Windows, na síťových zařízeních různých výrobců).

Na druhou stranu, velkou nevýhodou GRE je, že neprovádí šifrování, proto se ve skutečnosti nejedná o „pravý“ VPN tunel. V praxi je GRE zapouzdřován do tunelů VPN (například v kombinaci s IPSec), aby byla konverzace zároveň šifrována.

Konfigurace v Linuxu je jednodušší než IPSec, vystačíme si s vestavěnými nástroji (příkaz `ip tunnel` a další varianty příkazu `ip`).

V příloze na straně 271 je ukázka nastavení jednoduchého GRE tunelu mezi dvěma vzdálenými sítěmi (mezi routery s nainstalovaným Linuxem).



Další informace:

- <https://www.ietf.org/rfc/rfc2784.txt>
- <https://www.root.cz/clanky/tuneluji-tunelujes-tunelujeme-jak-a-k-cemu/>
- <https://www.cloudflare.com/learning/network-layer/what-is-gre-tunneling/>
- https://www.cisco.com/c/en/us/td/docs/iosxr/ncs5500/interfaces/61x/b-ncs5500-interfaces-configuration-guide-61x/b-ncs5500-interfaces-configuration-guide-61x_chapter_01101.pdf
- <https://networklessons.com/cisco/ccie-routing-switching/how-to-configure-gre-tunnel-on-cisco-ios-router>
- https://www.juniper.net/documentation/en_US/junos/topics/topic-map/switches-interface-gre.html



Úkol

V příloze si projděte postup konfigurace GRE tunelu.



8.5.3 L2TP tunely



Protokol L2TP je potomkem starších protokolů PPTP (Microsoft) a L2F (Cisco). Je standardizován organizací IETF – verze L2TPv3 z roku 2005 je standard RFC 3931.

Jak název napovídá, tento protokol pracuje na vrstvě L2 (Layer 2 Tunneling Protocol). VPN protokoly linkové vrstvy jsou zajímavé tím, že obvykle dělají „smyčku“ na vyšší vrstvu – zapouzdřují provoz z nadřízené vrstvy, ale samy se zapouzdřují do TCP nebo UDP segmentu.

Hlavním účelem protokolu *L2TP* je tunelování (zapouzdřování) paketů protokolu PPP a jeho potomků (včetně ADSL/VDSL provozu), je velmi oblíbený u různých ISP. Zapouzdřuje se do UDP segmentů, následně do IP paketu.

Stejně jako GRE, ani L2TP neprovádí šifrování, tedy je obvykle kombinován s IPSec v transportním módu, což je standardizováno v RFC 3193. Takže k paketu, který má být takto přenesen:

- se nejdřív na L2 přidá PPP záhlaví a zápatí (PPP zajistí navázání a kontrolu průběhu spojení, základní autentizaci, v záhlaví je identifikace zapouzdřeného protokolu, v zápatí kontrolní součet), ale tento krok není povinný (L2TP dokáže zapouzdřit i IP pakety),
- pak se přidá záhlaví L2TP (v záhlaví jsou parametry pro tunel a relaci a další údaje),
- tento L2TP rámec se zapouzdří do UDP (takže smyčka nahoru),

- dále se může přidat IPSec záhlaví a záhlaví nosného protokolu, obvykle IP.

L2TP je podporován v Linuxu i ve Windows. V Linuxu se dnes doporučuje použití OpenL2TP (s využitím IPSec), ve Windows se nachází klient L2TP/IPSec pro transportní mód.



Poznámka:

Ve Windows se také můžeme setkat s protokolem SSTP (Secure Socket Tunneling Protocol), který přenáší PPP nebo L2TP šifrované s využitím protokolu SSL.



8.5.4 OpenVPN



Zajímavou implementací VPN je projekt *OpenVPN*. Běží na většině známých UNIXových systémů včetně Linuxu, existuje i varianta pro Windows. Je to otevřené řešení pro remote access, které zvládá jak vytvoření tunelu, tak i šifrování, šifrují se pouze přenášená data.

Toto řešení pracuje na vyšších vrstvách (nad L3), což znamená, že například není třeba řešit problémy s NAT či jinými překlady IP adres. Pro šifrování, autentizaci a správu klíčů se používá protokol SSL nebo TLS, na transportní vrstvě se používá většinou UDP, ale je možné použít i TCP. Protokol TLS je sice standardizován, ale celé řešení OpenVPN nikoliv.

OpenVPN se velmi často používá i ve firmách pro implementaci remote-access přístupu zaměstnanců na home office nebo na pracovní cestě. Stačí mít na hranici sítě koncentrátor (server) a na straně klientů klientskou část OpenVPN.

8.5.5 SSH

O SSH už něco víme z kapitoly o některých aplikačních protokolech, zde se pouze soustředíme na SSH spojení, které taktéž může být chápáno jako tunel.

SSH je protokolem vyšších vrstev a komunikuje přes TCP. Na rozdíl od TLS a některých dalších podobných protokolů SSH-2 nabízí správu více navázaných relací najednou (TLS pouze jedno spojení).

Pomocí SSH lze vytvářet tunely typu remote access, point-to-point, na vyšších vrstvách. Tunelování funguje na principu předávání (forwardování) provozu na porty, na které je napojen SSH tunel. SSH server naslouchá na TCP portu 22.



Existuje více různých implementací SSH. K nejoblíbenějším patří nástroj *OpenSSH*, pro který existuje klientská i serverová varianta (spíše pro UNIXové systémy). Instalace a používání jsou popsány v příloze na straně 294. Pro Windows existují mezi volně šířitelným softwarem spíše klienty, například *PuTTY*.

Když používáme SSH, měli bychom využívat bezpečné verze softwaru pro přenos dat. V UNIXu (včetně Linuxu) používáme příkazy *scp* (kopírování, místo *cp*) a *sftp* (místo *ftp*). Existuje také souborový manažer *WinSCP* implementující obdobu těchto příkazů.



Další informace:

- <https://blog.tracks.com/2014/05/17/ssh-tunnel-local-and-remote-port-forwarding-explained-with-examples.html>
- <https://www.ietf.org/rfc/rfc4254.txt>
- <https://www.debianhelp.co.uk/ssh.htm>





Úkol

Ve zmíněné příloze si projděte instalaci a používání SSH. Jsou tam také odkazy na další informace.



8.5.6 MPLS VPN

MPLS síť se používá spíše pro implementaci propojení firemních poboček tunelem (příp. firmy a obchodního partnera), tedy VPN typu site-to-site, spoje jsou typu point-to-point. S MPLS sítěmi jsme se již seznámili a víme, že se používá systém záhlaví, z nichž jedno je určeno právě pro virtuální síť. Řešení pracuje na rozhraní vrstev L2 a L3.



V síti MPLS VPN rozeznáváme tato zařízení:

- CE (Customer Edge) – hraniční uzel sítě zákazníka, přes který se data přenášejí do MPLS tunelu,
- PE (Provider Edge) – hraniční uzel sítě poskytovatele řešení, je přímo spojen s uzlem CE,
- P (Provider) – vnitřní uzly MPLS sítě poskytovatele, v podstatě Core MPLS směrovače.

Uzly CE nejsou součástí MPLS sítě, jsou na nich vyžadovány pouze protokoly TCP/IP (konkrétně hlavně IP). Uzly PE jsou hraničními (LER) uzly MPLS sítě a implementují protokol BGP, příp. iBGP (distribuce tabulek mezi uzly). K jednomu PE uzlu může být připojeno více CE uzlů.

IP paket přicházející z CE na uzlu PE je opatřen MPLS záhlavím se dvěma MPLS labely:

- vnitřní (B=1) bude využit až na druhém konci tunelu, určuje CE příjemce,
- vnější (B=0) určuje cestu v MPLS síti přes P uzly.

VRF tabulka (VPN Routing and Forwarding Table) je tabulka směrovacích informací pro konkrétní VPN a síť. Nachází se na uzlech PE, jeden PE může obsahovat i více VRF.



V tabulce je konkrétní port (rozhraní) asociován s konkrétní VRF (patří do příslušné VPN), pokud paket přijde z rozhraní napojeného na určitou VRF, podle dané VRF se rozhoduje, co s ním. To znamená, že VRF určuje, ke kterému CE se má dostat paket příchozí z daného rozhraní.



Pro MPLS VPN existují dvě základní řešení – MPLS Layer-3 VPN, MPLS Layer-2 VPN. Zapouzdřovat se mohou jak IP pakety, tak i ethernetové rámce (řešení Ethernet over MPLS – EoMPLS), případně PDU jiných protokolů na těchto vrstvách.

Nevýhodou MPLS VPN je, že potřebujeme externího poskytovatele této služby (obvykle ISP, přes jehož WAN síť náš tunel povede), tyto tunely si nemůžeme vytvořit sami. Naopak výhodou je, že pakety přenášené tunelem neprocházejí přes nezabezpečený Internet, konfigurace je na straně poskytovatele služby a součástí služby je i QoS, takže i přes poměrně vysoké finanční náklady si tato služba své zákazníky nachází.

8.5.7 Další možnosti řešení VPN

IP-in-IP je jednoduše zapouzdření IP paketu do IP paketu. Může jít o stejnou nebo různou verzi, často se například tímto způsobem řeší transport IPv6 paketů přes úsek sítě, který „nerozumí“ IPv6 (takže IPv6 paket zapouzdříme do IPv4 paketu).

Důležitým prvkem IP-in-IP je změna IP adresy v záhlaví, což je vlastně jedna z vlastností tunelování: ve vnějším záhlaví se jako zdrojová použije IP adresa zařízení (routeru) na začátku tunelu, jako cílová

se dosadí adresa konce tunelu. Na rozdíl od NAT (což je vlastně také překlad adres) se původní záhlaví zachovává (NAT zasahuje do původního záhlaví, nové nevytváří).

V případě Ipv6 je přímo v tomto mechanismu vestavěna možnost šifrování (ve vnějším IP záhlaví by pak bylo volitelné bezpečnostní záhlaví), u IPv4 bychom šifrování museli zajistit jinak (třeba pomocí protokolu IPSec).

VPLS (Virtual Private LAN Services) je multipoint řešení pracující na linkové vrstvě, v podstatě multipointová obdoba EoMPLS určená pro aplikace, které vyžadují funkčnost vzdálené multicast a broadcast komunikace. Poskytovatel této služby vlastně z pohledu klienta simuluje switch (jsme na L2) propojující vzdálené LAN sítě.

DMVPN (Dynamic Multipoint VPN) je řešení postavené na protokolu NHRP (Next Hop Resolution Protocol), což je rozšíření protokolu GRE pro využití na multipoint spojích. Pracuje na síťové vrstvě a nevyžaduje externího poskytovatele (konfiguraci si můžeme provést u sebe, přenos probíhá běžně přes Internet, ovšem v tunelu). Dá se říct, že jde o multipoint alternativu ke GRE a IPSec.



Další informace:

- <http://www.ipsec-howto.org/ipsec-howto.pdf>
- <https://openvpn.net/>
- <https://www.root.cz/clanky/openvpn-vpn-jednoduse/>
- <https://www.root.cz/clanky/openvpn-vpn-jednoduse-2/>
- https://www.cisco.com/c/en/us/td/docs/net_mgmt/vpn_solutions_center/2-0/installation/guide/IG20_1.pdf
- <https://www.samuraj-cz.com/clanek/vpn-1-ipsec-vpn-a-cisco/>
- <https://www.soom.cz/clanky/319-Protokol-sem-protokol-tam>
- <http://www.cs.vsb.cz/grygarek/TPS/projekty/0607Z/L2TP.pdf>
- http://www.cs.vsb.cz/grygarek/TPS/projekty/0506Z/krs008_TPS_projekt.pdf
- <https://www.cisco.com/c/en/us/support/docs/multiprotocol-label-switching-mpls/mpls/29828-mplsvpnte.html>



Souhrn kapitoly

V této kapitole jsme se zaměřili na teorii managementu sítě – probrali jsme stručně ISO model řízení sítě s protokolem CMIP a pak jsme se více jmenovali používanějším protokolu SNMP pocházejícím z TCP/IP. Následovalo téma logování provozu pomocí standardu Syslog a monitorování sítě s využitím RMON, Snortu a NetFlow. Poslední sekce kapitoly je věnována zabezpečení komunikace přes Internet pomocí VPN, přičemž jsme probrali několik různých řešení: zejména IPSec VPN, GRE tunely, L2TP, OpenVPN, SSH a MPLS tunely.



Bezpečnost a správa

 **Rychlý náhled:** V této kapitole se budeme zabývat různými tématy z oblasti bezpečnosti. Podíváme se na některé typy kybernetických útoků, dále na české bezpečnostní struktury a bezpečnostní týmy, následuje související hardware (síťové analyzátory, firewally), dohledové systémy a SIEM.

 **Klíčová slova:** Mapping, sniffing, SQL injection, spoofing, DoS, DDoS, Man-in-the-Middle (MitM), hijacking, útoky na heslo, sociální inženýrství, CZ.NIC, CESNET2, bezpečnostní týmy, CSIRT, síťový analyzátor, firewall, DMZ, paketový filtr, ACL, SPI, IDS/IPS, DPI, proxy, dohledový systém, SIEM.

 **Cíle studia:** Cílem této kapitoly je získat přehled v oblasti kybernetické bezpečnosti, a to co se týče zařízení používaných v zabezpečení a monitoringu sítě, tak i u protokolů, metod a technologií používaných pro zabezpečení komunikace.

9.1 Pojmy z oblasti kyberbezpečnosti

Protože zdroje o kyberbezpečnosti jsou obvykle v anglickém jazyce, budou nás i zde zajímat především pojmy v angličtině (v češtině se obvykle používají také).

 **Asset** (překládá se také jako aktivum, majetek) je subjekt nebo objekt, který potřebujeme chránit: lidé, data, zařízení, atd., cokoliv ve vlastnictví dané společnosti, za co je odpovědná a co souvisí s její činností.

 **Threat** (hrozba, ohrožení) je potenciální riziko, které může hrozit společnosti a assetům. Pokud daná hrozba existuje, tak ještě není řečeno, že daná situace opravdu může nastat, ale pravděpodobnost není nulová. Například s určitou pravděpodobností může dojít k úniku dat z důvodu nedostatečné konfigurace firewallu nebo upovídané sekretářky, k poškození a úniku dat může dojít nesprávnou reakcí zaměstnance na podvržený e-mail, atd., nicméně s těmito riziky se dá do určité míry pracovat a snižovat je.

Rizika mohou souviset s únikem citlivých dat nebo krádeží drahého zařízení, poškozením dat, zařízení či infrastruktury, případně s konkurenčním bojem.

Pojem *Threat Intelligence* (nebo taky Cyber Threat Intelligence) označuje souhrn metod, informací a jiných prostředků používaných pro eliminaci nebo alespoň snižování rizika těchto hrozeb, také hovoříme o metodách založených na důkazech, abychom zdůraznili, jak důležité je mít veškeré bezpečnostní

informace dobře podložené.

Také se můžeme setkat se zkratkou *APT* (Advanced Persistent Threat), která označuje pokročilou trvalou hrozbu, již musí daná společnost dlouhodobě čelit a kterou nedokáže zcela eliminovat. Tyto hrozby dnes často bývají způsobeny útočníky či jejich skupinami sponzorovanými některým státem, jejich cílem je získat konkurenční výhodu vlastnímu průmyslu a naopak oslabit průmysl a infrastrukturu jiného státu.

 *Vulnerability* (zranitelnost) je slabina v některém systému (aplikaci, operačním systému, ale může být i ve firmwaru nebo v integrovaných obvodech, dokonce také v metodice – nemusí být technologická). Existence zranitelnosti znamená možnost jejího zneužití a tedy hrozbu (threat).

Z principu by mělo být důležité zranitelnosti odstraňovat, například aktualizacemi softwaru.

 *Exploit* je mechanismus/postup, který lze použít ke zneužití některé zranitelnosti. Exploity jsou vytvářeny nejen „zlými“ hackery, ale také white a blue hackery, kteří v rámci Bug Bounty programů hledají zranitelnosti, pomocí exploitů ověřují jejich rizikovost a výsledek hlásí společnosti, jejíž produktům riziko hrozí.

 *Attack Surface* (snad by mohlo být přeloženo jako oblast či plocha útoku) je souhrn míst, přes která by se útočník mohl dostat do naší sítě a dalších míst zájmu, včetně cloudu, a tedy jde o místa, na která se soustředíme při ochraně assetů. Nemusí jít nutně o fyzická místa.

 *Perimeter* je víceméně zabezpečená hranice mezi sítí dané společnosti a veřejnou sítí. Zejména na perimetru se nejvíce řeší monitoring a ochrana.

9.2 Kybernetické útoky

V počítačové síti je třeba čelit různým typům útoků, z nichž některé jsou vzájemně provázány.

9.2.1 Průzkumné útoky

Průzkumné útoky (reconnaissance attacks, recon útoky, mapování) – jde vlastně o jakousi přípravu k následujícím útokům. Hacker získává co nejvíc informací (otevřené porty, použité rozsahy IP adres, operační systémy a jejich verze, atd.), aby vlastní útok mohl být co nejúspěšnější.

Obvyklé nástroje jsou *nslookup*, *dig*, *whois*, *wireshark*, *nmap* a jeho frontendy jako například *Zenmap*, *kismet* a další.

 *Network sniffing* (naslouchání na síti) probíhá tak, že packet sniffer (zachytávač nebo odposlouchávač paketů) odkloní provoz na síti, zachytává pakety a získává z nich informace (pak je posílá dál k cíli, nedoručené pakety by znamenaly jeho odhalení).

Účelem sniffingu je především získávat hesla přenášená v textovém tvaru a také další citlivé informace. Také je možné zachytávat a pozměňovat obsah paketů anebo dokonce číst a pozměňovat informace na uzlech v síti (včetně konfiguračních souborů nebo hesel).

Účinnou obranou je především spolehlivé šifrování (včetně autentizace po síti), a také fyzické zabezpečení sítě, protože tento typ útoku vyžaduje fyzický přístup k síti. To může být problém u bezdrátových typů připojení.

Packet sniffer (třeba Wireshark) však může být zcela legálně využíván administrátorem ke sledování provozu na síti.

 *ICMP útoky* mohou sloužit ke zjišťování topologie sítě (podsítě, zařízení), ale také mohou sloužit k provádění DoS útoků nebo podvržení informací. K průzkumným útokům mohou sloužit například tyto ICMP zprávy:

- ICMP Echo request a reply,
- ICMP unreachable,
- ICMP mask reply,
- ICMP router discovery, atd.

 *Odposlech* (eavesdropping) je zachytávání a odposlouchávání síťového provozu a extrakce informací z tohoto provozu. Tento provoz je buď posílán dále regulérnímu adresátovi, nebo sice posílán, ale pozměněný, nebo zahazován.

 *Man-in-the-Middle* (MitM) – hacker se dostane mezi dva komunikující počítače a odposlouchává nebo dokonce pozměňuje komunikaci mezi nimi. Tento typ útoku souvisí s některými jinými únos spojení, zjišťování hesla apod.



Další informace:

- <https://www.sciencedirect.com/topics/computer-science/reconnaissance>
- https://owasp.org/www-pdf-archive/ICMP_Attacks.pdf



9.2.2 Přístupové útoky

Přístupové útoky (access attacks) slouží k získání přístupu k určitému systému, navýšení (eskalaci) oprávnění, a následně například ke krádeži nebo poškození dat, zneužití některé zranitelnosti apod. Často jde o podvržení adresy v paketu či rámci nebo podvržení údaje v některé tabulce na síťovém zařízení.

 *IP Spoofing* je podvržení IP adresy v odesílaném paketu. Komunikující uzel předstírá, že je vlastníkem IP adresy, která ve skutečnosti patří jinému uzlu (nebo nepatří žádnému uzlu), ale hlavně patří do sítě či VLAN s požadovaným oprávněním nebo je jí jinak důvěřováno. Obvykle jde o případ, kdy se útočník snaží předstírat, že je řádným členem privátní sítě používající určitý rozsah IP adres, případně se pokouší vydávat za konkrétní uzel v síti.

Častým způsobem použití je zneužití některých protokolů, včetně protokolu SMTP pro odesílání e-mailů. Používá se také jako prostředek pro sofistikovanější útoky, například DoS (v paketech, které jsou zasílány na napadené zařízení, je zdrojová adresa podvržená). Hacker se musí vypořádat pouze s jedním problémem: pokud potřebuje obdržet i odpověď na odesílaný paket (což někdy není nutné), musí tu odpověď nějak zachytit.

 *Další typy spoofingu*: podvržení identity se dá „spáchat“ i jinak, často se jedná o hacknutí některých tabulek s adresami, například:

- ARP Spoofing znamená podvržení záznamů v ARP tabulkách na zařízeních (je napaden mechanismus překladu IPv4 adresy na MAC adresu), obdoba pro IPv6 je ND Spoofing,
- DNS Spoofing je podvržení záznamů v zónovém souboru na DNS serveru (je napaden mechanismus překladu jmenné adresy na IP adresu).

 *Hijacking* (únos spojení) – jde především o útoky související s vytáčeným spojením, ale také obecně s jakýmkoliv placeným spojením vyžadujícím autentizaci (například soukromé Wi-Fi sítě). Účinnou obranou je použití vhodného šifrovacího protokolu při autentizaci.

 *SQL Injection* – hackerovi se podaří podstrčit (inject) SQL příkaz někam, kde nemá co dělat, případně vhodnou volbou řetězců „vyrobit“ a poslat serveru SQL příkaz. Většinou se jedná o zneužití formulářů na webu nebo řetězce předávaného přes HTTP GET, kdy útočník vyplní místo běžného řetězce „škodlivý“ řetězec obsahující symboly se speciálním významem. Typický zneužívaný symbol je třeba apostrof.

Proti tomuto typu útoku se celkem dá bránit, pokud programátor webové aplikace (tedy webového rozhraní k databázi) zajistí pečlivou analýzu vstupů od uživatele, a především je třeba mít správně nastavená oprávnění u databáze (například požadavky na destruktivní operace typu DROP, DELETE, ale také změny dat nebo vyžádání si chráněných informací nemají být prováděny běžným uživatelem z nechráněné části sítě, potažmo Internetu).

 *Útoky na zjištění hesla* – heslo lze zjistit například brute-force útokem (hrubá síla), použitím trojského koně a některými výše popsány metodami. Další možností je sociální inženýrství, které je v současné době na vzestupu (uživatel defacto tento údaj prozradí sám a dobrovolně, je z něho podvodně vylákan).

Brute-force útok může znamenat zkoušení všech možných kombinací znaků v řetězci o „vhodné“ délce, ale může být veden jako *slovníkový*, tj. automaticky jsou zkoušeny všechny řetězce z vytvořeného slovníku (nemusí jít jen o anglická slova!). Proti tomu se lze bránit nastavením maximálního počtu neúspěšných přihlášení, po překročení tohoto limitu je přístup zcela zablokován. Dále je možné po uživateli vyžadovat používání tzv. „silného“ hesla (dostatečně dlouhého, obsahujícího jak písmena, tak i číslice a další znaky – dvojtečka, procento, zavináč, apod.), s čímž ale bývají problémy (uživatelé raději volí takové heslo, které si dokážou zapamatovat).

 *Lidský faktor* – „nepřítel“ může být i uvnitř sítě, a to dokonce i takový, který to o sobě netuší. Zvláště v poslední době jsou pro útoky často voleny *metody sociálního inženýrství*. Sociální inženýrství je metoda, jak z uživatele vytáhnout potřebné informace třeba i bez použití techniky, a to jeho uvedením v omyl (obelstěním). Uživatel je přesvědčen, že informace předává důvěryhodné osobě z naprosto nutných důvodů.

Útočník se vydává například za zaměstnance bezpečnostního oddělení, opraváře, zaměstnance telefonní společnosti, nového kolegu, apod. a nenápadně ze svých obětí vytáhne vše potřebné (hlavně hesla), případně si „vyzkouší“ nový skvělý počítač nebo se jinak dostane k technice na pracovišti. Stává se to především ve větších firmách, kde se nepočítá s tím, že by se všichni zaměstnanci znali, ale kontakt může probíhat i „neosobně“ přes telefon nebo mail. Právě do telefonu jsou zaměstnanci schopni říci o své firmě neuvěřitelné věci včetně těch, které jsou obchodním tajemstvím.

Popřípadě mnoho uživatelů nepochopitelně důvěřuje e-mailu: stačí například poslat e-mail s informací, že zřejmě došlo ke kompromitaci účtu, přičemž pro kontrolu je třeba zpět odeslat přihlašovací údaje, aby administrátor ověřil, zda je vše v pořádku (navíc e-mail musí být „správně“ graficky vyveden), a spousta uživatelů slepě poslechne.

Sociální inženýrství může mít i „materiální“ povahu – například volně „pohozená“ přenosná zařízení, která jsou pro mnoho lidí neodolatelná. Podle DHS¹ je kolem 60 % běžných uživatelů schopno náhodně nalezený USB flash disk připojit ke svému počítači, třebaže netuší, zda se na něm nenachází škodlivý

¹DHS (Department of Homeland Security, <http://www.dhs.gov>)

software. Podobně lze zneužít i jiná přenosná zařízení, zde je nejlepší ochranou poctivost, tedy odevzdat nalezený předmět do ztrát a nálezů.

Sociální inženýrství je v současné době jedna z nejpoužívanějších (a taky nejefektivnějších) metod získávání utajených informací. Na to by měli myslet zejména administrátoři firemních sítí a zajistit patřičné školení všech, kdo se do firemní sítě připojují.



Další informace:

- <https://www.ciscopress.com/articles/article.asp?p=1728833&seqNum=3>
- http://www.cis.syr.edu/~wedu/seed/Book/book_sample_tcp.pdf
- <https://www.utc.edu/center-academic-excellence-cyber-defense/pdfs/course-paper-5620-attacktcpip.pdf>
- <https://unit42.paloaltonetworks.com/dns-tunneling-how-dns-can-be-abused-by-malicious-actors/>



9.2.3 Útoky způsobující odmítnutí služby

Útoky způsobující odmítnutí služby mají za cíl odstavit určité zařízení (většinou server), tedy znemožnit mu poskytovat službu nebo alespoň velmi zpomalit jeho činnost.



DoS (Denial of Service) – jde o vynucení odmítnutí služby legitimním uživatelům. Tento útok probíhá buď využitím chyby v kódu (některého protokolu nebo operačního systému), vyvolaným přetížením sítě nebo podvrženými zprávami – pakety o stavu sítě. Server je zahlcen žádostmi o spojení (TCP handshake) nebo žádostmi o data, které není schopen vyřídit, a proto i následný provoz je zadržen. Případně může jít o útok způsobující zacyklení nebo přílišné zaneprázdnění procesů běžících na napadeném zařízení (tedy přetížení procesoru), nebo zahlcení jeho operační paměti.



DDoS (Distributed DoS) – velmi nebezpečná varianta předchozího typu útoku, proti které se téměř nelze bránit (jde to, ale za cenu velkých nákladů, existují například ochranné cloudové služby, přes které je příchozí provoz filtrován). Častým nedobrovolným „účastníkem“ DDoS útoků jsou sítě botů. Bot je napadený počítač ovládaný na dálku hackerem (bez dovolení a většinou také bez vědomí právoplatného majitele). Sítě botů, a to i velmi rozsáhlé (resp. jejich služby), jsou „prodávány“ na černém trhu kromě jiného právě k DDoS útokům.

DDoS útok se dá provést i pomocí jinak celkem užitečných mechanismů. Například typ útoku *Ping Flood* (čte se [flad], znamená „záplava“) probíhá tak, že napadený stroj je zahlcen zprávami ICMP Echo Request (to jsou ty, které posílá příkaz `ping`). Tím se zahlcuje jak jeho vstupní kapacita (příjem), tak i jeho výstupní kapacita (odesílání) – když se pokouší odpovídat. Aby pachatel skryl svou identitu, v IP packetech zapouzdřujících tyto zprávy falšuje zdrojovou IP adresu. DDoS varianta útoku znamená, že ICMP zprávy posílají zařízení zapojená do botnetu.

Smurf Attack („šmoulí“ útok) má ještě vyšší míru distribuovanosti. Spočívá v tom, že velké množství zařízení (typicky serverů a routerů) dostává zprávy ICMP Echo Request, jejichž zdrojová adresa je podvržená – nastavená na adresu oběti. Tato zařízení odpovídají zprávami ICMP Echo Reply na adresu oběti, čímž zahlcují její vstupní kapacitu.



Poznámka:

Existuje mnohem více typů útoků, zde byl pouze přehled těch nejznámějších.





Další informace:

- <https://www.cloudflare.com/learning/ddos/what-is-a-ddos-attack/>
- <https://www.imperva.com/blog/security-glossary-top-12-ddos-attack-types-need-know/>



Úkol

Podívejte se na <https://www.cloudflare.com/learning/ddos/ping-icmp-flood-ddos-attack/>. Tato stránka popisuje útok typu ICMP Ping flood. Nahoře v menu je položka *Common DDoS Attacks* – projděte si zde uvedené útoky, jestli se zde nevyskytuje některý, o kterém jste ještě neslyšeli.



9.3 České organizace

9.3.1 CZ.NIC

Organizace CZ.NIC z.s.p.o. (zájmové sdružení právnických osob) je sdružení společností zainteresovaných především v provozování síťové infrastruktury a poskytování služeb přístupu k Internetu. Toto sdružení především zabezpečuje funkčnost české národní domény **.cz** (včetně DNSSEC), provozuje registr doménových jmen příslušejících do této domény, angažuje se v různých projektech obvykle spojených s bezpečností. Také je provozovatelem národního bezpečnostního týmu CSIRT.CZ (ale má také svůj vlastní bezpečnostní tým CZ.NIC-CSIRT).

Pro správu registru domén CZ.NIC používá vlastní open-source systém FRED (Free Registry for ENUM and Domains) dostupný každému pod licencí GNU GPL, pro tento systém se rozhodli i další registrátoři ve světě.

Důležitou součástí činnosti CZ.NIC je osvětová činnost, ať už pro běžné uživatele informačních technologií, nebo pro odbornou veřejnost. Na webu sdružení také najdeme volně ke stažení elektronické knihy, zejména o kybernetické bezpečnosti.



<https://www.nic.cz>

9.3.2 CESNET2

CESNET2 (Czech Education and Scientific NETwork) je národní vysokorychlostní počítačová síť určená pro vzdělávání, vědu, výzkum a vývoj, založená roku 1996. Propojuje především vysoké školy a výzkumné ústavy včetně všech ústavů Akademie věd České republiky, jsou připojeny také některé nemocnice a střední školy. Je jedním z členů CZ.NIC.

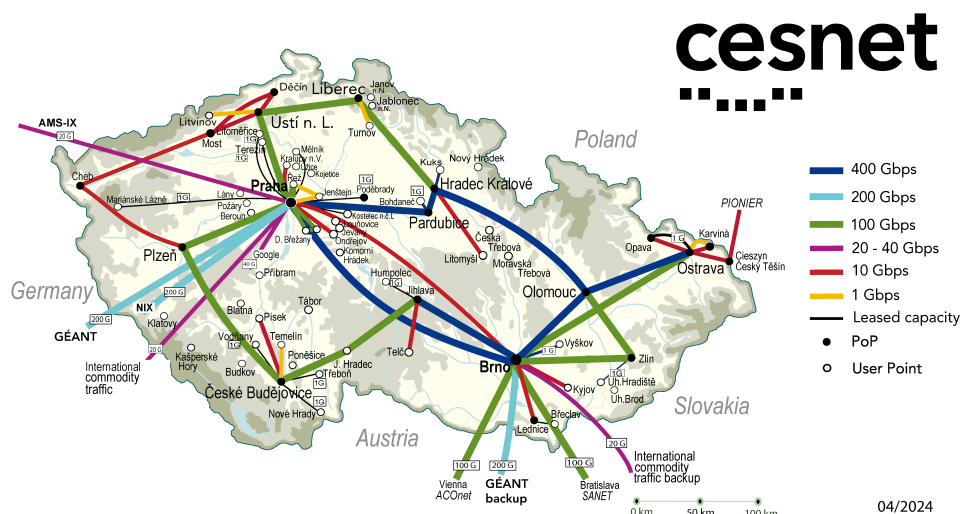
Jde už o druhou generaci této sítě. Původní síť CESNET první generace dosahovala rychlostí v desítkách až stovkách Mb/s, zatímco CESNET2 dosahuje rychlostí kolem 10 Gb/s, na některých místech až 100 Gb/s. V polovině 90. let se hodně používala síť ATM, později kolem roku 2000 byla postupně zaváděna technologie PoS (Packet over SDH) a MPLS a přenosová soustava byla postavena na DWDM.



<http://www.cesnet.cz>

Topologie sítě z roku 2015 je naznačena na obrázku 9.1.

²Zdroj: <http://www.cesnet.cz>

Obrázek 9.1: Topologie sítě CESNET (rok 2024)²

Účelem sítě CESNET2 je poskytování vysokorychlostních přenosů pro účely výzkumu a vzdělání, zabývá se také různými typy využití přenosové kapacity (VoIP, multimediální přenosy a videokonference, QoS, vysokorychlostní bezdrátové připojení, atd.). V současné době se vzhledem k momentální situaci v oblasti sítí hodně angažuje v bezpečnostních technologiích.

Zajímavý projekt je také *EduRoam.cz*, který umožňuje připojit se do bezdrátové sítě kteréhokoli člena CESNETu (většinou jde o podporu mobility zaměstnanců vysokých škol a výzkumných ústavů).

CESNET2 je také napojena na mezinárodní sítě s podobným zaměřením, například na evropskou síť GÉANT, slovenskou SANET, rakouskou ACONET a polskou PIONIER.

 *Metacentrum* je aktivita CESNETu věnovaná provozu národního výpočetního gridu (distribuované výpočetní sítě). V gridu je v současné době (2015) několik výkonných výpočetních center zahrnujících celkem 10 712 procesorů. Implementuje různé metody pro BigData, včetně algoritmu Hadoop od společnosti Google. Grid je využíván především akademickými pracovníky pro náročné výpočty ve výzkumu a vývoji. Členství je bezplatné pro akademické pracovníky, kteří jsou zaměstnanci nebo studenty členů CESNETu.

9.3.3 Bezpečnostní týmy

 Svůj bezpečnostní tým zasahující (nejen) proti kyberkriminalitě by dnes měla mít každá větší organizace. Setkáme se zde se dvěma důležitými zkratkami:

- CERT (Computer Emergency Response Team),
- CSIRT (Computer Security Response Team).

V obou případech jde víceméně o totéž – tým odborníků určený k boji proti kyberkriminalitě, tedy lidé aktivně se zabývající bezpečností v oblasti ICT. Obě zkratky se používají pro označení bezpečnostních týmů jak na národní úrovni, tak i na úrovni soukromých společností. Ještě něco mají tyto týmy společného – spolupráci.

CERT je ve skutečnosti chráněnou značkou ve vlastnictví Carnegie-Mellon University, proto se v praxi setkáme častěji se zkratkou CSIRT nebo nějakou obdobou zkratky CERT.

 **CSIRT.CZ** je CSIRT týmem provozovaným sdružením CZ.NIC (to je správce domény .CZ – české národní domény). Jedná se o hlavní CSIRT tým České republiky, a to od roku 2012, kdy bylo uzavřeno memorandum mezi CZ.NIC a Národním centrem kybernetické bezpečnosti (resp. Ministerstvem vnitra ČR), které má ze zákona povinnost zajistit provoz bezpečnostního týmu bojujícího proti kyberkriminalitě. Nicméně tento tým ve skutečnosti existuje už od roku 2007 (na státní úrovni se totiž o bezpečnostních týmech dlouho jen jednalo, a teprve potom, co vznikly iniciativou v soukromém sektoru, byla jejich činnost posvěcena zákonem).

Tým CSIRT.CZ plní tyto úkoly:

- řešení zjištěných bezpečnostních incidentů ve spolupráci s českými provozovateli sítí a zahraničními partnery,
- pomoc provozovatelům sítí při zřizování vlastních CSIRT týmů,
- poskytuje reaktivní i proaktivní služby, poradenství,
- provozuje *honeypoty* („vějičky“ pro útočníky – zařízení záměrně vypadající jako slabě zabezpečená),
- mohou zjišťovat možné autory DoS útoků a předávat informace orgánům, které pak na toto oznámení reagují, nebo informovat provozovatele sítě o napadení této sítě.

Na mezinárodní bezpečnostní scéně tým spolupracuje především s organizacemi TERENA (Trans-European Research and Education Networking Association) a FIRST (Forum for Incident Response and Security Team).

Proaktivními službami se rozumí služby v oblasti bezpečnostní prevence, například informační pomoc, nabídky školení, provoz IDS/IPS (systémů detekujících možnost nebo pravděpodobnost napadení) apod. *Reaktivní služby* naopak reagují na již proběhlý nebo probíhající bezpečnostní incident, v tomto případě například přijímání oznámení a podnětů o incidentech a reakce na ně.

K novějším službám týmu patří služba bezplatného penetračního testování webu (tzv. Skener webu), a dále taktéž bezplatná služba zátěžových testů odpovídajících svou intenzitou DDoS útokům, které na několik dnů roku 2013 silně zpomalily české sítě.

 **CESNET-CERTS** je bezpečnostní tým společnosti CESNET, provozovatele sítě propojující akademické a výzkumné organizace (především vysoké školy a ústavy Akademie věd). Pracuje už od roku 2004 a nese odpovědnost za bezpečnost především v sítích nějak souvisejících s akademickou a vědeckou sférou (včetně domény eduroam.cz). Povinnosti a pravomoci jsou v podstatě podobné jako u CSIRT.CZ, jen vztažené na sítě, kde má tým pravomoc.



Poznámka:

CSIRT týmy (ani celostátní) nemají žádné pravomoce k řešení bezpečnostních incidentů na jiné než technické úrovni, nejde o represivní orgány. Mohou pouze monitorovat a radit, pravomoci k razantnějším reakcím mají pouze ve své vlastní síti.



 Další: svůj bezpečnostní tým by měly mít všechny větší organizace. Vlastní CSIRT týmy působí například v sítích ISP (poskytovatelů Internetu), bank, poskytovatelů energií, provozovatelů energetických společností (jako je u nás například ČEZ), velkých automobilek, velkých poskytovatelů internetových služeb (Google, Seznam, Facebook), atd. Svůj bezpečnostní tým má také Slezská univerzita.

**Poznámka:**

Každý takový tým se určitým způsobem prezentuje – musí mít někde uvedeny kontaktní informace, vymezení pole působnosti a odpovědnosti, seznam služeb, které poskytuje (především v rámci sítě svého provozovatele). Musí být zřejmé, kdo se na tým může obrátit, s jakým problémem k řešení a jakým způsobem (může to být třeba formulář na webu). Pro tyto účely obvykle tým pořádá i různá školení (to je v podstatě jedno z proaktivních opatření).

**Další informace:**

- Kybernetická bezpečnost na webu NÚKIB (Národní úřad kybernetické bezpečnosti):
<https://www.nukib.cz/cs/kyberneticka-bezpecnost/>
- <https://www.csirt.cz/>
- <http://www.cesnet.cz/sluzby/reseni-bezpecnostnich-incidentu/>
- <https://csirt.cesnet.cz/>
- <http://www.earchiv.cz/b08/b0408002.php3>



9.4 Bezpečnostní zařízení v síti

9.4.1 Monitorování provozu

 **Zrcadlení provozu** (mirroring) nám dává možnost sledovat a analyzovat provoz na určitém segmentu sítě, případně při určité konfiguraci i provoz na jiném segmentu, pokud dokážeme sledovaný provoz přeposílat například přes k tomu určenou VLAN síť.

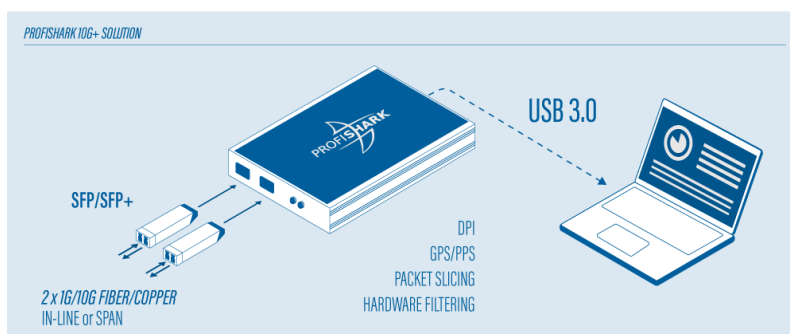
 **SPAN** (Switch Port ANalyzer) je vlastnost switchů, kterou také nazýváme zrcadlení portů (port mirroring). Provoz přicházející na určitý port (tzv. zdrojový SPAN port) je zrcadlen (tedy přeposlán, kopírován) na jiný port (monitor/destination SPAN port), ke kterému je připojeno zařízení pro zobrazení, uložení či analýzu paketů (v jednodušším případě počítač s Wiresharkem, případně IDS). Provoz je samozřejmě posílán také na ten port, pro který je určen.

Local SPAN znamená, že zdrojový i cílový SPAN port jsou na jednom switchi. RSPAN (Remote SPAN) oproti tomu znamená, že každý z těchto portů je na jiném switchi a mezi nimi existuje zajištěná cesta, po které sledovaný provoz přeposíláme (typicky je zařazen do k tomu určené VLAN). Cílový SPAN port je tedy na switchi, ke kterému je připojen IDS nebo jiné zařízení, na kterém probíhá ukládání nebo analýza sledovaného provozu.

 **Network Tap** (síťový odposlech) je speciální analyzátor, jehož úkolem je „sbírat“ provoz na daném síťovém segmentu a (obvykle) přeposílat na jeden či dva vybrané porty (u novějších i například na USB výstup). Může fungovat jako IDS, procházet všechny pakety, nebo může sledovat jenom konkrétní typ paketů (třeba multimediální, případně pro konkrétní protokol), bývá začleněn do systému ochrany sítě.

První podmínka se zdá jednoduchá, ale je třeba si uvědomit, že požadavek „v obou směrech“ znamená vlastně dvojnásobek běžného provozu, protože dnes se komunikuje v plném duplexu. Plně duplexní provoz

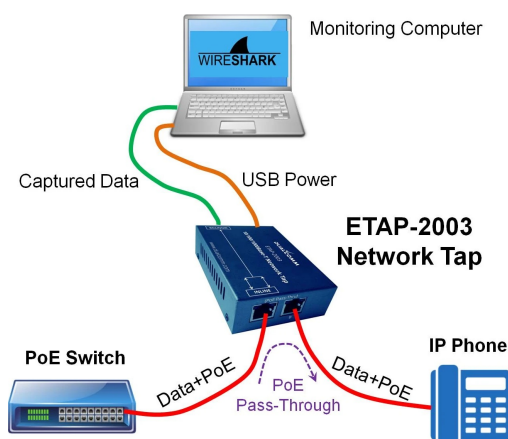
³Zdroj: <https://www.profitap.com/profishark-10g-plus/>

Obrázek 9.2: Schéma zapojení síťového odposlechu³

musíme do odposlechu dostat „v jednom směru“ (směrem k odposlouchajícímu počítači). Zařízení určená pro použití v sítích s velkým provozem mívají běžně pro tento směr minimálně dva porty, pro každý původní směr komunikace jeden.

Na obrázku 9.2 vidíme příklad zapojení síťového odposlechu v síti. Toto zařízení má:

- vstupy vedoucí ze sousedících zařízení, na obrázku jde o SFP moduly (například pro ethernetové optické kabely),

Obrázek 9.3: Schéma zapojení síťového odposlechu s PoE⁴

- výstup(-y), to může být jedno nebo několik síťových rozhraní, u novějších obvykle USB rozhraní vyšší verze (novější verze USB už mají dostatečnou propustnost, ale samozřejmě záleží na datovém toku mezi sledovanými zařízeními),
- napájení – obvykle bývá realizováno přes výstup (pokud jde o kroucenou dvojlinku, pak PoE, pokud jde o USB, pak se používají napájecí vodiče v USB kabelu, případně lze použít přídatný USB kabel pro napájení).

Na spoji vedoucím mezi sledovanými zařízeními taky může být používáno PoE. V tom případě by se tap neměl „přiživovat“, napájení by měl transparentně přeposílat, jak vidíme na obrázku 9.3.



Další informace:

Další informace o tapech najdeme na

- <https://www.profitap.com/profishark-10g-plus/>

⁴Zdroj: <https://www.dualcomm.com/collections/network-tap>

- <https://www.garlandtechnology.com/network-tap>
- <https://www.dualcomm.com/collections/network-tap>
- <https://www.netscout.com/product/enterprise/netscout-taps>



Obrázek 9.4: Dualcomm ETAP 2306: Dual-Mode GbE Tap pro kroucenou dvojlinku a optiku (4×RJ-45, 2×SFP s optikou) a Garland Single-mode pasivní TAP pro optiku⁵

9.4.2 Firewall

 Firewall je síťový prvek (hardwarový nebo softwarový), který slouží k řízení provozu mezi sítěmi s různou úrovní důvěryhodnosti či zabezpečení. Ve firewallu jsou definována *pravidla* pro komunikaci mezi sítěmi, podle kterých reaguje buď povolením komunikace, jejím zakázáním a nebo žádostí o vyjádření uživatele. Pravidla se obvykle týkají těchto údajů:

- zdroj a cíl komunikace, může být zadán IP adresou,
- číslo portu, přes který se komunikuje (tj. můžeme zablokovat používání některého portu), může jít o zdrojový i cílový port,
- používaný protokol,
- stav spojení, atd.

Firewall může být buď jen *jednosměrný* (filtruje pouze příchozí pakety) nebo *obousměrný* (filtruje příchozí i odchozí pakety). Lepší je samozřejmě obousměrný, dokáže efektivněji zachytit případné vynášení citlivých informací.

 Kromě softwarových firewallů existují také *hardwarové firewally* fungující jako mezilehlé prvky sítě. Dnes jsou většinou součástí routerů nebo jiných běžných síťových prvků. Hardwarový firewall má velkou výhodu v nižším riziku napadení (není závislý na operačním systému klientského počítače). Další výhodou je nezávislost na výkonu počítače – pokud je procesor počítače hodně zatížen, má to vliv i na činnost (především rychlost) softwarového firewallu na tomto počítači nainstalovaného. Hardwarový firewall (třeba vestavěný v jiném síťovém zařízení) takto ovlivněn není.

Z hlediska bezpečnosti je v domácnostech a menších firmách za ideální považována kombinace jednoduchého hardwarového firewallu ve směrovači (například ADSL router) a softwarového routeru na počítači, každý z nich jiného typu.

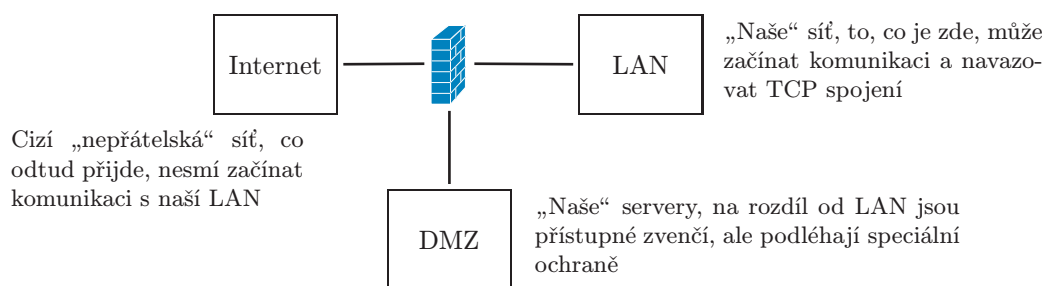
Softwarové firewally slouží buď k ochraně koncových uzlů, anebo mohou běžet v operačním systému nainstalovaném na (téměř) jakémkoliv síťovém prvku.

⁵Zdroje: <https://www.dualcomm.com/>, <https://www.garlandtechnology.com/>

 **Oblasti podle firewallu.** Z pohledu firewallu členíme síť do několika oblastí:

- *vnitřní síť* – do této oblasti patří vše, co je „naše“, čemu můžeme důvěřovat, uzel z vnitřní sítě může (s ohledem na přístupová oprávnění) iniciovat spojení k vnitřní i vnější síti,
- *vnější síť* – typicky Internet,
- *demilitarizovaná zóna (DMZ)* – zóna „nikoho“, z bezpečnostního hlediska někde mezi vnitřní a vnější sítí.

Do DMZ lze umístit například to, co je sice „naše“, ale má být přístupné z vnější sítě (například web server, mail server, DNS server). Servery v DMZ by měly být zabezpečeny tak, jako by byly opravdu přímo na Internetu, třebaže jsou od Internetu odděleny firewallem.



Obrázek 9.5: Oblasti z pohledu firewallu

Další možnost využití DMZ je propojení k třetí straně, které víceméně důvěřujeme, typicky dodavateli služby, kterou si nemůžeme zajistit vlastními silami. Obě možnosti můžeme zkombinovat a používat dvě demilitarizované zóny (nebo více).

Na síťovém zařízení podporujícím DMZ máme obvykle jeden nebo více (hardwarových) portů takto označených, případně můžeme některé porty sami nakonfigurovat tak, aby se s nimi zacházelo jako s DMZ (například tehdy, když chceme starší počítač využít jako hardwarový firewall, ukázky najdeme v příloze).

 **ZPS (Zone-based Policy Firewall).** Na obrázku 9.5 vidíme tři zóny – důvěryhodnou LAN (itranet), nedůvěryhodný Internet a DMZ. Zone-based firewallly (předpokládáme že mají více portů než jen tři) dokážou definovat tolik zón, kolik si administrátor vzpomene, přičemž každá zóna má svá pravidla pro komunikaci s jinými zónami a každé rozhraní přísluší do konkrétní zóny. Můžeme mít například pět různých zón (a do nich přiřazena jednotlivá síťová rozhraní) a každá bude mít trochu jinak nastavena firewallová pravidla.

 **OpenWRT** je zajímavý open-source projekt embedded Linuxu (tj. Linuxu upraveného pro konkrétní typ zařízení, v podstatě firmware daného zařízení) pro routery některých výrobců, například společnosti Mikrotik (RouterBoard). Jeho výhodou oproti „originálnímu“ operačnímu systému (například u desek RouterBoard je to RouterOS taktéž založený na Linuxu) je to, že má k běžným linuxovým distribucím mnohem blíže než jiné embedded Linuxy. Zatímco obvykle se firmware aktualizuje jako celek (to může být celkem nebezpečné, dotyčné zařízení se při poruše při aktualizaci může zcela odrovnat), OpenWRT je pružný systém, který lze aktualizovat, konfigurovat a jakkoliv měnit prakticky za provozu a můžeme při tom používat jak běžný balíčkovací systém, tak při konfiguraci třeba textový editor.

Jak bylo výše zmíněno, OpenWRT používá balíčkovací systém a tedy lze doinstalovat různé balíčky podle potřeby, výběr je veliký. Do distribuce patří samozřejmě firewall NetFilter s přístupovým programem iptables, podporuje ftp, samba, telnet, ssh, lze doinstalovat X Window a nějakého vhodného správce oken,

konfigurujeme obvykle přes ssh (případně lze přes ssh bezpečně provozovat grafické prostředí). Ve výchozím stavu je tato distribuce vybavena typicky pro Wi-fi router, což ale můžeme změnit podle svých vlastních potřeb.

9.4.3 Filtrování

Rozlišujeme různé typy filtrování podle toho, na kterou vrstvu ISO/OSI lze danou metodu zařadit (a tedy na které informace v záhlavích „dosáhneme“). Obvykle se jedná především o filtrování na síťové vrstvě (L3), protože tam lze pracovat s IP adresami. Následující stránky souvisejí s funkcí firewallu nebo mohou být použity samy o sobě (například ACL).

 **Paketový filtr.** Jedná se o nejjednodušší filtrování na L3 a částečně L4, v pravidlech se uvádějí jen IP adresy a čísla portů. Je to jednoduché a rychlé řešení (provoz je zdržován jen minimálně), které se dříve běžně uplatňovalo především na mezilehlých síťových prvcích (například starší verze operačního systému IOS pro směrovače), dnes je najdeme v některých nejlevnějších směrovačích. Nevýhodou je neschopnost nahlížet do komunikace probíhající v složitějších protokolech.

 **ACL na síťové vrstvě.** ACL (Access Control List – seznam řízení přístupu, přístupový seznam) je metoda široce používaná jak v desktopových a serverových operačních systémech, tak i v síťových zařízeních. Účelem je určovat pravidla přístupu pro různé typy provozu.

Přístupový seznam je vlastně seznam položek označených kategorií:

- *deny* (nepustit) – to, co nemá projít,
- *permit* (pustit) – to, co může projít,
- defaultní akce pro provoz, pro který nebylo pravidlo, obvykle *deny all*

ACL je většinou implementován ve formě „co není dovoleno, je zakázáno“, takže najdeme spíše jen položky *permit* (a co není v těchto položkách, to neprojde). Případně se můžeme setkat s celou strukturou navzájem provázaných ACL. Položky se procházejí sekvenčně, jedna po druhé, a hledá se shoda. Také pořadí položek je důležité.

Filtruje se obvykle podle cílové IP adresy, podle zdrojové IP adresy, podle jejich kombinace, a nebo případně podle dalších kritérií (například podle položek v PDU síťové vrstvy – protokolu IP, ICMP, podle TCP/UDP portu, se kterým se navazuje spojení ze síťové vrstvy, apod.). Adresa obvykle bývá adresou podsítě, tedy je uložen prefix a jeho délka (resp. maska, aby bylo zřejmé, u kolika bitů adresy se má hledat shoda).



Další informace:

Ukázky ACL najdeme například na

- http://www.cs.vsb.cz/grygarek/PS/projekt0304/acl_priklad.pdf
- <https://www.samuraj-cz.com/clanek/cisco-ios-18-inter-vlan-routing-a-acl-smerovani-mezivlany/>



 **Stavová inspekce paketů (SPI).** Přesuneme se o vrstvu výše – na transportní vrstvě (L4, konkrétněji teď jsme na L3 a L4) se provádí filtrování *SPI* (Stateful Packet Inspection, stavová inspekce paketů, také stavový paketový filtr). Zatímco na L3 se pakety berou jako „jedináčci“ bez vzájemné vazby, na L4 je možné brát při filtrování v úvahu jejich vzájemné vztahy. například můžeme odlišit pakety, které navazují spojení, od paketů, které patří do již existujícího spojení.

Paket navazující spojení je důkladně prověřen (údaje z vrstev L3 a L4, například zdrojová a cílová adresa, protokoly, zdrojové a cílové porty, příznaky nastavené podle jednotlivých protokolů, cokoliv, co je v záhlavích PDU) a pokud projde, vytvoří se v *stavové tabulce* záznam povolující dané spojení. Pokud paket kontrolou neprojde úspěšně, záznam se nevytvoří.

Pakety, které patří do již navázaného spojení (nebo se za takové vydávají), se filtrují podle toho, zda jejich spojení je zaznamenáno ve stavové tabulce.

V současné době je SPI v podstatě standardem kvalitních firewallů. Oproti běžnému paketovému filtru nabízí větší bezpečnost při relativně malém zpomalení provozu při filtrování.

 **Proxy na aplikační vrstvě.** Posuneme se ještě o vrstvu výše – na aplikační vrstvě TCP/IP pracují proxy firewally (také aplikační brány). Pracují s aplikačními protokoly, například rozumí protokolu HTTP, FTP, IMAP, dokážou pracovat s pakety obsahujícími prvky pro ActiveX, apod. Tento typ firewallu odděluje síť až do té míry, že počítač (server) ve vnější síti nezná IP adresu počítače ve vnitřní síti, se kterým komunikuje.

Pokud takový firewall vidí i do datové části zpráv a poskytuje další analytické funkce, jde o *Next-Generation Firewall* (NGFW). Pokud umí vstupovat i do šifrovaného provozu, hovoříme o *SSL Visibility Appliance*.

Veškerá komunikace je zpracovávána pomocí tzv. *proxy* – softwarových bran buď naprogramovaných pro konkrétní typ komunikace (protokol) s poměrně vysokým stupněm zabezpečení (například pro protokol FTP nebo HTTP), nebo pomocí *generické proxy* použitelné obecně pro různé protokoly. Jedno z nejznámějších řešení generické proxy je open-source projekt stunnel.

Proxy defacto zcela odděluje vnitřní a vnější síť (v podobném smyslu jako NAT) a při filtrování využívá informace prakticky ze všech vrstev TCP/IP, protože při „prokopávání“ k záhlaví protokolů vrstvy L7 prochází přes záhlaví předchozích vrstev. Rozlišujeme dva typy proxy:

1. *Běžný (standardní) proxy* – pro něj platí vše, co bylo dříve o proxy napsáno. Filtruje všechny pakety podle údajů z PDU protokolů aplikační vrstvy a nižších vrstev.
2. *Dynamický proxy* – chová se odlišně k různým paketům; k zahajujícím spojení se chová stejně jako první typ proxy, ale k paketům patřícím do již vytvořeného spojení se chová jako SPI (tj. na nižší vrstvě TCP/IP). Důsledkem je zrychlení odbavování příchozího i odchozího provozu.

9.4.4 IDS/IPS a hloubková inspekce paketů

SPI můžeme brát jako základ pro IDS/IPS systémy.

 *IDS* (Intrusion Detection System, systém pro detekci útoků) je pasivní prvek (tj. provoz neplyne přes toto rozhraní, ale je k němu zrcadlen). Nijak neovlivňuje provoz, pouze ho sleduje a informuje o případných problémech v provozu.

 *IPS* (Intrusion Prevention System) dokáže aktivně reagovat na detekovaný problém. Reakce má útoku zabránit, tedy pokud IPS zjistí počínající útok, tak například může zasáhnout do konfigurace jiných síťových zařízení za účelem odvrácení útoku, zahodit podezřelé pakety nebo je tzv. označit (mark).

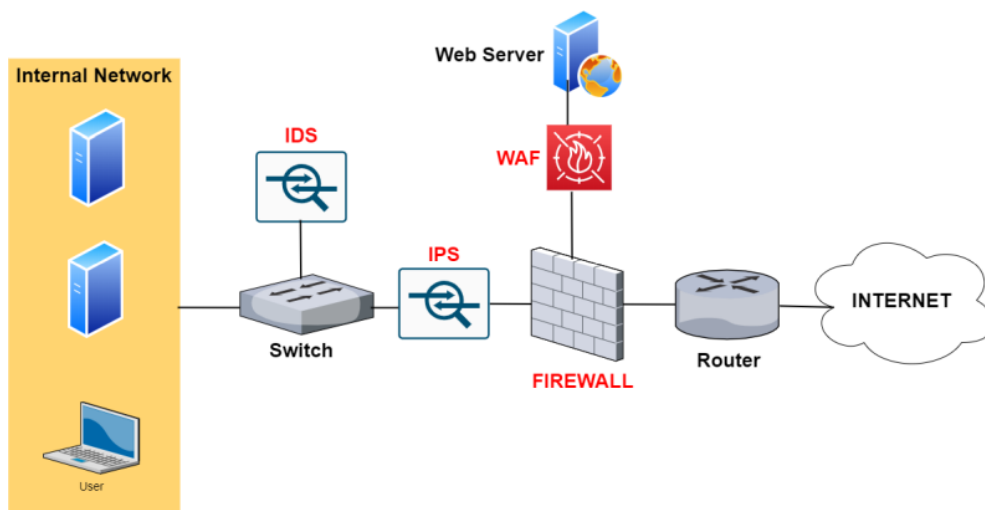
Hlavní problém IPS je, že provoz plyne přes něj (jinak by nemohl zahazovat nebo označovat podezřelé pakety), tedy může zpomalovat síťový provoz či zhoršovat latenci, a to dokonce i v případech, kdy zrovna k žádnému útoku nedochází.

IDS/IPS může být součástí firewallu nebo hardwarový modul v routeru, nebo jde o samostatné zařízení, nebo to může být softwarový modul či aplikace.

 IDS a IPS pracují podobně jako antivirus, tedy používají

- signatury útoků pro rozpoznání známých typů útoků (vede si jejich databázi, kterou je nutné aktualizovat z důvěryhodného zdroje nebo „mít k dispozici v cloudu“),
- heuristiky (využívá statistické metody vyhodnocující provoz na síti) – hledá v provozu na síti podezřelé pakety,
- detekci neobvyklého chování sítě (zjišťují se odchylky od běžného provozu na síti).

 Detekce odchylek od běžného chování sítě se označuje jako *anomaly-based detection*, znamená porovnávání skutečného chování sítě s naučenou běžnou úrovní a typem provozu, což je *baseline*. Baseline je model pro běžné chování sítě či zařízení. Typicky zahrnuje běžné zatížení sítě, frekvenci provozu, obvyklé protokoly v provozu, rozsah vyskytujících se IP adres, dynamiku změn v provozu sítě, atd.



Obrázek 9.6: Schéma možného zapojení IDS/IPS⁶

Na obrázku 9.6 vidíme jedno z možných schémat zapojení firewallů a IDS/IPS ve větší síti. Celá síť je chráněna firewallem, za kterým je IPS. Veškerá komunikace do intranetu prochází IPS. Firewall má tři zóny, z toho jedna je DMZ s webovým serverem chráněným zařízením WAF (Web Application Firewall). Switch zrcadlí veškerý provoz na DMZ.

Poznámka:

Všimněte si, že IDS není „přímo na cestě“, ale vede k němu „odbočka“. Provoz směřující do DMZ (nebo opačně) je zrcadlen (kopírován) na port vedoucí k IDS zařízení. Proč tomu tak je? Protože nechceme, aby zařízení IDS „bylo vidět“. Pokud se útočníkovi podaří vloudit se do sítě a začne skenovat naši síť (zjišťovat, kde co máme, jakou co má adresu apod.), IDS pro něj zůstane neviditelným ze dvou důvodů – pouze detekuje (nedává o sobě vědět žádnými aktivními reakcemi) a zároveň není na cestě k žádnému síťovému prvku. Takže máme v síti skrytý prostředek, který můžeme následně použít proti útočníkovi, resp. sledovat jeho činnost bez toho, aby to útočník zpozoroval.



⁶Zdroj: <https://www.networkstraining.com/firewall-vs-ips-vs-ids-vs-waf/>

 **HIDS, HIPS** (Host-based IDS/IPS) je modul na koncovém zařízení (serveru, počítači apod.) určený k ochraně daného zařízení. Také hovoříme o agent-based systému (protože na zařízení běží software v roli agenta, který má určitou úlohu).

Příklady HIDS/HIPS: OSSEC (Open Source HIDS Security), Safetica, AlienVault USM, Tripwire a jeho open-source varianta AIDE, Cisco AMP. OSSEC používá centrální server pro IDS/IPS management a agenty instalované na jednotlivých zařízeních.

 **NIDS, NIPS** (Network-based IDS/IPS) je implementován jako senzor, přes který prochází síťový provoz (NIPS) nebo je provoz na toto zařízení zrcadlen (IDS).

Příklady NIDS/NIPS: Snort, Squil, Suricata, Fortigate, FortiOS, Zeek (Bro), OpenWIPS-ng. Většina z nich je volně dostupná, kromě Fortigate a FortiOS. Většinou tyto systémy dokážou plnit obě role: IDS i IPS, podle konfigurace.

 Setkáváme se také s názvem *Stavový paketový filtr s hloubkovou kontrolou* (Deep Packet Inspection, DPI), to je právě firewall s integrovaným modulem IDS/IPS. Tento typ firewallu přidává další možnosti definování pravidel související s obvyklými vlastnostmi komunikace s danou (známou) aplikací či protokolem. Pokud například zjistí, že protokol HTTP je používán pro jiný typ komunikace než s WWW serverem, tento požadavek zablokuje jako podezřelý.



Další informace:

- <https://www.networkstraining.com/firewall-vs-ips-vs-ids-vs-waf/>
- <https://community.broadcom.com/symantecenterprise/communities/community-home/librarydocuments/viewdocument?DocumentKey=92f2ea26-cdf3-48a9-85b6-1143a52fee82>
- <https://community.broadcom.com/symantecenterprise/communities/community-home/librarydocuments/viewdocument?DocumentKey=09a07639-026b-4045-ae8d-d0cb1fe42a9b>
- <https://www.systemonline.cz/clanky/systemy-pro-detekci-neopravneneho-pruniku.htm>



Další informace:

V přílohách najdeme sekce o firewallu ve Windows a v Linuxu včetně potřebných příkazů pro konfiguraci. Na straně 252 je sekce o firewallu ve Windows, na straně 280 je sekce o firewallu v Linuxu (velmi podrobná včetně příkladů pravidel).

Další informace o firewallech:

- <https://www.stunnel.org/>
- <https://www.root.cz/serialy/openwrt/>
- http://www.vutbr.cz/www_base/zav_prace_soubor_verejne.php?file_id=15131



Úkol

V příloze od strany 280 je sekce o firewallu *NetFilter*, ve které jsou na příkladech vysvětleny nejdůležitější funkce firewallů. Projděte si dotyčnou sekci a ujasněte si princip filtrování, SPI, překladu adres a dalších



9.5 Dohledové systémy

 *Dohledový systém* (Network Management Software) je systém, který dokáže monitorovat stav sítě (různé veličiny), tedy sbírat informace z různých zdrojů na síti, generovat reporty a v případě problémů (abnormální hodnoty sledovaných veličin, velké výkyvy apod.) vhodně reagovat (například odeslat informaci adminovi).

Když si vybíráme dohledový systém, bereme v úvahu:

- co vše může být monitorováno (množství a typy monitorovaných služeb),
- konkrétnost hlášení (některé poruchy jsou „hierarchické“, tedy pokud něco selže v důsledku selhání něčeho jiného, neměl by report zahrnovat vše, ale pouze „kořen problému“),
- uživatelské rozhraní (dnes je velmi běžné webové rozhraní), mělo by být dostatečně přehledné a vhodně hierarchicky členěné, hodně záleží na tom, co si lze zobrazit a v jaké formě, možnost vizualizovat nejrozličnější datové toky či cokoliv dalšího, co se může s časem měnit,
- způsoby zasílání hlášení – přes SMTP server (příp. může být i integrovaný), SMS,
- možnosti nastavení uživatelských oprávnění pro přístup k evidovaným informacím,
- funkční omezení – způsob komunikace s hlídanými zařízeními (přes který protokol, např. SNMP), způsob logování, atd.

Existují jak open-source, tak i komerční dohledové systémy. Ke komerčním patří například WhatsUp Gold,⁷ SonicWALL,⁸ a další. Na některé open-source produkty se podíváme v následujícím textu.

 Další informace:

- <https://www.networksecuritytoolkit.org/nst/index.html>
- <https://www.networkmanagementsoftware.com/network-monitoring-tools/>
- <https://www.monitortools.com/>



9.5.1 Nagios

Nagios je velmi oblíbený open-source monitorovací (dohledový) systém. Zvládá nejen samotné monitorování, ale i vizualizaci. Dokáže monitorovat různé typy síťových služeb (HTTP, POP3, ICMP, SMTP), nemá problémy s šifrováním, monitoruje také využívání prostředků na uzlech sítě (rozumí si s různými operačními systémy včetně Windows), zvládá vizualizaci stavu sítě (také dokáže vykreslit topologii sítě včetně 3D grafu). Pro konfiguraci lze použít *webové rozhraní* (ale je možné využívat jiná rozhraní, například Centreon v kombinaci s MySQL).

Vznikl nový fork (odnož) Nagiosu, který se nazývá *Icinga*. Podle vývojářů tohoto forku je účelem především zrychlit a zprůžnit vývoj.

 Další informace:

- <https://www.abclinuxu.cz/clanky/site/nagios-plus-centreon-plus-mysql-instalace-a-zakladni-konfigurace>

⁷<http://www.annexnet.cz/ipswitch-whatsupgold-popis-produktu/>

⁸<http://www.sonicwall.com/>

- <https://www.nagiosbook.org/>
- <https://assets.nagios.com/downloads/nagioscore/docs/nagioscore/3/en/quickstart.html>
- <https://icinga.com/>



9.5.2 Zabbix

Zabbix je dalším oblíbeným open-source monitorovacím systémem. Potřebujeme Zabbix Server (ten instalujeme na dohledový server, měl by na něm běžet některý UNIXový systém včetně Linuxu) a Zabbix agenty (na klientských zařízeních, různé operační systémy včetně Windows). Konfigurace se provádí opět přes webové rozhraní, tedy po základní orientaci nic moc složitého.

Funkčnost je podobná, snad mírně nižší, než v případě Nagiosu, obvykle je Zabbix považován za jednodušší ve vybavenosti a konfiguraci (Nagios je zase považován za stabilnější). Taktéž podporuje SNMP a agenty je možné instalovat na různé operační systémy.



Další informace:

- <https://www.root.cz/serialy/dohledovy-system-zabbix/>
- <https://www.zabbix.com/documentation/current/>



9.5.3 OpenNMS

OpenNMS⁹ je open-source dohledový systém, který je narozdíl od předchozích napsán v Javě. Je určen pro monitorování rozsáhlých sítí s velkým množstvím sledovaných zařízení. Monitoring může být distribuovaný (na více serverech), díky tomu lze monitorovat tak velké množství zařízení.

Jde o dobře rozšiřitelný systém. Monitoring různých služeb je řešen pomocí pluginů (podporu monitorování další služby lze přidat jednoduše přidáním pluginu).

Konfigurace se opět provádí přes webové rozhraní. Práci se systémem je možné si vyzkoušet na demo aplikaci na stránkách projektu.



Další informace:

- https://www.howtoforge.com/opennms_network_management
- <https://www.opennms.com/documentation/>



9.5.4 Zenoss

Zenoss¹⁰ je další open-source dohledový systém (jedna varianta je open-source a volně ke stažení, druhá – Enterprise – komerční). Je postaven na myšlence využití existujících open-source projektů.

Jeho jádrem je objektový webový server *Zope* napsaný v jazyce Python, a také v Pythonu je možné Zenoss dále rozšiřovat (dokonce lze údajně používat i pluginy z Nagiosu, který je také napsán v Pythonu).

⁹Informace o OpenNMS najdeme na http://www.opennms.org/wiki/Main_Page.

¹⁰*Zenoss* je dostupný na <http://www.zenoss.com/>.

Dále se přidává databázový systém *MySQL* a celý systém běží nad *Twisted*, což je událostmi řízené rozhraní pro programování síťových aplikací, které si rozumí s mnoha různými protokoly na více vrstvách ISO/OSI.

 Důležitou součástí celého systému je *RRDTool* (Round Robin Database Tool). Tento nástroj slouží k ukládání, zpracovávání a vytváření grafické reprezentace (grafů) jakýchkoliv čísel, která mu předáme. Základem je databáze organizovaná jako kruhová fronta (odtud název Round Robin Database) se statickou velikostí, ve které jsou starší data přepisována novějšími.

Zenoss si rozumí s protokoly SNMP, ICMP, protokoly rodiny TCP/IP, komunikuje s UNIXovými systémy (včetně syslogu) a také s Windows (podporuje také WMI). Na webu firmy najdeme také informaci o tom, že jsou podporována také cloud zařízení (pojem Cloud Computing je velmi populární). Konfigurace probíhá přes přehledné webové rozhraní.



Další informace:

- <https://www.zenoss.com/product>
- <https://www.root.cz/serialy/prechadzame-na-rrdtool/>
- <https://oss.oetiker.ch/rrdtool/tut/>
- <https://help.ubuntu.com/community/Zenoss>



9.5.5 Cacti

Cacti¹¹ je další open-source nástavba nástroje RRDTool, podobně jako Zenoss. Také v Cacti jde o to, jak získat data, vhodně je uložit a následně analyzovat a zobrazit. Data mohou být získávána z různých zdrojů, například od SNMP (Net-SNMP), skriptů v nejrůznějších skriptovacích jazycích nebo WMI, a uložena buď do SQL databáze nebo pomocí RRDTool. Systém je určen pro menší a střední sítě (stovky, možná tisíce, uzlů sítě).

Komunita kolem systému Cacti se utěšeně rozrůstá, a tak existuje hodně pluginů a také šablon pro různé typy hodnot. Stejně jako u předchozích nástrojů, také zde máme k dispozici přehledné webové rozhraní pro konfiguraci a manipulaci s daty.



Další informace:

- <https://www.root.cz/clanky/cacti-vse-dulezite-v-jednom-monitoru/>
- <https://www.ubuntugeek.com/install-and-configure-cacti-monitoring-tool-in-ubuntu-9-10-karmic-server.html>
- <https://www.root.cz/clanky/cacti-ziskavani-vlastnich-dat-pomoci-ssh/>
- <https://www.samuraj-cz.com/clanek/cacti-snmp-monitoring-a-grafy/>
- <https://forums.cacti.net/about30438.html>



¹¹ Cacti najdeme na <http://www.cacti.net/>.

¹² Zdroj: <https://lintut.com/best-free-monitoring-system-for-linux/>

Obrázek 9.7: Dohledový systém Cacti¹²

Úkol

Vyberte si jeden ze zmiňovaných dohledových systémů a najděte o něm co nejvíce informací. Pokud máte možnost, vyzkoušejte ho.



Další informace:

A ještě pár odkazů týkajících se zabezpečení:

- http://www.scycore.com/papers/low_linux_intro.pdf
- <https://www.networksecuritytoolkit.org/nst/index.html>
- <https://www.abclinuxu.cz/clanky/bezpecnost/linuxove-dmz-i>



9.6 SIEM

SIEM (Security Information and Event Management) může svou definicí připomínat dohledový systém, je však více orientován na bezpečnost. Tyto systémy slouží jako nástroje pro shromažďování, analýzu a prezentaci informací získaných z různých napojených zařízení, ať už běžných součástí sítě nebo zařízení bezpečnostního charakteru (kamery, nejrozličnější senzory, IDS/IPS). Obvyklé úlohy:

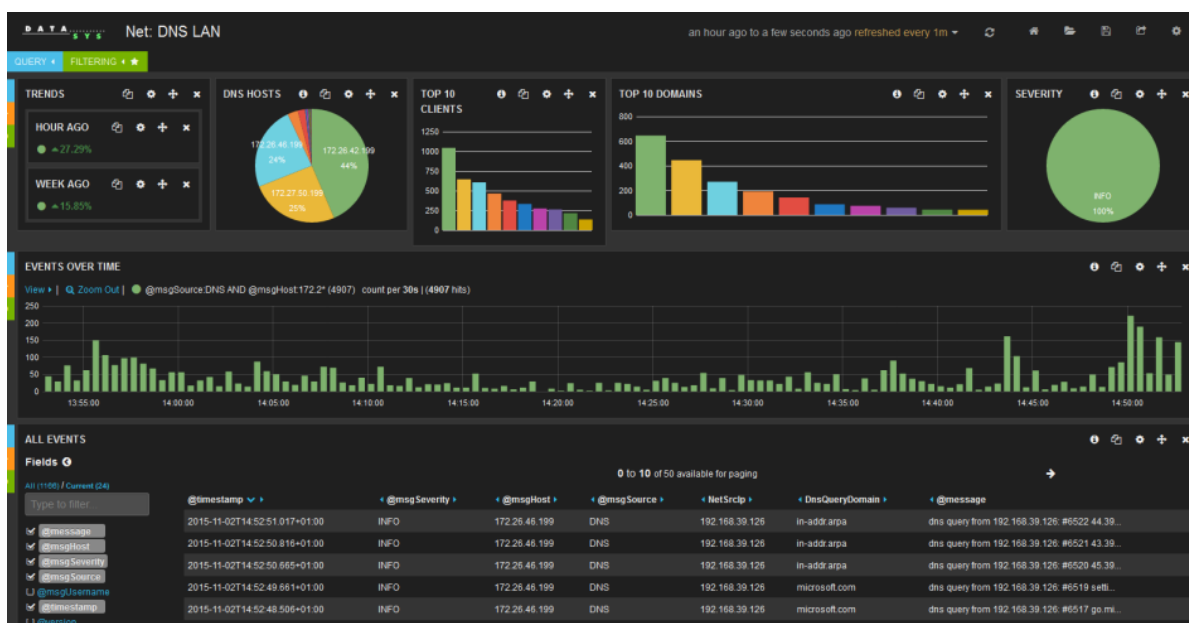
- shromažďování dat z různých zdrojů, práce s logy, normalizace údajů (převod do jednotného formátu),
- korelace – hledání a vyhodnocování vztahů mezi daty, např. mezi různými událostmi v síti,

- okamžité varování při výskytu problému,
- prezentace informací – vytváření přehledů, sestav, grafů, statistik.

Účelem je mít na jednom místě všechny informace potřebné při řešení bezpečnostních incidentů a mít možnost co nejrychleji na ně reagovat.

Můžeme najít buď kompletní SIEM řešení, nebo zvlášť SIM (Security Information Management, tj. práce s logy, třídění informací, reportování, vizualizace apod.) a SEM (Security Event Management, tj. především práce s logy, analýza v reálném čase, hledání závislosti mezi událostmi a rychlá reakce na události). SIEM používají při analýze jak běžné metody z dohledových systémů (jako je práce s logy), tak i metody podobné těm, které používají antivirové programy (práce se signaturami, heuristická analýza apod.).

Nabídka SIEM je poměrně rozsáhlá. Můžeme si vybrat open-source řešení nebo komerční produkt, ovšem vybírat bychom měli především podle vlastních potřeb. Důležitým kritériem je počet zařízení, která v síti máme a na jejichž logy bude SIEM napojen, mělo by nás zajímat, s kterými typy zařízení dokáže SIEM nativně pracovat a naopak která zařízení bude problém napojit, s velikostí sítě souvisí údaj EPS (Events per Second), tedy kolik událostí za sekundu bude muset SIEM zvládat, atd.



Obrázek 9.8: Kibana – rozhraní používané v SIEM Elisa¹³

 Vlastní SIEM nabízí Cisco, z dalších je oblíbený Splunk (má více variant), Helix Security Platform (od FireEye), Qradar SIEM (od IBM), AlienVault OSSIM (od AT&T Cybersecurity), AlienVault USM (od téhož výrobce), FortiSIEM (od Fortinetu), McAfee Enterprise Security Manager, LogRhythm NextGen SIEM Platform, Trustwave SIEM, Rapid7 InsightIDR, Juniper JSA SIEM, nebo český produkt ELISA (od DataSys, staví na Zabbixu, Xlogu a NetFlow).



Další informace:

- <https://www.monitortools.com/>

¹³Zdroj: <https://docplayer.cz/36323951-Datasys-elisa-log-management-rizeny-zabbixem-lukas-maly-dis-it-konzultant-bezpecnost-a-monitoring.html>

- KŘÍŽ, Lukáš. *SIEM: Náročná cesta k pokročilému zajištění kybernetické bezpečnosti* [online], 2019. Dostupné na: https://ictrevue.ihned.cz/c3-66593380-0ICT00_d-66593380-siem-narocna-cesta-k-pokrocilemu-zajisteni-kyberneticke-bezpecnosti
- <https://searchsecurity.techtarget.com/feature/Comparing-the-best-SIEM-systems-on-the-market>
- <https://www.itcentralstation.com/categories/security-information-and-event-management-siem>
- <https://www.comparitech.com/net-admin/siem-tools/>
- <https://www.selecthub.com/siem-solutions/>
- <https://www.softwaretestinghelp.com/siem-tools/>



Souhrn kapitoly

V této kapitole jsme navázali na předchozí kapitolu a více se ponořili do tématu zabezpečení počítačové sítě. Nejdřív jsme si ujasnili některé důležité pojmy z kybernetické bezpečnosti, následoval přehled nejběžnějších kybernetických útoků a informace o organizacích, které se u nás těmito problémy zabývají. Následovala sekce o bezpečnostních zařízeních v síti: síťových analyzátorech, nástrojích na monitorování provozu, firewallech, různých možnostech filtrování provozu a IDS/IPS zařízeních. Dalším probíraným tématem jsou dohledové systémy, které nám mohou zprostředkovat informace získané v rámci monitoringu v síti, a jejich ještě více bezpečnostně zaměřená varianta – SIEM.



Seznam doporučené literatury

- [1] CESNET2. URL: <http://www.cesnet.cz/>
[cit. 20. 5. 2010]
- [2] DONAHUE, G. A. *Kompletní průvodce síťového experta*. Brno, Computer Press, 2009.
- [3] EMPSON, S. *CCNA: Kompletní přehled příkazů. Autorizovaný výukový průvodce*. Brno, Computer Press, 2009.
- [4] HARPER, A. et al. *Hacking – manuál hackera*. Praha, Grada, 2008.
- [5] HORÁK, J. — KERŠLÁGER, M. *Počítačové sítě pro začínající správce*. Brno, Computer Press, 2008.
- [6] HUBERT, B. et al. *Linux Advanced routing & Traffic Control* [online].
URL: <http://lartc.org/howto/>, další na <http://lartc.org/>
[cit. 1. 7. 2011]
- [7] JANEČEK, J. *Distribuované systémy* [online].
URL: https://dsn.felk.cvut.cz/wiki/_media/vyuka/cviceni/x36pko/skripta_dsy.pdf
[cit. 20. 5. 2010]
- [8] JONES, D. *Automatizace správy a skriptování Microsoft Windows*. Brno, Computer Press, 2006.
- [9] *Linux Home Networking* [online]. Knihy o správě počítačových sítí v Linuxu.
URL: <http://www.linuxhomenetworking.com/wiki/index.php>
[cit. 1. 7. 2011]
- [10] MINASI, M. — YORK, D. *Linux pro administrátory Windows*. Brno, Computer Press, 2004.
- [11] OULEHLA, Milan a Roman JAŠEK. *Moderní kryptografie*. Praha: IFP Publishing, 2017. ISBN 978-80-87383-67-4.
- [12] PETERKA, J. *eArchiv: Co je čím v počítačových sítích* [online].
URL: http://www.earchiv.cz/i_coje.php3
[cit. 20. 5. 2010]

- [13] PETERKA, J. *eArchiv: Principy počítačových sítí* [online]. URL: http://www.earchiv.cz/i_pri.php3 [cit. 20. 5. 2010]
- [14] PRICE, B. *Active Directory*. Brno, Computer Press, 2005.
- [15] PUŽMANOVÁ, R. *Moderní komunikační sítě od A do Z*. 2. aktualizované vydání. Brno, Computer Press, 2006.
- [16] PUŽMANOVÁ, R. *TCP/IP v kostce*. 2., upr. a rozš. vyd. České Budějovice: Kopp, 2009. ISBN 978-80-7232-388-3.
- [17] RUEST, D. — RUEST, N. *Virtualizace: Podrobný průvodce*. Brno, Computer Press, 2010.
- [18] SATRAPA, P. *Internetový protokol IPv6*. Praha, CZ.NIC, z.s.p.o., 2008. 358 stran. Kniha je dostupná v elektronické formě na http://knihy.nic.cz/files/nic/edice/pavel_satrapa_ipv6_2008.pdf [cit. 1. 7. 2011]
- [19] SCOTT, C. — WOLFE, P. — ERWIN, M. *Virtual Private Networks*. 2nd edition. O'Reilly, USA, 1999. Náhled většiny stránek je dostupný na Google Books.
- [20] SCHRODER, C. *Linux: Kuchařka administrátora sítě*. Brno, Computer Press, 2009. Z originálu Linux Networking Cookbook vydaného nakladatelstvím O'Reilly Media.
- [21] *Svět sítí, tutoriály* [online]. URL: <http://www.svetsiti.cz/tutorialy.asp> [cit. 20. 5. 2010]
- [22] *System × Virtualization Strategies Development Space* [online]. IBM Redbooks. URL: <http://www-01.ibm.com/redbooks/community/display/REDP4480/Home> [cit. 20. 5. 2010]
- [23] THOMAS, T.M. *Zabezpečení počítačových sítí*. Autorizovaný průvodce. Brno, Computer Press, 2005.
- [24] TRULOVE, J.. *Sítě LAN: hardware, instalace a zapojení*. Praha: Grada, 2009, 384 s. ISBN 978-80-247-2098-2.
- [25] *Výukové materiály Cisco* [online]. URL: http://www.cisco.com/en/US/docs/internetworking/technology/handbook/ito_doc.html [cit. 20. 5. 2010]
- [26] WALLACE, Kevin. *CCNP routing and switching route 300–101: official cert guide*. Second Printing. Indianapolis, IN: Cisco Press, 2015. ISBN 978-1-58720-559-0.

Přílohy

Práce s adresami a sítěmi ve Windows

V této příloze jsou popsány a demonstrovány příkazy, které se používají ve Windows při práci se sítěmi, Linuxu se věnujeme v další příloze. Můžeme je zatím brát jako inspiraci pro procvičování a taky jako předeihu k zatím neexistujícím cvičením z tohoto předmětu.

A.1 Problémy při používání příkazů ve Windows

V serverových Windows podobně jako v desktopových variantách je možné pro téměř vše použít nástroje grafického rozhraní. Velmi často je však praktičtější použít Příkazový řádek – příkazy někdy umějí více než nástroje grafického režimu, lze je snadněji používat pro vzdálenou správu a hlavně příkazy je možné skriptovat a tedy automatizovat zpracování některých úloh.

Ve Windows je s příkazy pro Příkazový řádek podobný problém jako v grafickém rozhraní: různé verze Windows se liší vybaveností a některé příkazy pracují v různých verzích odlišně. Komplikace se projevují především v heterogenních sítích, ve kterých jsou uzly s různými verzemi a případně edicemi Windows (když píšeme skript, musíme dbát na to, aby byl použitelný na všech uzlech, kde ho potřebujeme spouštět).

Jistou komplikací je také změna funkčnosti *RunAs mechanismu* – od verze Vista sice pořád existuje příkaz **runas** pro spuštění procesu s odlišnými přístupovými oprávněními, ale je problematický. Pokud chceme ve Windows od Visty výše spustit příkaz s jinými (typicky administrátorskými) oprávněními než pod kterými právě pracujeme, zvolíme v kontextovém menu ikony programu Příkazového řádku volbu „Spustit jako správce“ a tedy spustíme s vyššími oprávněními už program **cmd.exe**. Svým způsobem to může být bezpečnostní problém, protože někteří uživatelé zřejmě budou na Příkazovém řádku s vyššími oprávněními pracovat i tehdy, když to není nutné.

Pokud chceme některé úlohy provádět vzdáleně, také můžeme narazit na přístupová oprávnění. Obvykle pomůže používat doménový správcovský účet, který lze využít na všech zařízeních v doméně.

A.2 Soubory související se sítěmi

 Se sítí souvisí využívání některých (textových) souborů – naštěstí není vše v registru, některé z níže uvedených souborů budou zmíněny u popisovaných příkazů.

Především z důvodu kompatibility se síťovými protokoly a ne-windows síťovými klienty existují ve složce `...\system32\drivers\etc` tyto soubory:

- **networks** – obsahuje doménové a IP adresy lokálních sítí,
- **hosts** – urychluje mapování IP adres na známé doménové adresy (hostitele, anglicky *hosts*),
- **services** – informace o známých síťových službách,
- **protocol** – informace o známých síťových protokolech.

V 64bitových Windows je třeba hledat pomocí 64bitového správce souborů, případně Průzkumníka; v 32bitové aplikaci se totiž některé složky (včetně této) nezobrazují.

Vzpomeňte si, že v UNIXových systémech jsou v adresáři `/etc` nejružnější konfigurační soubory včetně síťových, zde si Microsoft bral inspiraci.



Úkol

Prohlédněte si obsah souborů **networks**, **hosts**, **services** a **protocol** zmíněných v této sekci.



A.3 Práce s adresami a síťovými kartami

A.3.1 Základní práce s IP adresami – ipconfig



Pro zobrazení informací o adresách můžeme použít příkaz **ipconfig**.

ipconfig /? (nebo **ipconfig -?**) zobrazí seznam parametrů

ipconfig /all vypíše podrobné informace (IP a MAC adresu, adresu brány, masku podsítě)

ipconfig /release název_karty uvolnění IP adresy přidělené zadané síťové kartě (když neuvedeme název karty, jsou uvolněny IP adresy všech karet, je možné také použít hvězdičkovou konvenci)

ipconfig /renew název_karty obnovení přidělení IP adresy pro zadanou síťovou kartu (když není název karty uveden, provede se pro všechny dostupné karty)

ipconfig /displaydns zobrazí záznamy adres přidělených v systému DNS (informace k příslušným DNS jménům včetně hodnoty TTL), a to jak ve směru doménové jméno → IP adresa, tak i ve směru opačném (reverzní adresy, ty jsou v záhlaví označeny řetězcem „in-addr.arpa“); vždy jsou zobrazeny alespoň dva záznamy (localhost a jeho reverzní adresa)

ipconfig /flushdns vymaže DNS cache (tj. dotazy na doménové adresy, které takto z cache smažeme, se musejí provádět znovu), použijeme, pokud DNS překlad funguje nekorektně (máme v cache chybné či neaktuální záznamy)



Příklad

Výstup příkazu může vypadat například takto (záleží na verzi Windows, nejdřív ve Windows XP):

```
C:\Windows> ipconfig /all
```

Konfigurace protokolu IP systému Windows

```
Název hostitele . . . . . : PCPRACE
Primární přípona DNS. . . . . :
Typ uzlu . . . . . : neznámý
Povoleno směrování IP . . . . . : Ne
```



```
WINS Proxy povoleno . . . . . : Ne
```

Adaptér sítě Ethernet Připojení k místní síti:

```
Přípona DNS podle připojení . . . . :
Popis . . . . . : Broadcom NetXtreme Gigabit Ethernet
Fyzická Adresa. . . . . : 00-0F-FE-12-34-56
Protokol DHCP povolen . . . . . : Ne
Adresa IP . . . . . : 193.84.195.15
Maska podsítě . . . . . : 255.255.255.128
Výchozí brána . . . . . : 193.84.195.1
Servery DNS . . . . . : 193.84.192.10
```

Výstup téhož příkazu ve Windows 7 je hodně dlouhý, vyjmeleme pouze část výstupu odpovídající jedné ethernetové kartě:

...

Adaptér sítě Ethernet Připojení k místní síti:

```
Přípona DNS podle připojení . . . : fpf.slu.cz
Popis . . . . . : Atheros AR8131 PCI-E Gigabit Ethernet
Controller (NDIS 6.20)
Fyzická Adresa. . . . . : E0-CB-4E-45-67-89
Protokol DHCP povolen . . . . . : Ano
Automatická konfigurace povolena : Ano
Místní IPv6 adresa v rámci propojení . . . : fe80::ec8b:a3a6:68bc:9b2f%11
(Preferované)
Adresa IPv4 . . . . . : 193.84.194.100(Preferované)
Maska podsítě . . . . . : 255.255.255.128
Zapůjčeno . . . . . : 9. června 2011 15:25:14
Zápůjčka vyprší . . . . . : 9. června 2011 21:25:14
Výchozí brána . . . . . : 193.84.194.129
Server DHCP . . . . . : 193.84.192.10
IAID DHCPv6 . . . . . : 248613215
DUID klienta DHCPv6. . . . . : 00-01-00-01-12-F4-7B-66-E0-CB-4E-45-67-89

Servery DNS . . . . . : 193.84.192.10
Rozhraní NetBios nad protokolem TCP/IP. . . . . : Povoleno
...
```

Řetězec DUID (DHCP Unique Identifier) jednoznačně identifikuje klienta serveru DHCPv6 a naopak, je jedinečný pro danou konverzaci s DHCP serverem při stanovení konfiguračních parametrů. Všimněte si vazby na MAC adresu.

Číslo IAID (Identity Association Identifier) jednoznačně identifikuje sadu adres přidělených DHCP klientovi. Klient může mít k danému síťovému rozhraní více sad IPv6 adres (plus další potřebné adresy, například musí znát adresu DNS serveru), každá taková sada musí mít jiné IAID.



Příklad

Pokud se nedaří navázat síťové spojení nebo spojení nefunguje korektně (například tehdy, když naše síťová karta omylem získala IP adresu, která je již přidělena jinému zařízení, a nebo nebyly korektně získány adresy DNS serverů), může pomoci znovuzískání IP adresy a souvisejících informací (adresa brány, adresy DNS serverů). Použijeme postupně tyto příkazy:

`ipconfig /flushdns` pročistíme DNS cache, tím se zbavíme starších potenciálně chybných záznamů (tento krok nemusí být nutný)

`ipconfig /release` uvolnění IP adresy přidělené všem síťovým kartám, ale často nebývá nutné

`ipconfig /renew` obnovení přidělení IP adresy všech síťových karet

`ipconfig /all` ověříme si, zda je vše v pořádku

Tento postup samozřejmě nebudeme používat, pokud máme pevnou (statickou) IP adresu. Kdybychom se o to pokusili, obdržíme chybové hlášení sdělující, že daný adaptér (tj. síťová karta) není ve službě DHCP povolen (proto nemůže žádat na DHCP serveru jinou adresu).

Když budeme chtít takto obnovit adresu pro jedinou konkrétní kartu, zadáme v příkazech název (můžeme používat i hvězdičkovou konvenci pro zkrácení názvu), například

`ipconfig /release Připojení*1`

(pokud máme síťovou kartu s názvem „Připojení k místní síti 1“).



A.3.2 Překlad adres – arp a nslookup

 *ARP* je protokol převádějící IP adresy na fyzické (MAC) adresy. Fyzická adresa je v síti zjišťována odesíláním tzv. ARP paketů s dotazy. Aby nebylo nutné opakovaně posílat tyto pakety, udržuje každé síťové zařízení (tj. také zvlášť každá síťová karta) mezipaměť, do které ukládá dosud zjištěné dvojice IP a fyzických adres (ARP tabulky). ARP tabulku je možné prohlížet a přidávat i mazat záznamy. Ukážeme si využití příkazu `arp`, který slouží právě k práci s ARP tabulkou.

`arp -a` zobrazí ARP tabulku, také vidíme, které záznamy jsou dynamické (automaticky vytvořené z komunikace) a které statické (ručně vložené)

`arp -a -n 10.0.128.10` zobrazí ARP tabulku pouze pro síťovou kartu, která má přidělenou zadanou IP adresu

`arp -s 123.123.123.123 00-12-34-56-78-9A` přidá se do ARP tabulky statický záznam se vztahem zadané IP adresy a MAC (fyzické) adresy, MAC adresa se zadává v hexadecimálních číslech s pomlčkami

`arp -d 123.123.123.123` odstraní z ARP tabulky záznam pro zadanou IP adresu

`arp -d *` odstraní z ARP tabulky všechny záznamy

Lze přidat parametr s určením IP adresy síťové karty, pro kterou daná ARP tabulka platí (každá karta má svou tabulku), to má smysl jen v případě, kdy je v provozu víc než jedna síťová karta.



Příklad

ARP tabulka může vypadat takto (na Windows 7):

C:\Windows> arp -a

Rozhraní: 193.84.194.100 -- 0xb

internetová adresa	fyzická adresa	typ
193.84.194.129	00-0d-66-22-11-00	dynamická
193.84.194.162	00-1e-4f-33-22-11	dynamická
193.84.194.185	00-1e-4f-dd-bb-00	dynamická
193.84.194.186	00-26-2d-ff-ee-11	dynamická
193.84.194.188	00-1f-d0-66-77-88	dynamická
193.84.194.212	00-1e-4f-aa-11-bb	dynamická
193.84.194.220	00-1e-4f-dd-cc-bb	dynamická
193.84.194.227	00-24-be-77-66-55	dynamická
193.84.194.229	00-1e-4f-ee-dd-44	dynamická

193.84.194.255	ff-ff-ff-ff-ff-ff	statická
224.0.0.2	01-00-5e-00-00-02	statická
224.0.0.22	01-00-5e-00-00-16	statická
224.0.0.252	01-00-5e-00-00-fc	statická
239.255.255.250	01-00-5e-7f-ff-fa	statická
255.255.255.255	ff-ff-ff-ff-ff-ff	statická
...		

Ve Windows nižších verzí je ARP tabulka výrazně kratší. V tabulce vidíme dynamické i statické záznamy, u všech je IP a MAC adresa. Všimněte si skupinových a všesměrových IP a MAC adres (ke skupinové IP adrese patří skupinová MAC adresa a naopak).



Příklad

Určitě jste si všimli, že ve výstupu příkazu `arp` jsou pouze IPv4 adresy. Pokud je třeba vypsát IPv6 sousedy, musíme použít příkaz `netsh`, vypisujeme tabulku protokolu NDP:

```
netsh interface ipv6 show neighbors
```



Zatímco protokol ARP slouží k překladu IP adresy na MAC adresu, protokol DNS slouží k překladu doménové adresy na IP adresu. Pro základní přístup k tomuto překladu lze i na klientských stanicích využít příkaz `nslookup`. Příkaz lze využívat běžným způsobem i interaktivně.

`nslookup www.seznam.cz` vypíše se nejdříve adresa DNS serveru, který poskytl odpověď, a pak samotná odpověď (IP adresy a aliasy pro zadaný doménový název)

`nslookup 77.76.75.3` reverzní překlad, takto zjistíme doménovou adresu k této IP adrese

`nslookup` přechod do interaktivního režimu, tento režim ukončíme příkazem `exit`

Plnou nápovědu k tomuto příkazu získáme pouze v interaktivním režimu: tedy výše uvedeným příkazem přejdeme do interaktivního režimu a pak zadáme `?` nebo `help`.



Příklad

Ukážeme si práci v interaktivním režimu:

`nslookup` spustíme program `nslookup`, prompt se změní na symbol `>`

`?` zobrazíme nápovědu, je značně podrobnější než `nslookup /?` vně interaktivního režimu; také lze použít vnitřní příkaz `help`

`www.google.com` napíšeme doménovou adresu, získáme IP adresu, výstup (v různých verzích Windows se ve výstupech setkáme s více či méně kostrbatou češtinou):

```
Server:  decsu.fpf.slu.cz
Address:  193.84.192.10

Neautorizovana odpoved:
Název:    www.l.google.com
Addresses: 209.85.149.147
          209.85.149.99
          209.85.149.103
          209.85.149.104
          209.85.149.105
          209.85.149.106
Aliases:  www.google.com
```

`set all` tato forma příkazu `set` nic nenastavuje, ale informuje nás o momentálním nastavení pro poslední zadanou adresu (v našem případě pro `www.google.com`; pokud jsme předtím nic nepřekládali, pak obecně platná nastavení), výstup:

```
Vychozí server:  decsu.fpf.slu.cz
Address:  193.84.192.10
```

```
Hostitel =  www.l.google.com
Addresses:  209.85.149.147
            209.85.149.99
            209.85.149.103
            209.85.149.104
            209.85.149.105
            209.85.149.106
Aliases:  www.google.com
```

```
Nastavení možnosti:
nodebug
defname
search
recurse
nod2
novc
noignoretc
port=53
typ=A+AAAA
trida=IN
doba vyprseni casoveho limitu=2
obnoveni=1
koren=A.ROOT-SERVERS.NET.
domena=fpf.slu.cz
MSxfr
verze IXFRversion=1
srchlist=fpf.slu.cz/slu.cz
```

Takže se například dozvíme, že se použil rekurzivní dotaz, jedná se o záznam adresy pro IPv4 i IPv6 (typ `A+AAAA`), atd.

`set norecurse` od této chvíle bude požadován iterativní (nerekurzivní) typ dotazů, zpět nastavíme příkazem `set recurse`

`ls fpf.slu.cz` chceme vypsat adresy (počítače) v zadané doméně, výpis je celkem dlouhý

`ls -t ns fpf.slu.cz` vypíší se pouze doménové servery (typ je `NS`, tedy name server)

`ls -t aaaa sltu.cz` vypíší se zařízení s IPv6 adresou (tj. záznamy typu `AAAA`) v zadané doméně



A.3.3 Směrování – route

Pokud posíláme data na určitou IP adresu, pakety obvykle (ne vždy) procházejí přes bránu do Internetu či jiné sítě. Proto každý uzel sítě včetně klientských počítačů zná adresu brány a případně další směrovací informace (tj. kterým směrem poslat data do určité (pod)sítě), vede si svou směrovací tabulku.

 Ve Windows se s touto tabulkou pracuje příkazem `route`.

`route` (nebo s jakýmkoliv nesprávným parametrem) zobrazí nápovědu k příkazu

`route print` vypíše směrovací tabulku, v nejnovějších verzích Windows vlastně dvě tabulky – zvlášť pro IPv4 a IPv6; na začátku výpisu také máme seznam všech síťových rozhraní včetně virtuálních, tento seznam může být také užitečný

`route print 193.*` vypíše pouze ty cíle ze směrovací tabulky, které odpovídají zadanému výrazu (můžeme používat hvězdičkovou konvenci, také symbol ? pro jeden znak)

`route add 193.84.197.0 mask 255.255.255.128 193.84.195.8 metric 4` do tabulky přidá statický záznam určující, že k síti s IP adresou 193.84.195.8 a maskou 255.255.255.128 se dá dostat přes bránu s IP adresou 193.84.195.63 a metrika (cena trasy) je 8

`route -p add 193.84.197.0 mask 255.255.255.128 193.84.195.8 metric 4` totéž jako předchozí, ale záznam zůstane v tabulce i po restartu systému (bez přepínače `-p` by byl záznam platný jen do konce činnosti systému)

`route change 193.84.197.0 mask 255.255.255.128 193.84.195.9 metric 3 if 2` změna existující položky tabulky (změna platí i po restartu), změnili jsme adresu brány, metriku a navíc jsme nastavili směrování na síťové rozhraní (karta) číslo 2, k danému cíli zadáváme vždy jen ty položky, které chceme změnit

`route del 193.84.197.0` odstraníme záznam z tabulky (lze použít i hvězdičkovou konvenci)

Záznamy ve směrovací tabulce lze tímto příkazem také mazat a měnit. Když máme více síťových karet, lze v příkazu pro přidání do tabulky použít i parametr určující síťové rozhraní.



Příklad

Příkaz `route print` ve Windows XP vypíše tento výstup:

```
=====
Seznam rozhraní
0x1 ..... MS TCP Loopback interface
0x2 ...00 0f fe 12 34 56 ..... Broadcom NetXtreme Gigabit Ethernet - Packet
Scheduler Miniport
=====
Aktivní směrování:
      Cíl v síti      Síťová maska      Brána      Rozhraní Metrika
      0.0.0.0         0.0.0.0         193.84.195.1 193.84.195.15    10
      127.0.0.0       255.0.0.0       127.0.0.1    127.0.0.1        1
      193.84.195.0    255.255.255.128 193.84.195.15 193.84.195.15    10
      193.84.195.15   255.255.255.255 127.0.0.1    127.0.0.1        10
      193.84.195.255  255.255.255.255 193.84.195.15 193.84.195.15    10
      224.0.0.0       240.0.0.0       193.84.195.15 193.84.195.15    10
      255.255.255.255 255.255.255.255 193.84.195.15 193.84.195.15    1
Výchozí brána:      193.84.195.1
=====
Trvalé trasy:
Žádné
```

V seznamu síťových rozhraní jsou pouze dvě, jedno je loopback (místní smyčka), druhé je ethernetová síťová karta. Ve vyšších verzích bychom zde viděli více záznamů (reálné síťové karty ethernetové, Wi-fi, Bluetooth, ale i virtuální zařízení), a také virtuální počítače si instalují vlastní virtuální síťové karty. Například pokud používáme virtuální počítač od VMWare, najdeme v seznamu rozhraní podobné záznamy:

```
20...00 50 56 c0 00 01 .....VMware Virtual Ethernet Adapter for VMnet1
21...00 50 56 c0 00 08 .....VMware Virtual Ethernet Adapter for VMnet8
```

Výpis byl pořízen na stroji, jehož síťová karta (jediná) má přiřazenu IP adresu 193.84.195.15, všimněte si, kde všude se tato adresa v tabulce vyskytuje (například u místních a multicast směrování).

Z výpisu vidíme, že žádné statické (trvalé) trasy nejsou definovány. Všimněte si, že poslední řádek tabulky obsahuje informaci o výchozí bráně – co není směrováno podle výše uvedených pravidel, to odchází na zadanou výchozí bránu.



Pokud zadáme síťové zařízení či název sítě symbolickým názvem, příkaz tento název musí přeložit na IP adresu. K tomu využije obsah souboru

- `...\windows\system32\drivers\etc\networks`, pokud jde o název cíle směrování (obvykle to bývá síť či podsíť),
- `...\windows\system32\drivers\etc\hosts`, pokud jde o bránu (brána je konkrétní uzel v síti, tedy hostitel protokolu, anglicky *host*).

Do těchto souborů můžeme přidávat vlastní symbolické názvy (je to celkem jednoduché, soubory jsou okomentovány). V souboru `hosts` je vždy nejméně jeden název uzlu, a to *localhost*.



Úkoly

1. Zjistěte svou IP adresu, masku, adresu brány, adresu DNS serveru.
2. Máte statickou IP adresu nebo ji získáváte dynamicky od DHCP serveru? Kde to zjistíte? Pokud máte dynamickou adresu, uvolněte ji (deaktivujte síťovou kartu) a získejte novou.
3. Zobrazte ARP tabulku.
4. Zjistěte IP adresu serveru `extrahardware.cz`.
5. Zjistěte adresu svého DNS serveru, zda při komunikaci s DNS serverem používáte rekurzivní nebo iterativní dotazy, na kterém portu komunikujete, v jaké jste doméně. Vypište seznam uzlů ve vaší doméně.
6. Vypište směrovací tabulku na vašem počítači.



A.3.4 Malá síť a skupina (workgroup)

Na systémech s Windows se setkáváme také se zkratkou *NBT* (nebo NetBT, NetBIOS over TCP/IP). Samotný NetBIOS funguje defacto jako jmenná, spojovací a komunikační služba v malých sítích s Windows (skupiny počítačů v síti typu peer-to-peer). Umožňuje ve skupinách sdílet různé prostředky (složky, tiskárny). Názvy počítačů, které zadáváme v grafickém rozhraní Windows, jsou vlastně názvy podle protokolu NetBIOS.

 Aby bylo možné názvy NetBIOSu používat i mimo malé sítě, pracuje dnes NetBIOS nad protokoly TCP/IP ve formě NBT. Pro práci s názvy v NBT můžeme využít příkazy `hostname` a `nbtstat`:

`hostname` zobrazí jméno našeho počítače v NBT

`nbtstat` zobrazí nápovědu příkazu `nbtstat`

`nbtstat -n` zjistíme všechny názvy podle protokolu NBT, které souvisejí s tímto zařízením (tj. obvykle jen název našeho počítače, pokud počítač patří do skupiny, tak i název této skupiny), jestliže máme

více síťových karet (tj. více rozhraní), pak se zobrazí zvlášť tabulky pro všechna aktivní rozhraní (která mají přidělenou IP adresu), i když u všech najdeme stejný NBT název

`nbtstat -c` vypíše se seznam ostatních počítačů ve skupině (NetBIOS názvy a IP adresy)

`nbtstat -a počítač` totéž, ale pro jiný (zadaný) počítač patřící do naší skupiny, zadáváme NetBIOS jméno

`nbtstat -A IP_adresa` totéž jako předchozí, ale počítač zadáváme jeho IP adresou

Také můžeme zobrazit tabulku relací s daným uzlem v síti (skupině).



Příklad

Předpokládejme, že máme dva počítače připojené ve skupině *mojeskupina*:

- *pcnotebook* s IP adresou 10.0.0.1,
- *pcdesktop* s IP adresou 10.0.0.3,

Na prvním z nich nejdřív vypíšeme tabulku místních NBT názvů:

`C:\Windows> nbtstat -n`

Připojení k místní síti:

Adresa IP uzlu: [10.0.0.1] ID oboru: []

Tabulka místních názvů systému NetBIOS

Název	Typ	Stav
PCNOTEBOOK	<00> UNIQUE	Registrovaný
MOJESKUPINA	<00> GROUP	Registrovaný
PCNOTEBOOK	<20> UNIQUE	Registrovaný

Teď na tomtéž počítači vypíšeme seznam dostupných počítačů:

`C:\Windows> nbtstat -c`

Připojení k místní síti:

Adresa IP uzlu: [10.0.0.1] ID oboru: []

Název vzdálené paměti NetBIOS Tabulka

Název	Typ	Adr. hostitele	Životnost [sek]
PCDESKTOP	<20> UNIQUE	10.0.0.3	92

Pokud dostaneme jen chybovou hlášku „V mezipaměti nejsou žádné názvy“, důvodem může být to, že počítače nejsou ve stejné skupině (tj. nevidí se), ale také se tento problém může/nemusí objevit, pokud na těchto počítačích máme různé verze Windows.

Můžeme také vypsat NBT tabulku druhého počítače (jsme pořád na tomtéž počítači):

`C:\Windows> nbtstat -a pcdesktop`

Připojení k místní síti:

Adresa IP uzlu: [10.0.0.1] ID oboru: []

Tabulka názvů vzdálených počítačů NetBIOS

Název	Typ	Stav
-------	-----	------

```
-----
PCDESKTOP      <00>  UNIQUE      Registrovaný
MOJESKUPINA    <00>  GROUP       Registrovaný
PCDESKTOP      <20>  UNIQUE      Registrovaný
```

Adresa MAC = 00-1D-0F-12-34-56

Jiný přepínač použijeme, pokud příkaz zadáme s využitím IP adresy:

```
nbtstat -A 10.0.0.3
```

Výstup by byl tentýž jako v předchozím případě. Příkaz bohužel nepřijímá regulární výrazy ani hvězdičkovou konvenci, tedy případnou automatizaci (například vypsání údajů pro více různých adres) bychom mohli řešit například příkazem `for` (viz skripta pro Windows z předmětu Operační systémy).



Úkol

Pokud jste v LAN s více počítači, které jsou ve stejné skupině, vyzkoušejte použití příkazu `nbtstat` podle ukázek v příkladech. Pokud ne, vyzkoušejte alespoň výpis NBT názvů na vlastním počítači.



A.4 Testování a statistiky

A.4.1 Testování cesty a průchodnosti

 To, zda je dostupná doménová nebo IP adresa (resp. příslušné zařízení na síti), zjistíme příkazem `ping`. Tento příkaz odesílá k danému zařízení ICMP pakety (zprávy *ICMP Echo Request*) a podle zaslaných odpovědí určuje, zda je počítač dostupný a jaká je odezva připojení (pokud má vzdálený počítač firewall nakonfigurovaný tak, aby požadavky `ping` byly ignorovány, může se počítač jevit jako nedostupný)

`ping www.google.com` zjistí, zda je zadaný server dostupný (je dostupný prakticky vždy, tedy pokud se zobrazí hláška, že tomu tak není, zřejmě nejsme připojeni k síti nebo je někde na cestě porucha)

`ping -n 2 www.google.com` tento parametr omezí (nebo u většího čísla navýší) počet odesílaných testovacích paketů, zde na 2; výchozí hodnota je 4 pakety

`ping -t www.google.com` počet odesílaných paketů není stanoven, posílají se opakovaně až do přerušení klávesovou zkratkou 



Příklad

Příkaz `ping` můžeme použít i pro otestování funkčnosti vlastní síťové karty:

```
ping loopback    nebo
```

```
ping 127.0.0.1
```



Příklad

Jestliže se síťové spojení podle parametrů zdá v pořádku, ale ve skutečnosti nefunguje, ověříme, zda je funkční DNS překlad adres (doménových adres na IP adresy):

`ping www.google.com` skončí chybovým hlášením, když spojení nefunguje (můžeme použít adresu jakéhokoli serveru na Internetu, u kterého máme jistotu, že odpovídá na žádosti příkazu `ping`)

`ping 74.125.79.104` provedeme totéž, ale použijeme IP adresu (zde se jedná o IP adresu serveru z předchozího příkazu)

Pokud první příkaz selhal a druhý byl úspěšný, pravděpodobně nefunguje DNS server. Měli bychom tedy reagovat buď zjištěním správné adresy (například u svého poskytovatele připojení) nebo použitím adresy veřejného DNS serveru.

 Takových serverů již existuje dostatek, například Google má veřejné DNS servery na adresách 8.8.8.8 a 8.8.4.4, případně můžeme použít DNS servery sdružení CZ.NIC 217.31.204.130 nebo 217.31.204.131. Ve Windows se adresy DNS serverů nastavují v grafickém režimu, v Síťových připojeních (vlastnostech protokolu TCP/IP).

 Návod najdeme například na adrese

<http://www.nic.cz/page/847/jak-nastavit-verejne-dns-servery-cz.nic/>.



 Program `tracert` používáme, když chceme znát cestu, po které procházejí naše pakety (tj. adresy směrovačů na cestě). Ve Windows se používá metoda ICMP paketů posílaných v IP paketech s postupně se zvyšující hodnotou TTL.

`tracert -?` vypíše nápovědu (pozor, UNIXová syntaxe, přepínače píšeme s pomlčkou)

`tracert www.google.com` chceme trasu k zadanému serveru, postupně (celkem pomalu) se vypíší všechny směrovače na cestě včetně odezvy v ms počítané od předchozího uzlu

Můžeme zadat také IP adresu. Metoda ICMP paketů je poměrně pomalá. Navíc některé směrovače jsou z bezpečnostních důvodů nastaveny tak, aby neodpovídaly na ICMP pakety (pak od dotyčného uzlu obdržíme informaci o vypršení časového limitu žádosti), také je obvyklé, že se takto nedostaneme přímo ke konkrétnímu serveru, ale jen ke směrovači na okraji sítě, ve které je daný server.

 Program `pathping` slouží k podobným účelům, také dokáže vypsát trasu k zadanému serveru (jednotlivé směrovače) a navíc pro každý tento úsek vypočítává statistiku stejnou jako příkaz `ping`. Pracuje na principu zpráv *ICMP Echo Request* a *ICMP Echo Reply* stejně jako příkaz `ping`.

`pathping -?` vypíšeme nápovědu

`pathping www.google.com` spustíme výpis směrovačů a výpočet statistik

Zjištění uzlů na cestě může být mírně rychlejší než u `tracert`, ale počítání statistiky naopak zabere více než 100 sekund a opět se nedostaneme moc daleko (dokonce můžeme „uváznout“ v síti našeho ISP). Pokud se nám zdá výpočet příliš dlouhý, program lze předčasně ukončit stisknutím .

A.4.2 Statistiky v netstatu

 Velice silný nástroj použitelný pro analýzu statistik souvisejících s protokoly rodiny TCP/IP je `netstat`. Stejně jako jiné příkazy pro práci se sítěmi vznikl podle podobně pojmenovaných nástrojů z UNIXu, u příkazu `netstat` se však setkáme se zřejmě největší podobností. Příkaz nejen zachovává UNIXovou syntaxi (přepínače se píší za pomlčku), ale také dokonce umožňuje sdružovat jednopísmenné přepínače k jedné pomlčce. Množina přepínačů je odlišná v různých verzích Windows!

`netstat` vypíše základní statistiku (nevšímá si aplikací, které jen naslouchají, vypisuje TCP spojení)

`netstat -a` vypíše celou statistiku

`netstat -a -o` vypíše celou statistiku, přidá sloupec s PID příslušného procesu

`netstat -ao` totéž (přepínače můžeme sdružovat do jediného řetězce)

`netstat -ano` navíc místo doménových vypisuje IP adresy, tato forma je velmi používaná a je dobré si ji zapamatovat

`netstat -ab > seznam.log` vypíše všechny procesy (názvy i PID), které právě komunikují se sítí (jakoliv), výsledek je směrován do zadaného souboru

`netstat -e` základní statistika pro Ethernet (počet odeslaných a přijatých paketů, počet chyb apod. – týká se protokolů rodiny TCP/IP)

`netstat -es` podrobná statistika *všech* protokolů z rodiny TCP/IP

`netstat -sp tcp` podrobná statistika, přepínač `-p` s přidáním parametrem `tcp` určuje, že chceme pouze statistiku protokolu TCP (můžeme použít parametr `tcp`, `udp` nebo `ip`)

`netstat 2` vypisuje výstup opakovaně každé dvě sekundy (můžeme zadat jakýkoliv interval a také kombinovat s jinými přepínači), činnost příkazu lze ukončit jen klávesovou zkratkou 

`netstat -r` zobrazí směrovací tabulku, stejný výstup jako příkaz `route print`



Příklad

Po několika prohlédnutých stránkách ve webovém prohlížeči může statistika IP vypadat takto:

C:\Windows> netstat -sp ip

Statistika protokolu IPv4

Počet přijatých paketů	= 30736
Přijato s chybami hlaviček	= 0
Přijato s chybami adresy	= 3509
Předané datagramy	= 0
Přijato s neznámým protokolem	= 0
Zahozené přijaté pakety	= 1078
Doručené přijaté pakety	= 26369
Požadavky na výstup	= 23914
Zahozená směrování	= 0
Zahozené výstupní pakety	= 2
Nesměrovatelné výstupní pakety	= 0
Vyžadováno nové sestavení	= 0
Úspěšná nová sestavení	= 0
Chybná nová sestavení	= 0
Úspěšně fragmentované datagramy	= 0
Neúspěšně fragmentované datagramy	= 0
Vytvořené fragmenty	= 0



Úkoly

1. V předchozím textu najdete IP adresy veřejných DNS serverů (jeden od Googlu, jeden od CZ.NIC) a vyzkoušejte na ně příkaz `ping`. Všimněte si časových údajů ve výpisech.
2. Všimněte si rozdílu (včetně časového) ve výstupech příkazů
`tracert www.google.com` `pathping www.google.com`
3. Zjistěte PID všech procesů, které právě komunikují po síti (případně naslouchají).
4. Vypíšte statistiku protokolů rodiny TCP/IP.

5. Vypíšte podrobnou statistiku protokolu IP.
6. Vypíšte podrobnou statistiku protokolu TCP, přitom zobrazte i sloupec s PID komunikujícího procesu (přepínač `-o`, například `netstat -osp tcp`). Zjistěte, kterým procesům patří řádky s navázaným spojením nebo čekající (time-wait) – může to být například Skype, mail klient, webový prohlížeč, antivirus, apod. O který proces jde u konkrétního PID, zjistíte například ve Správci úloh (zobrazte si sloupec s PID), Process Exploreru nebo ve výstupu příkazu `tasklist`.



A.5 Služba WMI a program wmic

 Program `wmic` (WMI Console) je dostupný od verze Windows XP/Server 2003. Umožňuje využívat rozhraní WMI na Příkazovém řádku. pracuje ve dvou režimech – klasickém (externím) a interaktivním. Při použití v klasickém režimu zadáváme příkaz začínající názvem programu (`wmic`) následovaný parametry, do interaktivního režimu se dostaneme zadáním příkazu `wmic` bez dalších parametrů (prompt se změní na `wmic:root\cli>` a zadáváme interní příkazy).

Nápovědu získáme buď v grafickém režimu, nebo příkazem `wmic /?`, a nebo v interaktivním režimu této konzoly příkazem `/?`. Ještě podrobnější nápovědu zobrazíme příkazem `/?:FULL`.

Příkaz `wmic` je velmi komplexní. Má tuto syntaxi:

WMIC *přepínače* *předmět* *sloveso* *parametry*, kde

- *přepínače* nastavují obecné chování příkazu, mohou být například
 - `/node:počítač` – příkaz bude proveden na zadaném počítači (před názvem počítače nebudeme dávat opačná lomítka)
 - `/user:uživatel` – při vyhodnocování bude použit zadaný uživatel s jeho přístupovými oprávněními (budeme dotázáni na heslo)
 - `/output:soubor` – výpis bude směřován do zadaného souboru
- *předmět* (také alias) určuje to, s čím chceme manipulovat nebo na co se dotazujeme, následuje *zkrácený* seznam (všechny možnosti zjistíme v nápovědě):

bios	diskdrive	memPhysical	onBoardDevice	registry
bootconfig	fsdir	netuse	OS	service
cpu	irq	NIC	pageFile	share
dcomApp	job	NICConfig	printer	temperature
desktop	memLogical	NTEvent	process	useraccount

- *sloveso* určuje požadovanou akci, která má být provedena s předmětem:
 - LIST – toto sloveso lze použít na všechny předměty, zobrazí obecnou informaci o předmětu, můžeme upřesnit, jak podrobné informace chceme (full – všechny, brief – základní v tabulce, writeable – které lze měnit, atd.),
 - GET – získání podrobnějších informací o všech nebo vybraných vlastnostech předmětu,
 - SET – změna vlastností předmětu,
 - ASSOC – vrátí instance zadaného objektu (tedy s ním asociované),
 - CREATE, DELETE – vytvoření nové instance, odstranění instance nebo třídy,
 - CALL – spuštění metody zadané třídy WMI.



Příkaz *v neinteraktivním režimu* používáme takto:

`wmic bios list` vypíše se informace o BIOSu (úplná)
`wmic os list brief` stručná informace o operačním systému
`wmic os list full` úplná informace o operačním systému
`wmic memorychip get` bude zobrazena tabulka s informacemi o paměťových čípech
`wmic memorychip get capacity,devicelocator,name,partnumber` vypíše vybrané informace o paměťových čípech (oproti předchozímu příkazu jen některé sloupce)
`wmic diskdrive get model,size,interfacetype,mediatype` příkaz zobrazí seznam všech paměťových zařízení (i výměnných), a to vlastnosti z výčtu oddělené čárkou
`wmic /output:D:\procinfo.txt cpu get` informace o procesoru (v tabulce) budou uloženy do zadaného souboru
`wmic cpu get > D:\procinfo.txt` totéž
`wmic service where state="running" get caption, name` vypíše seznam běžících služeb (pouze ty vlastnosti, které byly specifikovány), všimněte si syntaxe podobné SQL (pracujeme s databází)
`wmic process call create notepad.exe` spustí zadaný proces (zavolá proceduru create třídy process pro vytvoření procesu)
`wmic /node:počítač process call create notepad.exe` totéž, ale na jiném počítači (to příkaz `start` neumí), název počítače se zadává bez úvodních opačných lomítek
`wmic os call reboot` restart systému
`wmic /node:počítač os call shutdown` vzdálené vypnutí systému
`wmic service "spooler" call startservice` spustí zadanou službu (Zařazování tisku)
`wmic /node:ucetni process where name="explorer.exe" call terminate`
zadaný proces bude ukončen, a to na uvedeném počítači (vypadá to, že účetnímu bude ukončeno grafické prostředí, ale tento proces se obvykle po ukončení znovu spustí, tedy jde vlastně o restart grafického prostředí)
`wmic /node:ucetni /user:dadmin process where name="explorer.exe" call terminate` jako předchozí, ale v příkazu na zadaném počítači vystupujeme s oprávněními zadaného uživatele

Vidíme, že můžeme používat i dotazy filtrované podle vlastností (where) podobně jako v SQL. Pokud se jedná o řetězcovou hodnotu, musíme ji uzavřít do uvozovek (například `name="explorer.exe"`), ale čísla nebo hodnoty `true/false` do uvozovek neuzavíráme.



Příklad

Ukážeme si práci v interaktivním módu.

`wmic` spustíme interaktivní režim, prompt je `wmic:root\cli>`
`os /?` dotážeme se, co lze použít na předmět `os` (operační systém)
`os list /?` chceme upřesnit parametry pro sloveso `list`
`os list brief` získáme tabulku s několika základními informacemi
`os list full` získáme seznam (už ne tabulku) s dvojicemi vlastnost=hodnota
`os list free` tabulka s hodnotami volného místa (ve fyzické paměti, v stránkovacích souborech, ve virtuální paměti)

`/output:D:\procinfo.txt cpu get` do souboru se uloží hodně široká tabulka vlastností procesoru

`/output:D:\procinfo.txt cpu list full` totéž, ale místo široké tabulky máme v souboru seznam položek vlastnost=hodnota

`cpu get /?` zeptáme se, co se dá zjistit o procesoru

`cpu get NumberOfLogicalProcessors` získáme údaj o počtu logických procesorů (počet jader nebo jeho dvojnásobek, pokud procesor podporuje hyperthreading)

`cpu get AddressWidth,Caption,CurrentClockSpeed,DataWidth,Description,ExtClock`
vypíše se tabulka s vybranými sloupci

`process get` vypíše se seznam běžících procesů, u každého je příkaz, kterým byl spuštěn (všimněte si, že v seznamu je i `wmiprvse.exe`, který zajišťuje vyhodnocování dotazů na WMI)

`process list brief` získáme tabulku procesů s některými informacemi

`process where ThreadCount>8 list brief` podobný výstup, ale vypíšou se pouze procesy, které mají více než 8 vláken

`service get name,state,serviceType` tabulka s názvy, stavy a typy služeb

`service where desktopInteract=true get name,state` takto získáme seznam interaktivních služeb (všimněte si, že u neběžících služeb je `PID=0`)

`service where (desktopInteract=true and startMode="auto") get name,processid` výběrová kritéria můžeme kombinovat (jako v SQL), ale v tom případě je uzavřeme do závorky

`service where desktopInteract=true get name,processid,status /every:5` dotaz bude automaticky opakován v intervalu 5 sekund (stisknutím některé klávesy opakování přerušíme)

`quit` ukončíme interaktivní režim (funguje také příkaz `exit`)



Mechanismus WMI se hodně používá právě ve skriptu, a to typicky v jazyce VBScript, JScript, Perl nebo PowerShell (je třeba mít nainstalován interpret daného jazyka, což u druhého a zejména třetího uvedeného musíme zajistit sami).

Při psaní WMI skriptu bychom měli především zajistit příslušná oprávnění, zvláště když je skript spouštěn na vzdáleném počítači v LAN. Jde o oprávnění v modelu DCOM na daném počítači (WMI je totiž postaveno na COM objektech). Problém lze zjednodušit spouštěním pod některým doménovým správcovským účtem, pak není třeba řešit oprávnění zvlášť na každém počítači.



Další informace:

S mnoha WMI skripty se můžeme seznámit například ve zdroji [8] ze seznamu literatury o automatizaci správy a skriptování. Další informace o WMI (včetně WMI skriptů) najdeme například na

- <http://64.4.11.252/en-us/library/aa561208%28BTS.10%29.aspx>
- <http://www.computerperformance.co.uk/vbscript/wmi.htm>
- <http://etutorials.org/Server+Administration/Active+directory/Part+III+Scripting+Active+Directory+with+ADSI+ADO+and+WMI/Chapter+26.+Scripting+with+WMI/>
- <http://flylib.com/books/en/2.679.1.8/1/>





Úkol

Pokud máte dostatečná přístupová oprávnění, vyzkoušejte si alespoň „vypisovací“ příkazy z výše uvedeného příkladu.



A.6 Firewall ve Windows

Firewall ve Windows XP je pouze jednosměrný. Ve verzi Vista a Server 2008 se objevil obousměrný firewall, ale filtrování odchozích paketů je v desktopové verzi standardně vypnuto (dá se ale zapnout), tedy funguje jako jednosměrný. Ve vyšších verzích už firewall pracuje tak, jak má, je obousměrný.



Příklad

Prohlédneme si konfiguraci firewallu. K firewallu budeme přistupovat ve Windows (zde verze XP SP2, ve vyšších je to obdobné), a to přes *NetShell*.

```
netsh    spustíme NetShell
firewall přesun do kontextu firewall, následující příkazy provádíme v tomto kontextu
show     nejdřív si ověříme, co vše lze zobrazit
show config    zobrazíme konfiguraci firewallu (je celkem obsáhlá)
show state     zobrazí se momentální stav firewallu
show opmode    zobrazí se provozní (operační) režim (porovnejte, co se rozumí stavem z předchozího příkazu a provozním režimem z tohoto příkazu)
set ?         tímto příkazem zjistíme, co vše lze nastavovat příkazem set (tj. daná vlastnost už je definována, ale my ji můžeme změnit)
set opmode enable    nastavíme provozní mód na dostupný
show allowedprogram  zobrazí aplikace, které mají povoleno přistupovat ven do sítě
add allowedprogram ?   zajímá nás, jak lze do seznamu přidat další aplikaci
show portopening      zjistíme porty s nastavenou výjimkou (otevřené)
delete portopening ?   ze seznamu zjistíme, že je v něm port, který by tam neměl být, tedy nás zajímá, jak ho uzavřít
show service          vypíšu se služby přistupující na síť
exit                 skončíme
```



Příklad

Od verze Vista výše se s firewallem zachází trochu jinak. Předně se už nerozlišují povolené/zakázané aplikace, protokoly a porty, ale vše jsou pravidla (rules). Tedy zobrazujeme pravidla, přidáváme pravidla, odstraňujeme je.

```
advfirewall firewall    přesuneme se do kontextu advfirewall a hned v tomtéž příkazu do jeho podkontextu firewall
```

`show` ověříme, co vše lze zobrazit, zjistíme, že použitelné je `show rule`

`show rule ?` požádáme o podrobnosti

`show rule name=all` pokusíme si nechat vypsat všechna nastavená pravidla, ale seznam je příliš dlouhý, pravidla vypadají nějak takto:

Název pravidla:	Windows Live Messenger (UPnP-In)

Povoleno:	Ano
Směr:	In
Profily:	Doména,Privátní,Veřejná
Seskupení:	Windows Live Messenger
LocalIP:	Any
Vzdálená IP adresa:	LocalSubnet
Protokol:	TCP
Místní port:	2869
Vzdálený port:	Any
Funkce Edge traversal:	No
Akce:	Allow

Podobně vypadají i další pravidla. Pravidlo je povoleno, směr je „In“, tedy se jedná pravidlo pro příchozí provoz (pro odchozí by bylo „Out“), jsou stanoveny síťové profily, pro které pravidlo platí, pravidla jsou seskupována (group) pro snazší správu, dále jsou určeny adresy, protokol, port. Pokud nám nestačí konec výpisu (záleží na velikosti vyrovnávací paměti Příkazového řádku, kolik údajů bude viditelných), pak musíme příkaz spustit mimo prostředí NetShellu a výstup směřovat do souboru:

```
netsh advfirewall firewall show rule name=all > d:\vystupfirewall.txt
```

(ale pak je třeba použít editor, který zvládá znakovou sadu Latin2, například PSPad (je třeba nastavit v menu *Zobrazit ; Zobrazení znaků MSDOS*)

Bohužel musíme zadat buď „all“ nebo konkrétní název pravidla, dále lze filtrovat – omezovat výpis (ale moc možností nemáme)

`show rule name=all type=dynamic dir=in profile=public` zkonkrétnili jsme dotaz, chceme všechna pravidla, která jsou dynamicky vytvořená, pro příchozí provoz a veřejnou síť

`add rule name="povolit udp 5000" dir=out protocol=udp localport=5000`
`action=allow` povolili jsme odchozí komunikaci na UDP portu 5000

`add rule name="šifrovat 80" protocol=TCP dir=in localport=80`
`security=authdynenc action=allow` takto jsme si vynutili šifrování příchozí komunikace na TCP portu 80 (o jaký port se asi tak jedná?)



A.7 Další příkazy

Z dalších užitečných příkazů (s mnohými jsme se seznámili v Operačních systémech):

`net` příkaz pro základní správu lokální sítě (sdílení zdrojů v LAN, případně nastavení NTP serveru) a další (včetně správy uživatelů, skupin, zásad účtů apod.)

ftp FTP klient pro přenášení souborů po síti, ve Windows také najdeme jednoduchou variantu TFTP

ipseccmd konfigurace protokolu IPSec, na klientech nemusí být dostupný

openfiles slouží ke správě otevřených souborů, můžeme si nechat vypsat soubory otevřené na tomto nebo jiném počítači v síti, případně otevřené uzavřít, na klientské stanici lze pro podobné účely používat také varianty příkazu **net**

Na serverových verzích Windows najdeme mnoho dalších příkazů, které na klientech nejsou instalovány, především jde o příkazy pro správu Active Directory (např. **ntdsutil** nebo **netdom**), či pro správu DNS serveru (např. **dnsutil**) a další. Seznam s komentáři najdeme v jednom z níže uvedených odkazů. U těchto příkazů bychom měli dát pozor především na to, že jejich syntaxe se může v různých verzích systému mírně lišit.



Další informace:

Seznam příkazů pro Příkazový řádek ve Windows (ne všech) najdeme například na

- <http://technet.microsoft.com/en-us/library/cc772390%28WS.10%29.aspx>
- <http://www.microsoft.com/resources/documentation/windows/xp/all/proddocs/en-us/ntcmds.mspx?mfr=true>.
- <http://commandwindows.com/>



Práce s adresami a sítěmi v Linuxu

V této příloze jsou popsány a demonstrovány příkazy, které se používají v Linuxu při práci se sítěmi. Podobně jako ve Windows, i tuto přílohu můžeme brát jako inspiraci pro procvičování a jako předehru k zatím neexistujícím cvičením z tohoto předmětu.

B.1 Specializované distribuce

 Na routerech nižší a střední třídy se běžně setkáváme s některou distribucí *embedded* (*vestavěného*) *Linuxu* (často to ani není znát, pokud při správě využíváme webové rozhraní), protože tento systém je velmi dobře vybaven nástroji pro práci v síti. Obecně – embedded systém je systém určen k oživení zařízení, která nejsou běžnými „univerzálními“ počítači, například zmíněných síťových zařízení, ale také průmyslových strojů, spotřební elektroniky, „chytrých“ zařízení, atd.

Podobný router či bránu nebo switch s Linuxem si můžeme vyrobit sami. Lze využít starší počítač (ovšem musíme počítat s poměrně vysokou spotřebou a hlučností a potřebujeme vhodnou síťovou kartu s více rozhraními), nebo můžeme pořídit malý počítač průmyslového typu nebo přímo zařízení určené k práci v síti.

Zbývá zvolit vhodnou distribuci Linuxu. Použitelná je v podstatě jakákoliv (až na ty typicky multimediální), typicky Slackware nebo Debian (Debian je výhodný pro svou univerzálnost co se týče hardwarových platform), ale můžeme si stáhnout distribuci specializovanou na využití na síťovém zařízení. Sice budeme zřejmě „ochuzeni“ o rozbujele grafické rozhraní, ale zato budeme mít méně práce s optimalizací distribuce pro naše využití, distribuce bude mít menší nároky na zdroje (především paměť, většinou je spodní hranice 12 nebo 16 MB), instalace bývá někdy možná i přes USB (například IPCop) a pravděpodobněji bude podporovat hardwarové platformy, které se na specializovaných zařízeních používají (záleží na tom, jaký máme hardware, i podle toho vybíráme operační systém).

 Příklady vhodných distribucí:

- *IPCop*¹ je předpřipraven pro instalaci na firewallu a internetové bráně (podporuje i VLAN), lze ho spravovat přes webové rozhraní nebo konzolu. Nabízí poměrně dost služeb – DHCP klient a server, NTP klient a server, Dynamic DNS, IDS (program Snort), SSH server, HTTP/FTP proxy (program Squid), stavový firewall, IPSec VPN, atd.

¹<http://www.ipcop.org/>

- *SEntry Firewall CD*² je zajímavý tím, že se spouští z bootovacího CD (tudíž je použitelný především na zařízeních vybavených optickou mechanikou, třeba starších počítačích), „proměnlivou“ konfiguraci zapisujeme na výměnné médium (třeba na disketu, když použijeme starší počítač, disketa má výhodu velmi snadné a rychlé ochrany proti zápisu). Může sloužit jako firewall, brána (software ebttables), IDS (Snort), router (Zebra a IPRoute2), server, podporuje OpenVPN, atd.
- *Freesco*³ (ze slov „Free Cisco“) je volně šiřitelný systém obdobný systémům pro zařízení Cisco, Nortel, 3-Com apod. Může plnit roli routeru (až 10 ethernetových portů), firewallu, serveru pro protokoly DHCP, NTP (čas), DNS, FTP, HTTP, SSH, RAS, PPPoE, PPTP, apod., protože jde o Linux, tak samozřejmě i firewall a NAT.
- *Pyramid Linux*⁴ je určen především pro bezdrátové přístupové body nebo může plnit roli firewallu, je odvozen z Ubuntu. Typickou hardwarovou platformou je buď x86 nebo jednodesková zařízení.
- *Bering uClibc*⁵ je mimořádně malá distribuce určená pro routery s bránou. Její zdrojový kód je natolik upraven směrem k minimalizaci, že není kompatibilní s jinými distribucemi, tedy se při doinstalovávání dalších balíčků musíme spolehnout na repozitář této distribuce (nicméně ten je zcela dostatečně vybaven). Obsahuje celkem běžnou sadu softwaru včetně firewallu, IDS, sledování sítě, podpora bezdrátu, VPN, atd.
- *Voyage Linux*⁶ je distribuce odvozená z Debianu, určená pro hardwarové platformy založené na x86 (včetně jednodeskových zařízení). Používá se na firewallech, bezdrátových access pointech, NAS, případně routerech, má také vestavěnou podporu pro VoIP bránu (software Asterisk, ve variantě Voyage ONE).
- *Embedded Debian* (EmDebian)⁷ je Debian upravený pro použití na embedded zařízeních, odlehčený (spíše osekáný), schopný běžet na zařízeních s velmi nízkou spotřebou. Předpokládají se úpravy systému křížovou kompilací (cross-compilation).
- *Embedded Gentoo*⁸ je derivát Gentoo, opět se počítá s cross-kompilací.
- *QPlus*⁹ je embedded distribuce Linuxu pro různé typy zařízení včetně routerů.

Pokud se rozhodneme použít některou běžnou distribuci, měli bychom zajistit co nejlepší zabezpečení našeho zařízení. Velmi dobrou volbou je instalace nástroje *Bastille Linux*,¹⁰ což není distribuce, ale sada skriptů, které formou průvodce zjistí informace od systému a uživatele (s uživatelem konverzuje formou otázek a odpovědí) a pak na pokyn (se souhlasem uživatele) provede potřebná bezpečnostní nastavení. Skripty je možné procházet postupně a případně volby rušit či se k nim vracet.



Další informace:

Velmi důkladný popis přizpůsobení či vylepšení (jakékoliv) distribuce Linuxu pro použití na síťovém

²<http://www.sentryfirewall.com/>

³<http://www.freesco.org/>

⁴<http://code.google.com/p/pyramidlinux/>

⁵<http://leaf.sourceforge.net/bering-uclibc/>,

http://www.csg.ethz.ch/education/lectures/pps_firewall/SS2006/doc/bering_installation_guide.pdf

⁶<http://linux.voyage.hk/>

⁷<http://www.emdebian.org/>

⁸<http://www.gentoo.org/proj/en/base/embedded/>, <http://www.gentoo.org/proj/en/base/embedded/handbook/>

⁹<http://sourceforge.net/projects/qplus/>

¹⁰<http://bastille-linux.sourceforge.net/>

zařízení najdeme především ve zdroji [20] ze seznamu literatury. Další informace o embedded Linuxu:

- *Embedded Linux Platform Specification*. LinuxFoundation.org. URL: <http://www.linuxfoundation.org/en/ELC>
- *Embedded Linux Interfacing*. URL: <http://www.embeddedlinuxinterfacing.com/>
- *eLinux (Embedded Linux) wiki*. URL: http://elinux.org/Main_Page
- *Embedded Linux*. Felk.cvut.cz. URL: <http://support.dce.felk.cvut.cz/osp/cviceni/2/>
- THELIN, J. *Introduction: a Typical Embedded System* [online]. Linux Journal, 2009. WWW: <http://www.linuxjournal.com/magazine/introduction-typical-embedded-system>



Poznámka:

Mnohé z dále uvedených příkazů vyžadují vyšší přístupová oprávnění. Může se také stát, že příkaz bude fungovat i bez vyšších přístupových oprávnění, ale jeho výstup bude jiný (obvykle kratší nebo prázdný), to se týká například příkazů skupiny `ip neighbor`. Proto je lepší alespoň „vypisovací“ příkazy zkoušet s vyššími oprávněními.



B.2 Soubory související se sítěmi

Vybavenost UNIXových systémů nástroji pro práci se sítěmi je obecně vysoká, ovšem v každém UNIXovém systému svým způsobem specifická. Některé firmy distribuující UNIXové systémy přidávají své vlastní typické nástroje.

 Odlišnosti jsou také v souborech a adresářích, které se sítěmi souvisejí. Většina konfiguračních skriptů obvykle bývá v adresáři `/etc/sysconfig/network` a jeho podadresářích, v některých linuxových distribucích to je adresář `/etc/rc.d/rc.inet1` (Slackware), `/etc/conf.d/net` (Gentoo) nebo `/etc/network` (Debian).



Příklad

Předpokládejme, že konfigurace sítě je v adresáři `/etc/sysconfig/network-scripts`. Přesuneme se do tohoto adresáře. Měl by tam být soubor `ifcfg-eth0`, ve kterém je uložena konfigurace prvního síťového rozhraní (karty) – IP adresa, maska, způsob získání IP adresy (pokud dynamicky, tak zde IP adresu nenajdeme) apod. Pokud máme více síťových rozhraní, pro každé z nich zde bude jeden takový soubor.



V souboru `/etc/resolv.conf` nastavujeme kromě jiného adresu DNS serveru (tj. když nevíme, kde tuto adresu zadat, podíváme se do tohoto souboru). Jde o řádek (nebo o více řádků, když chceme zadat i záložní DNS servery) ve formátu

`nameserver adresa`, například

`nameserver 193.84.192.10`

Zadaný DNS server bude využíván okamžitě po uložení souboru, není třeba restartovat systém ani žádný proces.



Příklad

Předpokládejme, že z počítače chceme vytvořit router. To není až tak těžké, stačí mít zařízení s více síťovými rozhraními a provést několik nastavení. Jedno z nich je povolení forwardingu (přeposílání). To se provede menší změnou v souboru `/etc/sysctl.conf`:

- najdeme řádek `net.ipv4.ip_forward=0` (je možné, že místo teček budou lomítka, pravidla syntaxe umožňují obojí)
- změníme na `net.ipv4.ip_forward=1`
- aby změna začala platit, musíme buď restartovat systém a nebo jednoduše přimět démona `sysctl`, aby znovu načetl své konfigurační soubory (včetně toho, který jsme pozměnili a samozřejmě uložili):
`sysctl -p`

Ovšem zapnutí forwardingu nestačí, je třeba na počítačích v dotyčných sítích (které router propojuje) nastavit toto zařízení jako bránu (například pokud má karta `eth0` IP adresu `193.84.200.1`, u počítačů v síti, do které je připojena, nastavíme tuto adresu jako adresu brány, IP adresu karty `eth1` zase použijeme v druhé síti, samozřejmě vždy tu adresu, která je v dané síti viditelná).

A pak samozřejmě musíme nakonfigurovat firewall a další bezpečnostní mechanismy.

Všimněte si, že není třeba restartovat počítač, třebaže jsme provedli podstatnou změnu v nastavení systému.



V adresáři `/etc` najdeme hodně užitečných souborů a adresářů, jak víme, například `/etc/networks`, `/etc/hosts`, `/etc/ethers`.

Dynamické (momentální) nastavení sítě obvykle najdeme v podadresářích a souborech v systému `/proc`, například `/proc/net/arp`.



Příklad

Do některých souborů v `/proc` lze za určitých okolností zapisovat, například příkazem

```
echo 1 > /proc/sys/net/ipv4/ip_forward
```

zapneme forwardování (přeposílání mezi sítěmi) na takovém zařízení, které má dvě síťová rozhraní (třeba dvě síťové karty) a je tedy připojeno do dvou různých sítí. Pokud je v tomto souboru hodnota 0, nepřeposílá se, hodnota 1 znamená, že se pakety mají přeposílat mezi rozhraními.

Toto nastavení však platí jen do restartu počítače (to platí obecně o čemkoliv, co změníme v `/proc`), proto je vhodné tento příkaz uložit také do některého skriptu, který se spouští při startu systému, například do `/etc/rc.d/rc.local`, podle konkrétní distribuce.

Zatím jsme si ukázali dvě možnosti, jak zapnout forwardování (jiný způsob je v předchozím příkladu). Rozdíl mezi nimi je v tom, že zápis do souboru v souborovém systému `proc` platí jen do restartu počítače, zatímco změna v souboru `/etc/sysctl.conf` je trvalá.



Úkol

Projděte si zde probírané soubory na umístěních platných ve vaší distribuci, včetně „síťových“ částí souborového systému `/proc`.



B.3 Starší příkazy pro práci s adresami

Nejdřív se podíváme na příkazy, které se v UNIXových systémech používají pro práci s adresami již desítky let. V novějších distribucích Linuxu se však místo `ifconfig`, `route` a `arp` doporučuje používat nový příkaz `ip`, z důvodu kompatibility a podpory nových technologií.

 Příkaz `ifconfig` slouží ke konfiguraci síťových rozhraní (interfaces) jádra (na rozdíl od Windows jsou v Linuxu síťová rozhraní součástí jádra).

`ifconfig` zobrazí informaci o všech síťových rozhraních

`ifconfig eth0` zobrazí informaci o síťovém rozhraní `eth0` (to obvykle bývá ethernetová karta)

`ifconfig eth0 down` „shodí“ zařízení `eth0` (zneaktivní tuto kartu), provádíme například tehdy, když chceme tomuto rozhraní přiřadit jinou adresu

`ifconfig eth0 up` aktivuje zařízení `eth0`

`ifconfig eth0 down hw ether 00:00:00:00:00:02` kromě toho, že zneaktivní zařízení `eth0`, také mu přiřadí MAC adresu; vnitřní příkaz `hw ether adresa` znamená, že jde o ethernetovou kartu, které přiřazujeme danou adresu (před podobnými změnami je vždy nutné kartu zneaktivnit)

`ifconfig eth0 172.19.124.104 netmask 255.255.255.0 up` nastaví IP adresu a masku podsítě pro `eth0`, pak je zaktivní (předpokládáme, že předtím byla tato karta neaktivní)

`ifconfig eth0 promisc` zapne promiskuitní mód (tj. karta přijímá/odposlouchává všechny pakety, nejen ty, které jsou jí adresované)

`ifconfig eth0 -promisc` vypne promiskuitní mód

Tento příkaz má další volby, kterými lze například určovat multicast a broadcast adresy, přidělovat zdroje (I/O paměť, IRQ apod.), nastavovat metriky, atd.



Příklad

Příkaz `ifconfig` na počítači s jednou síťovou kartou bude vypadat takto:

```
sarka@notebook:~$ ifconfig
```

```
eth0      Zapouzdření:Ethernet  HWadr 00:1D:72:31:0A:A0
          inet adr:10.0.0.2  Věsměr:10.0.0.255  Maska:255.255.255.0
          inet6-adr: fe80::21d:72ff:fe31:aa0/64  Rozsah:Linka
          AKTIVOVÁNO VŠESMĚROVÉ_VYSÍLÁNÍ BĚŽÍ MULTICAST  MTU:1500  Metrika:1
          RX packets:561 errors:0 dropped:0 overruns:0 frame:0
          TX packets:556 errors:0 dropped:0 overruns:0 carrier:0
          kolizí:0 délka odchozí fronty:1000
          RX bytes:379302 (370.4 KiB)  TX bytes:100735 (98.3 KiB)
          Přerušení:169

lo        Zapouzdření:Místní smyčka
          inet adr:127.0.0.1  Maska:255.0.0.0
          inet6-adr: ::1/128  Rozsah:Počítač
          AKTIVOVÁNO SMYČKA BĚŽÍ  MTU:16436  Metrika:1
          RX packets:270 errors:0 dropped:0 overruns:0 frame:0
          TX packets:270 errors:0 dropped:0 overruns:0 carrier:0
          kolizí:0 délka odchozí fronty:0
          RX bytes:24914 (24.3 KiB)  TX bytes:24914 (24.3 KiB)
```

Výpis samozřejmě může být celý v angličtině, například místo `délka odchozí fronty` bychom pak našli řetězec `txqueue_len`, nebo místo `AKTIVOVÁNO VŠESMĚROVÉ_VYSÍLÁNÍ BĚŽÍ MULTICAST` by bylo `UP BROADCAST NOTRAILERS RUNNING MULTICAST`.

Srovnajte s výstupem příkazů `ip link show` a `ip addr show dev eth0` v příkladu na straně 262, všechny tyto příkazy byly spuštěny na tomtéž počítači, ale v jiné situaci (tj. například nebudou sedět počty RX a TX bytes/packets).



Pro rychlé odpojení a aktivování síťového rozhraní existují také zkrácené verze příkazu:

`ifdown eth0` deaktivace síťové karty

`ifup eth0` aktivace karty

 Pro práci se směrovací tabulkou lze využít příkaz `route`.

`route` zobrazí hlavní směrovací tabulku (k jiným směrovacím tabulkám se tímto příkazem nedostaneme)

`route -n -A inet` první parametr způsobí vypisování číselných adres místo doménových (u některých příkazů také ve formě `-no-dns`, ale většina rozumí pouze `-n`), druhý stanoví, že se mají vypsat pouze IPv4 adresy (pro IPv6 by bylo `-A inet6`)

`route add -net 193.90.100.0 netmask 255.255.255.128 gw 193.90.220.3 dev eth0` přidáme do směrovací tabulky nový záznam (jde o adresu sítě se zadanou maskou, posledním parametrem je brána, na kterou se mají posílat pakety určené pro danou síť)

`route add default gw 193.88.50.10 dev eth0` určíme výchozí bránu pro zadané síťové rozhraní

`route del -net 193.90.100.0 netmask 255.255.255.128 gw 193.90.220.3` odstraníme řádek směrovací tabulky (obvykle stačí zadat jen údaj o IP adrese, kterou chceme odstranit z tabulky)

`route del default` odstraníme směrovací informaci o výchozí bráně, typicky když chceme nastavit jinou

`route -Cn` zobrazí se cache pro směrování, a to s číselnými IP adresami

Místo příkazu `route` se doporučuje používat příkaz `ip route`.

 V Linuxu se také používá příkaz `arp`, a to s podobnou syntaxí jako ve Windows:

`arp -a -n` zobrazí ARP tabulku, druhý přepínač znamená, že všechny adresy se zobrazí v číselném tvaru (místo doménových názvů), ale můžeme použít samozřejmě `arp -a` (ovšem když dochází k překladu IP adres na doménové, příkaz pracuje pomaleji)

`arp -n -i eth1` zobrazí ARP tabulku zadaného síťového rozhraní (to následuje za přepínačem `-i`)

`arp -s 123.123.123.123 00:12:34:56:01:08 -i eth1` přidání statického záznamu do ARP tabulky pro kartu `eth1`

`arp -d 123.123.123.123 -i eth1` odebrání záznamu z ARP tabulky pro kartu `eth1`

Je možné, že v některých distribucích nebude v nejnovějších verzích některý ze starších příkazů fungovat. Z toho a z dalších důvodů je lepší zvyknout si na příkaz `ip`.



Úkoly

1. Zjistěte, jaká síťová rozhraní na počítači máte. Zjistěte příslušné MAC adresy a přidělené IP adresy. Pokuste se některé síťové zařízení deaktivovat a pak znovu aktivovat (ne loopback!), pokud získáváte IP adresu od DHCP serveru, srovnajte adresy před deaktivací a po aktivaci.
2. Zjistěte, zda máte přidělenou také IPv6 adresu. Pokud ano: doplňte vynechané nulové sekce adresy (stačí jedna nula na sekci) a oddělte prefix od klientské části adresy. Byla klientská část získána autokonfigurací? (tj. srovnajte s MAC adresou, zda odpovídá EUI-64)

3. Zobrazte hlavní směrovací tabulku tak, aby se IP adresy nepřekládaly na doménové. Srovnajte s výpisem obsahujícím doménové adresy. Zjistěte, jakou IP adresu má výchozí brána v hlavní tabulce. Zobrazte směrovací cache (opět s IP adresami místo doménových).
4. Zobrazte ARP tabulku (opět tak, aby se IP adresy nepřekládaly na doménové, srovnajte s výpisem s doménovými adresami).



B.4 Mechanismus iproute2 (příkaz ip) – adresy, síť, směrování

 Ve všech novějších distribucích se setkáme s ještě komplexním příkazem `ip`, který je součástí balíčku `iproute2`. Tento příkaz byl zařazen jako náhrada příkazů `ifconfig`, `route` a `arp` (pro úplnost jsme se seznámili i s nimi), v současné době se místo těchto tří doporučuje používat spíše příkaz `ip`. Jeho konfigurační soubory najdeme v `/etc/iproute2`. Má specifické vnitřní příkazy, postupně probereme nejdůležitější varianty.

B.4.1 Konfigurace síťového rozhraní a adres

Nejdřív se podíváme na formu příkazu `ip` pro práci se síťovými rozhraními a IP adresami.

 `ip link ...` pracujeme na vrstvě L2 (linkové) ISO/OSI modelu, v podstatě i L1, tedy pracujeme s MAC adresami a síťovým rozhraním (například můžeme kartu odpojit či připojit)

`ip link show` zobrazí stav spojení (tj. stav síťového rozhraní), místo `show` je možné všude používat `list` nebo `ls`, na multihomed zařízeních může být výpis velmi dlouhý; na rozdíl od `ifconfig` se výpisem i *ta síťová rozhraní, která z nějakého důvodu nelze řádně aktivovat*

`ip -s link show eth0` zadali jsme název rozhraní (tj. nebudou se vypisovat informace pro všechna, pokud jich je více), a navíc parametr `-s` znamená podrobnější výpis

`ip link set dev eth0 up` aktivuje rozhraní `eth0`

`ip link set dev eth0 down` deaktivuje rozhraní `eth0`

`ip link set dev eth0 mtu 1500` změní hodnotu MTU na 1500 oktetů (Maximum transmission unit, maximální velikosti IP paketu, kterou je možné přijmout a odeslat)

`ip link set dev eth0 address 02:00:00:00:11:22` změníme MAC adresu síťové karty `eth0`, všimněte si prvního oktetu; pozor, změna MAC nemusí někdy dopadnout dobře, jde o potenciálně nebezpečnou operaci

 `ip addr ...` pracujeme na vrstvě L3 (síťové), tedy s IP adresami (místo `addr` lze použít `address` nebo `a`)

`ip addr show` takto zjistíme IP adresu a související informace, případně můžeme přidat i název karty, která nás zajímá, jedna karta může mít i více než jednu IPv6 adresu (jednu primární a ostatní sekundární)

`ip address show` totéž, také funguje `ip a show` nebo `ip a ls`

`ip addr show dev eth0 primary` zjistíme primární adresu zařízení `eth0`

`ip addr add 193.90.220.42/25 brd + dev eth0` nastavíme IP adresu pro rozhraní `eth0`,¹¹ broadcast adresu necháme nastavit na výchozí (pokud chceme určit jinou než výchozí, místo symbolu `+` napíšeme novou broadcast adresu), všimněte si, že místo masky zadáváme délku prefixu hned u adresy

`ip addr del 193.90.220.42/25 brd + dev eth0` odebereme kartě `eth0` zadanou adresu; pozor – pokud se jedná o primární adresu, odeberou se i všechny sekundární

`ip addr flush dev eth0` pročištění (flush, odstranění) všech adres na daném rozhraní, zůstane pouze IPv6 adresa získaná autokonfigurací, vnitřnímu příkazu `flush` se spíše vyhýbáme (můžeme omylem odstranit adresy, které by měly zůstat)

`ip -f inet6 addr flush dynamic` před vnitřním příkazem je použit parametr určující, že se bude pracovat pouze s IPv6 adresami (místo `-f inet6` je možné napsat zkráceně `-6`), odstraní se adresy IPv6 získané dynamicky (tj. kromě adresy získané autokonfigurací), pokud chceme provést totéž pro IPv4 adresy, napíšeme

`ip -f inet addr flush dynamic` nebo

`ip -4 addr flush dynamic` nebo

`ip -4 a flush dynamic`



Příklad

Podíváme se na výstupy několika příkazů:

```
sarka@notebook:~$ ip link show
```

```
1: lo: <LOOPBACK,UP,10000> mtu 16436 qdisc noqueue
    link/loopback 00:00:00:00:00:00 brd 00:00:00:00:00:00
2: eth0: <BROADCAST,MULTICAST,UP,10000> mtu 1500 qdisc pfifo_fast qlen 1000
    link/ether 00:1d:72:31:0a:a0 brd ff:ff:ff:ff:ff:ff
3: sit0: <NOARP> mtu 1480 qdisc noop
    link/sit 0.0.0.0 brd 0.0.0.0
```

```
sarka@notebook:~$ ip -s link show
```

```
1: lo: <LOOPBACK,UP,10000> mtu 16436 qdisc noqueue
    link/loopback 00:00:00:00:00:00 brd 00:00:00:00:00:00
    RX: bytes  packets  errors  dropped overrun mcast
    24914      270      0       0       0       0
    TX: bytes  packets  errors  dropped carrier collsns
    24914      270      0       0       0       0
2: eth0: <BROADCAST,MULTICAST,UP,10000> mtu 1500 qdisc pfifo_fast qlen 1000
    link/ether 00:1d:72:31:0a:a0 brd ff:ff:ff:ff:ff:ff
    RX: bytes  packets  errors  dropped overrun mcast
    379366     562      0       0       0       9
    TX: bytes  packets  errors  dropped carrier collsns
    100799     557      0       0       0       0
3: sit0: <NOARP> mtu 1480 qdisc noop
    link/sit 0.0.0.0 brd 0.0.0.0
    RX: bytes  packets  errors  dropped overrun mcast
    0          0        0       0       0       0
    TX: bytes  packets  errors  dropped carrier collsns
    0          0        0       0       0       0
```

¹¹Jedno síťové rozhraní může mít více IPv6 adres. Jestliže již má IPv6 adresu přiřazenu, další bude sekundární. Pokud chceme jednomu síťovému rozhraní přiřadit více IPv6 adres, měli bychom každé z těchto adres přiřadit label – viz `man ip`, především z důvodu jejich rozlišení v příkazech, například `ifconfig` s tím mívá problémy. Taktéž lze IP adresy jednoho rozhraní rozlišit pomocí aliasů, což je rozlišení na nižší úrovni než v případě labelů, například příkazem `ifconfig eth1:0 193.84.190.99 netmask 255.255.128.0 up` vytvoříme alias `eth1:0` k rozhraní `eth1` s vlastní IP adresou, který lze také deaktivovat bez ovlivnění rozhraní `eth1`.

Všimněte si, že v ostrých závorkách za názvem rozhraní jsou příznaky rozhraní (například zda jde o loopback, jestli je aktivní, povolení broadcastu, apod.).

```
sarka@notebook:~$ ip addr show
```

```
1: lo: <LOOPBACK,UP,10000> mtu 16436 qdisc noqueue
    link/loopback 00:00:00:00:00:00 brd 00:00:00:00:00:00
    inet 127.0.0.1/8 scope host lo
    inet6 ::1/128 scope host
        valid_lft forever preferred_lft forever
2: eth0: <BROADCAST,MULTICAST,UP,10000> mtu 1500 qdisc pfifo_fast qlen 1000
    link/ether 00:1d:72:31:0a:a0 brd ff:ff:ff:ff:ff:ff
    inet 10.0.0.2/24 brd 10.0.0.255 scope global eth0
    inet6 fe80::21d:72ff:fe31:aa0/64 scope link
        valid_lft forever preferred_lft forever
3: sit0: <NOARP> mtu 1480 qdisc noop
    link/sit 0.0.0.0 brd 0.0.0.0
```

Kdyby nás zajímalo pouze zařízení eth0, napsali bychom

```
ip addr show dev eth0
```



Úkoly

1. Vypište údaje o svém síťovém rozhraní, a pak je vypište včetně statistiky (podrobnější informace). Srovnajte oba výpisy. Určete, kolik máte síťových rozhraní (na koncové stanici to odpovídá počtu síťových karet plus případných virtuálních rozhraní), jakou máte velikost MTU, jaké máte MAC adresy, příp. jaká je broadcast MAC adresa, jaký je provoz na vstupu/výstupu rozhraní.
2. Zjistěte IP adresy (IPv4 i IPv6) na svých síťových rozhraních. Máte jen primární adresu, nebo i nějaké sekundární?
3. Pokud je na rozhraní definováno více adres, jak zobrazíte pouze tu primární? Co se stane, když síťovému rozhraní s více IP adresami odebereme primární/sekundární adresu?



B.4.2 Směrování a filtrování

Příkaz `ip` se používá při práci se směrovacími tabulkami a při definování pravidel pro filtrování obdobně jako u firewallu.

 V Linuxu neexistuje pouze jedna směrovací tabulka. Ve výchozím nastavení máme vždy alespoň hlavní (*main*) směrovací tabulku, která je zároveň výchozí (*default*), a *local* – lokální tabulku (v lokální směrovací tabulce najdeme především loopback (místní cesty) a směrování broadcastů a multicastů), další tabulky si můžeme dle potřeb vytvořit. Každá tabulka je identifikována svým jménem a číslem, v příkazech můžeme používat cokoliv z toho. Co se týče čísel směrovacích tabulek, tak u nově vytvořených můžeme použít čísla z intervalu 1 až 252, čísla předdefinovaných tabulek vidíme ve výpisu níže.

Zatímco příkaz `route`, se kterým jsme se již dříve seznámili, dovoluje přistupovat pouze k hlavní tabulce, příkaz `ip` umožňuje pracovat s jakoukoliv tabulkou.

Seznam směrovacích tabulek je uložen v souboru `/etc/iproute2/rt_tables`. Výchozím obsahem je

```
254     main
253     default
0       unspec
```

(plus řádky s komentáři, komentářový symbol je #).

 Novou směrovací tabulku vytvoříme jednoduše přidáním nového řádku do tohoto souboru (další informace najdeme u příkazu `ip rule` od strany 265). Obvykle soubor needitujeme přímo, ale použijeme přesměrování s přidáním na konec:

```
echo číslo název » /etc/iproute2/rt_tables
```

například

```
echo 15 novatab » /etc/iproute2/rt_tables
```

Všimněte si, že při přesměrování jsme použili dvojitou „šipku“, protože nechceme přepsat původní obsah (na to pozor!), ale přidat nový záznam na konec souboru.

Příkazy:

 `ip route ...` pracujeme se směrovací tabulkou, taktéž na vrstvě L3 (ale nad protokolem IP)

`ip route show` zobrazí hlavní směrovací tabulku (pozor, ne všechny)

`ip route show table local` zobrazí lokální směrovací tabulku

`ip route show table 12` zobrazí směrovací tabulku číslo 12 (na tabulku se odkazujeme jménem nebo číslem)

`ip route show cache` zobrazí vyrovnávací paměť směrování; velmi dlouhý výstup, je dobré přidat ještě IP adresu, která nás zajímá, a nebo výstup nějak filtrovat (třeba `grepem`)

`ip -s route show cache 193.90.102.36` zadali jsme IP adresu, která nás jako cíl zajímá, navíc přepínačem `-s` vyžadujeme podrobnější výstup (tj. vlastně statistiku směrování na danou adresu)

`ip route add default via 193.90.220.1` nastavení výchozí brány pro všechny cíle, které ve směrovací tabulce nejsou přímo uvedeny

`ip route add 193.90.100.0/25 via 193.90.220.3` přidání nového (statického) řádku do směrovací tabulky (zadááme adresu sítě s prefixem a dále adresu zařízení, přes které se k první zadané adrese dá dostat)

`ip route add 193.90.100.0/25 via 193.90.220.3 table administrativa` záznam jsme přidali do zadané směrovací tabulky `administrativa`, nikoliv do hlavní směrovací tabulky

`ip route delete 193.90.100.0/25` odstranění řádku tabulky (příp. lze opět zadat tabulku)

`ip route add prohibit 193.221.88.0/28` zakážeme směrování na zadanou adresu, tj. defacto tuto adresu znepřístupníme ze zařízení s touto tabulkou (používají například zaměstnavatelé k zamezení přístupu z daného počítače/podsítě k určitým oblastem), zadaná podsít' či uzel je ohlášen(a) jako „no route to host“ (ICMP zpráva)

`ip route add prohibit 193.221.88.0/28 from 193.90.102.29` podobně jako předchozí, ale zablokovali jsme přístup pouze z jednoho (zadaného) uzlu, z jiných uzlů bude směrování fungovat (klíčové slovo `from` použijeme typicky na směrovači s Linuxem, na koncové stanici nemá moc smysl)

`ip route add blackhole 69.63.189.11` pokud někdo z lokální sítě odešle paket na tuto adresu (mimochodem, je to jedna z IP adres serverů Facebooku), paket bude zahozen a dotyčný nebude informován

`ip route add unreachable 69.63.189.11` podobně jako předchozí, ale odesílatel je informován o tom, že paket nedorazil, ICMP zprávou „ICMP unreachable“;

router tedy může paket „záměrně“ zahodit třemi různými způsoby (prohibit, blackhole, unreachable).

`ip route add nat 192.205.120.56 via 193.90.102.29` stanoví NAT překlad pro příchozí pakety (přicházející do LAN): když tento router dostane paket s první („virtuální“) adresou jako cílovou, přepíše ji na druhou uvedenou (skutečnou v místní síti) a pošle dál



Příklad

Ukážeme si rozdíl mezi hlavní a lokální směrovací tabulkou na klientské stanici:

```
sarka@notebook:~$ ip route show
```

```
193.84.195.0/25 dev eth0 proto kernel scope link src 193.84.195.30
default via 193.84.195.1 dev eth0
```

```
sarka@notebook:~$ ip route show table local
```

```
broadcast 127.255.255.255 dev lo proto kernel scope link src 127.0.0.1
broadcast 193.84.195.0 dev eth0 proto kernel scope link src 193.84.195.30
local 193.84.195.30 dev eth0 proto kernel scope host src 193.84.195.30
broadcast 193.84.195.127 dev eth0 proto kernel scope link src 193.84.195.30
broadcast 127.0.0.0 dev lo proto kernel scope link src 127.0.0.1
local 127.0.0.1 dev lo proto kernel scope host src 127.0.0.1
local 127.0.0.0/8 dev lo proto kernel scope host src 127.0.0.1
```

Řádek začínající `default` v hlavní tabulce řádku určuje bránu lokální sítě.



Směrování lze spojit s filtrováním či podrobnějším řízením směrování podle zadaných *pravidel* (politik), čemuž říkáme *Advanced Routing* (pokročilé směrování). Zatímco příkaz `ip route` pracuje pouze s adresami, pomocí `ip rule` lze směrování ovlivnit také obsahem jiných polí IP paketu (směrování podle zásad), takto lze implementovat i mechanismus NAT.



`ip rule ...` pracujeme s pravidly (politikami, zásadami) pro směrování

`ip rule show` výpis pravidel (vlastně zásad) pro směrování, jsou tam minimálně tyto řádky:

```
0:      from all lookup local
32766:  from all lookup main
32767:  from all lookup default
```

(případně místo `default` může být přímo číslo výchozí směrovací tabulky 253), výchozí hodnoty pouze odkazují na lokální, hlavní a výchozí směrovací tabulku, číslo na začátku je tzv. priorita pravidla, toto číslo by mělo být jedinečné pro každé pravidlo (čím menší číslo, tím vyšší priorita)

`echo 12 pomocnatab > /etc/iproute2/rt_tables` trochu jsme odbočili, tímto příkazem jsme vytvořili směrovací tabulku s číslem 12 a názvem `pomocnatab`

`ip rule add from 195.200.94.0/24 table 12 priority 354` všechny pakety s uvedenou adresou jako zdrojovou (resp. adresou ze zadané sítě) budou směrovány podle tabulky 12 vytvořené v předchozím příkazu, hodnota pro prioritu pravidla by měla být u pravidel unikátní (pokud zvolíme již existující prioritu, jádro přidělí nejbližší menší volné číslo; když prioritu neuvedeme vůbec, jádro nějakou zvolí samo), čím menší číslo, tím vyšší priorita; pokud teď zadáme `ip rule show`, bude výstup následující (srovnejte s prvním výpisem pravidel):

```
0:      from all lookup local
354:    from 195.200.94.0/24 lookup pomocnatab
32766:  from all lookup main
32767:  from all lookup 253
```

`ip rule add from 195.80.94.0/24 to 194.200.24.0/24` určíme, že pakety ze sítě s první zadanou adresou do sítě s druhou zadanou adresou budou směrovány podle hlavní směrovací tabulky (případně můžeme zadat, která tabulka se má použít), prioritu pravidla zvolí jádro (všimněte si, které číslo jádro zvolilo: o něco nižší číslo než poslední uživatelské), ve výstupu příkazu `ip rule show` přibude tento řádek:

```
353:    from 195.200.94.0/24 to 194.200.24.0/24 lookup main
```

`ip rule add from 193.90.102.29 nat 192.205.120.56` NAT pravidlo pro odchozí pakety (tj. odcházející ze sítě): pokud má odejít ven ze sítě paket se zdrojovou adresou první uvedenou (skutečnou), bude zdrojová adresa nastavena na druhou uvedenou („virtuální“; všimněte si, jaký je vztah mezi adresami v podobném výše uvedeném příkazu `ip route`, musíme mít provedeny oba příkazy, aby to fungovalo)

`ip rule add iif lo table 10 priority 560` takto můžeme oddělit směrování provozu generovaného přímo na tomto zařízení (lo, loopback, poněkud netypické použití loopbacku) od směrování přeposílaného (forwarding) provozu, tedy provoz generovaný procesy na tomto zařízení bude směrován podle tabulky 10 (případně můžeme takto oddělit provoz směrovaný z konkrétního rozhraní, třeba eth2 coby třetí ethernetové síťové karty), do seznamu pravidel přibývá nový řádek (předpokládejme, že tabulka 10 se nazývá *dalsitabulka*):

```
560:    from all iif lo lookup dalsitabulka
```

`ip rule del priority 560` odstraníme pravidlo se zadanou prioritou (připomeňme si, že priorita je vlastně pořadové číslo v tabulce pravidel, pravidlo je takto jednoznačně identifikováno)



Příklad

Podíváme se na možnost vytvoření jednoduchého one-to-one stateless (bezstavového) NAT, tedy překládáme mezi jednou vnitřní a jednou vnější IP adresou. Toho lze docílit buď příkazem `ip`, nebo mechanismem *NetFilter* (to, jak víme, je modul jádra obvykle používaný jako firewall, jehož obrazem v uživatelském režimu je *iptables*). Pokud použijeme příkaz `ip`, bude náš router odpovídat na ARP dotazy na NAT IP adresu, což může být užitečné (mechanismus *NetFilter* tuto vlastnost nemá).

Předpokládejme, že chceme zpřístupnit stroj s IP adresou 195.159.200.28 ven ze sítě pod adresou 207.189.240.28, může se například jednat o některý server v DMZ:

`ip route add nat 207.189.240.28 via 195.159.200.28` nejdřív zajistíme překlad cílové adresy v přichozím (inbound) provozu, tedy DNAT

`ip rule add nat 207.189.240.28 from 195.159.200.28` zajistíme překlad zdrojové adresy v odchozím provozu, tedy SNAT (Source NAT, viz str. 289)

`ip route flush cache` aby překlad fungoval, musíme vyprázdnit směrovací cache se starými záznamy

`ip route show table all` ve výstupu by měl být kromě jiného řádek

```
nat 207.189.240.28 via 195.159.200.28 table local scope host
```

(řádků je tam obvykle hodně, výstup můžeme profiltrovat například příkazem `grep`)

`ip rule show` podobně zkontrolujeme seznam pravidel pro směrování, mezi nimi by měl být řádek
 32386: from 195.159.200.28 lookup main map-to 207.189.240.28
 (pravděpodobně s jinou hodnotou priority)

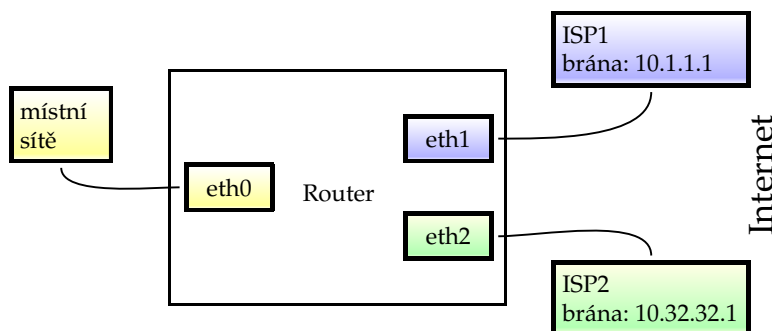
Totéž lze provést s použitím mechanismu *NetFilter* zároveň s určením dalších údajů pro překlad (například čísla TCP portu), ukázku máme na straně 291.



Příklad

Předpokládejme, že máme dva ISP (poskytovatele připojení k Internetu) připojené k témuž routeru s Linuxem. Tento případ není až tak exotický, můžeme se s ním setkat u některých středních/větších společností (respektive, nemusí jít ani o více různých ISP, stačí, když máme více bran).

Budeme chtít vytvořit dvě směrovací tabulky, zvlášť pro každého ISP. Situace je znázorněna na obrázku B.1.



Obrázek B.1: Situace pro více ISP přes jeden router

Nejdřív „vytvoříme“ dvě směrovací tabulky, pro každého ISP zvlášť:

```
echo 10 ISPprvni » /etc/iproute2/rt_tables
echo 20 ISPdruhy » /etc/iproute2/rt_tables
```

Teď se musíme rozhodnout, jak se bude dělit provoz mezi jednotlivé brány (a tedy i mezi jednotlivá rozhraní routeru). Předpokládejme, že provoz související s lokální sítí 10.192.0.0/24 půjde přes první bránu (rozhraní eth1) a bude směrován tabulkou 10, provoz sítě 10.192.128.0/24 půjde přes druhou bránu a bude směrován tabulkou 20, zbytek půjde také přes druhou bránu a bude směrován hlavní směrovací tabulkou.

Nakonfigurujeme směrovací tabulku a pravidla:

```
ip route add default via 10.32.32.1 dev eth2   nastavíme výchozí bránu pro pakety ze sítí neuvede-
ných jinde
ip route add default via 10.1.1.1 dev eth1 table 10   do tabulky 10 (ISPprvni) jsme přidali nový
záznam – pro vše, co bude směrováno podle této tabulky, platí výchozí brána 10.1.1.1, ke které se
dostaneme přes rozhraní eth1
ip route add default via 10.32.32.1 dev eth2 table 20   podobně pro další tabulku, určíme bránu
a port
ip rule add from 10.192.0.0/24 table 10   vše, co přijde ze sítě se zadanou adresou, bude směrováno
podle tabulky 10
ip rule add from 10.192.128.0/24 table 20   podobně pro další síť, určíme směrovací tabulku
```

Provoz lze směřovat i podle jiných kritérií.



Úkoly

1. Najděte seznam směrovacích tabulek (vypíšte soubor `/etc/iproute2/rt_tables`). Dále zjistěte, jak jsou nastavena přístupová oprávnění k tomuto souboru:
 - `ls -la /etc/iproute2/rt_tables` vypíše základní oprávnění,
 - `getfacl /etc/iproute2/rt_tables` vypíše seznamy POSIX ACL,
 - `lsattr /etc/iproute2/rt_tables` zobrazí atributy souboru v daném souborovém systému,
 - `getfattr /etc/iproute2/rt_tables` zobrazí rozšířené atributy souboru.
2. Vypíšte hlavní směrovací tabulku a porovnejte s výpisem pomocí příkazu `route` (je v předchozí sekci). Dále vypíšte lokální směrovací tabulku, a pokud existují další, také je vypíšte. Najděte adresu výchozí brány. Jak byste ji změnili?
3. Vypíšte seznam definovaných politik (zásad, pravidel) pro směrování. Všimněte si, na které směrovací tabulky tyto zásady odkazují.



B.4.3 Objevování sítě

Protokol IPv4 používá pro objevování sítě protokol ARP, u IPv6 to je NDP (Neighbour Discovery Protocol), respektive mechanismus NDis (Network Discovery). Pro práci s „okolím“ máme k dispozici další vnitřní příkaz příkazu `ip`.



Jednotlivé uzly sítě mohou být v některém z následujících stavů (označujeme NUD, Neighbour Unreachability Detection):

- *none* – stav „nevyplněn“
- *noarp* – soused je platný a dosažitelný, nemá být ověřován ARP dotazy, záznam má stanovenou dobu platnosti (po této době je vyřazen z ARP cache)
- *permanent* – podobně jako *noarp* je platný, dosažitelný a nemá být ověřován ARP dotazy, nemá stanovenou dobu platnosti (tj. platnost neomezená), z tabulky sousedů smí být odstraněn jen administrátorem
- *reachable* – soused je platný a dosažitelný, záznam má stanovenou dobu platnosti (životnost záznamu), po uplynutí této doby bude opětovně dynamicky zjišťován ARP dotazem (tento stav je běžný pro dynamicky zjišťované sousedy)
- *incomplete* – soused je „zjišťován“, proces objevování souseda není ukončen (obvykle znamená, že k dané IP adrese, na kterou se dotazuje některý protokol vyšší vrstvy, neexistuje žádný záznam)
- *stale* – soused je považován za platného, ale jeho dosažitelnost musí být neustále ověřována (po každém použití záznamu tento záznam přechází do stavu *delay*)
- *delay* – je třeba záznam ověřit, byl naplánován ARP dotaz na tuto adresu, pokud tento soused bude generovat provoz, stav se mění zpět na *stale*
- *probe* – sousedovi s příznakem *delay* vypršel limit na odpověď (nebo generování provozu), jádro odešle ARP dotaz

- *failed* – ARP dotaz selhal, záznam není platný

Většina sousedů je ve stavu *reachable*, *stale* nebo *permanent*.

Pokud je soused ve stavu *stale* a je připojen, pak v tomto stavu tráví většinu času (občas přechází do *delay*). Pokud je soused ve stavu *reachable*, obvykle v něm zůstává, ale když ho odpojíme, postupně přechází mezi stavy *reachable*–*stale*–*delay*–*probe*–*failed*, dotaz na danou IP adresu pak končí výpisem stavu *incomplete*. Příkazy:



`ip neighbour ...` pracujeme s vazbami mezi adresami L2 a L3 (tj. MAC a IP) ve stejné síti (tj. se „sousedý“, odtud klíčové slovo), v případě IPv4 jde o práci s ARP tabulkami, u IPv6 jsou to neighbour tabulky (slovo *neighbour* lze zkracovat na *neighbor*, *neigh* nebo *n*)

`ip neigh show` zobrazí momentální stav ARP nebo NDP cache (tabulky)

`ip neigh show dev eth0` zobrazí seznam sousedů připojených k rozhraní `eth0`

`ip neigh show to 193.168.200/24` zobrazí se info o sousedech ze zadané podsítě

`ip -s neigh show to 10.0.0.1` podrobnější statistika souseda se zadanou adresou (navíc počet uživatelů záznamu a dále před kolika sekundami byl záznam použit/potvrzen/aktualizován)

`ip neigh show nud permanent` zobrazí všechny sousedy, jejichž stav je „permanent“, obvykle jde o ručně přidávané sousedy

`ip neigh add 193.168.200.1 lladdr 00:0a:1b:2c:3d:4e dev eth2 nud permanent` přidali jsme do ARP tabulky statický záznam o sousedovi (zadáváme IP adresu, MAC adresu – Link Layer Addr, rozhraní, přes které je soused dosažitelný, a pak stav)

pokud přidáváme nový záznam, můžeme hodnotu `nud` nastavit na jednu z hodnot `noarp`, `reachable`, `permanent` nebo `stale`, čímž také určíme, zda záznam bude mít omezenou platnost a jestli má být ověřován

`ip neigh del 193.168.200.1 dev eth2` odstranili jsme záznam z tabulky



Příklad

Zatímco mezilehlé síťové zařízení (router, switch) má sousedů poměrně hodně, koncové zařízení připojené přes ethernet (konektor RJ-45) má obvykle jediného souseda – router nebo switch.

Na koncovém zařízení vypadají výpisy následovně:

```
sarka@notebook:/home/sarka# ip neigh show
```

```
10.0.0.138 dev eth0 lladdr 00:21:63:e2:0f:98 STALE
```

```
sarka@notebook:/home/sarka# ip -s neigh show
```

```
10.0.0.138 dev eth0 lladdr 00:21:63:e2:0f:98 ref 12 used 172/172/152 STALE
```

Druhý výpis obsahuje podrobnější údaje o sousedovi (což je router, jehož uvnitř viditelná IP adresa je 10.0.0.138). Na řádku postupně vidíme IP adresu, rozhraní, přes které jsme k sousedovi připojeni, dále počet odkazů na toto připojení a tři časové údaje (počet sekund od chvíle, kdy bylo připojení naposledy použito/potvrzeno/aktualizováno).





Úkoly

1. Zobrazte seznam svých sousedů (ARP/NDP cache). Pokud máte více aktivních síťových rozhraní, pak zvlášť pro několik rozhraní. Vyzkoušejte zobrazení se statistikou (podrobnějšími informacemi) o daném propojení.
2. Přidejte do ARP cache „virtuálního“ souseda, jehož MAC adresa je 03:00:01:a5:b6:c7, IP adresa 10.0.0.12, je připojen na rozhraní eth0 (i když v reálu není), stav spojení *permanent*. Pak vypište seznam sousedů a zkontrolujte, zda je v seznamu přidán záznam. Dále vypište seznam sousedů, jejichž připojení je ve stavu *permanent*. Potom nový záznam odstraňte a opět zkontrolujte, jestli záznam opravdu zmizel.
3. Zopakujte si celý postup aktivace síťového rozhraní (zkompletujte příkazy):
 - pokud je rozhraní aktivní, je třeba ho deaktivovat
 - nastavíme IP adresu (včetně délky prefixu)
 - nastavíme bránu
 - nastavíme adresu DNS serveru
 - aktivujeme rozhraní



B.4.4 Tunely

Mechanismus `iproute2` slouží také k vytváření IP/IP tunelů. Lze použít pro vytvoření VPN tunelu, zapouzdření IPv6 do IPv4 na cestě bez podpory IPv6, vytvoření mostu mezi více síťovými rozhraními na jednom zařízení patřícími do různých sítí, vytvoření vzdáleného mostu, apod.

Je možné pracovat s několika různými typy tunelů, ale předně musíme mít v jádře načten příslušný modul pro daný typ tunelu. Nejběžnější jsou

- *IP-IP tunely* (modul `ipip`) – dokážou zapouzdřit jen unicast IP pakety (tunely point-to-point),
- *GRE tunely* (modul `ip_gre`) – zapouzdřují také jiné typy paketů, a to spojení unicast i multicast, jsou poměrně hodně používány,
- *SIT tunely* (Simple Internet Transition, modul `ip6v6`) – k propojení IPv6 sítí přes IPv4 síť.

V příkazech pro práci s tunely se typ tunelu projevuje v povinném parametru `mode` (tedy mód tunelu), určuje způsob, jakým se má s transportovaným paketem zacházet (především jak a do čeho se má zapouzdřit). Používáme následující příkaz:



`ip tunnel ...` zabezpečený přenos, tunelování

`ip tunnel show mujtunnel` zobrazí informaci o vytvořeném tunelu s názvem `mujtunnel`, zobrazí se obvykle typ (mód) tunelu, vzdálená a místní IP adresa (konce tunelu), pak další zadané (nepovinné) vlastnosti jako název síťového rozhraní, ze kterého tunel vede, hodnota TTL pro pakety jdoucí do tunelu, hodnota TOS, apod. (vpodstatě se zobrazuje to, co se zadalo při vytvoření tunelu)

`ip -s tunnel show mujtunnel` kromě výše uvedených informací se zobrazí také statistika tunelu podobná výstupu příkazu `ip -s link show` (jde vlastně o stejný typ informace, jen se místo síťového rozhraní týká tunelu), viz výstup na straně 262


```
ip tunnel add mujtunnel mode ipip local 194.50.20.42 remote
195.84.152.140 ttl 32
```

vytvořili jsme nový tunel typu IP-IP se zadanou místní a vzdálenou adresou a hodnotou TTL

```
ip tunnel del mujtunnel
```

zrušili jsme tunel

```
ip tunnel change mujtunnel remote 195.35.84.15
```

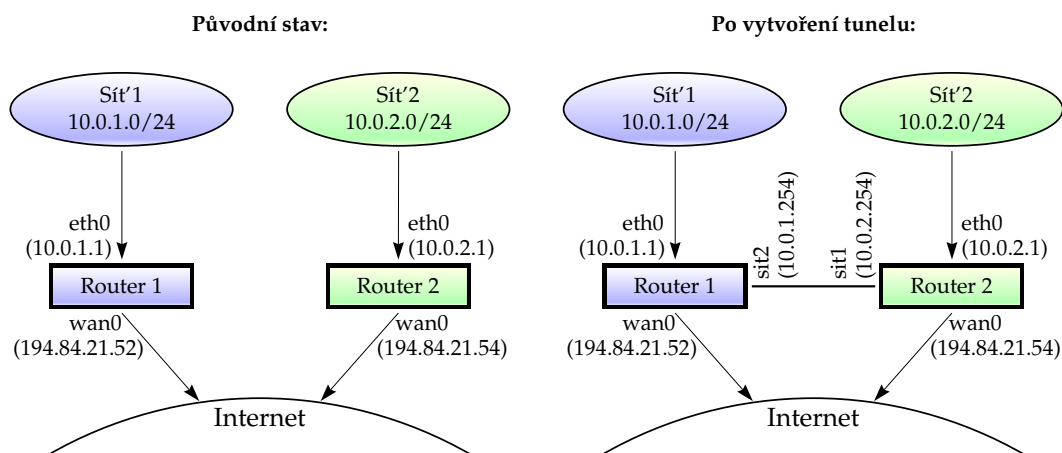
změna nastavení tunelu (změnili jsme vzdálenou adresu)

Abychom mohli tyto příkazy používat, musíme mít předně v jádře načten příslušný modul.



Příklad

Ukážeme si vytvoření GRE tunelu, který je zřejmě nejčastěji používán (protože dokáže zabalit cokoliv, zvládá i multicast a je všeobecně podporován prakticky v čemkoliv). Předpokládejme, že chceme propojit tunelem dvě vzdálené sítě (patřící stejné firmě), ve kterých jsou soukromé IP adresy. Rozsahy adres v těchto sítích by se neměly překrývat, aby nedocházelo k problémům při směrování.



Obrázek B.2: Vytvoření GRE tunelu

Situaci včetně adres máme znázorněnu na obrázku B.2 vlevo. Podle prefixů poznáme, že rozsahy adres v sítích se nepřekrývají. U routerů, přes které se sítě připojují do Internetu, vidíme také adresu viditelnou uvnitř (je z rozsahu adres vnitřní sítě) a adresu viditelnou vně (tu máme přidělenou ISP providerem). Aby bylo možné vybudovat funkční tunel, musíme všechny tyto adresy znát (zadávat se do příkazů).

Na routeru *Router1* zadáme tyto příkazy:

```
modprobe ip_gre
```

načteme do jádra modul pro podporu GRE tunelů

```
ip tunnel add sit2 mode gre local 194.84.21.52 remote 194.84.21.54 ttl 255
```

vytvoříme tunel, všimněte si vyšší hodnoty TTL (nevíme, přes kolik routerů tunel vede, tedy je lepší použít vyšší hodnotu)

```
ip link set sit2 up
```

aktivujeme nové virtuální rozhraní (tedy náš tunel)

```
ip addr add 10.0.1.254 dev sit2
```

přidělíme našemu konci tunelu IP adresu, měla by být viditelná v naší síti a neměla by být použita pro jiné zařízení (to si musíme pohlídat, případně upravit rozsahy adres)

```
ip route add 10.0.2.0/24 dev sit2
```

do směrovací tabulky přidáme pravidlo pro směrování do sítě za tunelem, bližší konec tunelu bude fungovat jako brána

Obdobně postupujeme na routeru *Router2*:

```
modprobe ip_gre    načteme do jádra modul pro podporu GRE tunelů
ip tunnel add sit1 mode gre local 194.84.21.54 remote 194.84.21.52 ttl 255
                   vytvoříme tunel
ip link set sit1 up   aktivujeme nové virtuální rozhraní (tedy náš tunel)
ip addr add 10.0.2.254 dev sit1   přidělíme našemu konci tunelu IP adresu
ip route add 10.0.1.0/24 dev sit1   do směrovací tabulky přidáme pravidlo pro směrování do sítě za tunelem
```

Na obrázku B.2 vpravo je znázorněna situace po vytvoření tunelu. Pak bychom měli vyzkoušet, zda jsou pakety doručovány správně, například na prvním routeru vyzkoušíme `ping` na druhý konec tunelu:

```
ping 10.0.2.154
```

Pokud nefunguje, je možné, že pakety byly zahozeny některým firewallem. Pokud jsou povoleny ICMP pakety a nemáme žádné restriktce na námi vytvořené virtuální rozhraní, měli bychom ještě zkontrolovat, zda mohou procházet GRE pakety (protokol číslo 47), například v chainu FORWARD bychom protokol 47 povolili takto (podrobnosti najdeme v samostatné sekci o firewallu v Linuxu):

```
iptables -A FORWARD -p 47 -j ACCEPT
```

Dále může být problém v zakázání čísla portu používaného tunelem, tedy pokud nepomůže povolení protokolu 47, můžeme vyzkoušet

```
iptables -A FORWARD -p tcp -dport 1723 -j ACCEPT
```

Případně v jiných tabulkách nebo s dalšími parametry, podle momentálního nastavení firewallu.



Měli bychom si uvědomit, že GRE tunely nejsou šifrovány.



Další informace:

- <http://www.root.cz/clanky/tuneluji-tunelujes-tunelujeme-jak-a-k-cemu/>
- <http://tier.cs.berkeley.edu/drupal/howto/ip-tunnel-using-gre-on-linux>
- <http://kovyrin.net/2006/03/17/how-to-create-ip-ip-tunnel-between-freebsd-and-linux/>
- <http://www.abclinuxu.cz/clanky/site/ipv6-konfigurace-site-tunely>
- IP-IP tunel: <http://fengnet.com/book/ICUNA/ch11lev1sec5.html>
GRE tunel: <http://fengnet.com/book/ICUNA/ch11lev1sec6.html>
- <http://www.informit.com/articles/article.aspx?p=29039&seqNum=3>
- <http://www.linuxfoundation.org/collaborate/workgroups/networking/tunneling>
- <http://www.policyrouting.org/iproute2.doc.html>



Poznámka:

V této sekci nejsou popsány všechny možnosti mechanismu `iproute2`. Existují například vnitřní příkazy pro monitoring, práci s multicast adresami, multicast směrování, práce s různými směrovacími protokoly včetně OSPF, atd.





Další informace:

Další informace o používání mechanismu `iproute2` (velmi užitečné tutoriály a jiné návody):

- <http://linux-ip.net/gl/ip-cref/>
- <http://linux-ip.net/html/index.html>
- <http://www.policyrouting.org/iproute2-toc.html>
- <http://lartc.org/howto/index.html>
- <http://www.faqs.org/docs/Linux-mini/Remote-Bridging.html>

A dále samozřejmě `man ip` použité v Linuxu nebo na Googlu.



B.5 Překlad názvů

B.5.1 Název počítače

 Také v Linuxu se pracuje s názvy počítačů. Předně je důležité vědět, jak zjistit název svého počítače. Pro tyto účely používáme příkaz `hostname` a jeho varianty:

`hostname` (bez parametrů) vypíše název našeho počítače; stejný údaj, jak víme, zjistíme také v souboru `/etc/sysconfig/network` nebo případně `/etc/hostname`, pokud existuje

`hostname novynazev` nastavíme název našeho počítače, potřebujeme vyšší přístupová oprávnění (tudíž použijeme příkaz `su` nebo `sudo`, podle toho, který naše distribuce podporuje)

`hostname -F soubor` nastavíme název našeho počítače, tento název je uložen v zadaném souboru, taktéž potřebujeme vyšší přístupová oprávnění

`domainname` (bez parametrů) vypíše název naší NIS domény (pozor, ne DNS; NIS je obdoba Active Directory pro UNIXové systémy, postupně nahrazován různými implementacemi LDAP), pokud ovšem máme NIS vytvořen

`dnsdomainname` (bez parametrů) vypíše název naší DNS domény (to je to, co v plném názvu počítače následuje za první tečkou), podobně `hostname -d`

B.5.2 Resolver a soubor `resolv.conf`

 Konfigurace *resolveru* (tj. části mechanismu překladu adres, která zadává dotazy, je na každém klientovi) je tedy v souboru `/etc/resolv.conf`. V tomto souboru mohou být záznamy několika typů:

- `nameserver adresa` doménový server, obvykle zadaný IP adresou, takových záznamů může být více (pak je ale důležité jejich pořadí, jako první dáme ten server, na který se chceme obracet nejčastěji, tj. obvykle nejbližší)
- `domain doména` název domény, zároveň s názvem našeho počítače tvoří plné jméno počítače, pokud nejsme v žádné síťové doméně, pak je zde záznam `domain localdomain`
- `search doména doména ...` seznam domén, které jsou při vyhledávání určitého názvu přidávány pro zkompletování plného jména (před odesláním dotazu DNS serveru), například pokud je v seznamu doména `fpf.slu.cz`, pak název uzlu `axpsu` lze doplnit na `axpsu.fpf.slu.cz`

- **options** volby volby pro konfiguraci resolveru (uvádějí se všechny na jednom řádku), např.
 - počet pokusů pro vyřízení DNS dotazu na každý známý DNS server se nastavuje takto:
`options attempts:3`
 (výchozí hodnota je 2, zvýšili jsme počet pokusů, zřejmě máme méně spolehlivé spojení k DNS serverům, obvykle není nutné měnit)
 - podporu IPv6 zapneme volbou `inet6`,
 - volba `no-check-names` se používá především tehdy, když máme v síti také počítače s nainstalovanými Windows; Windows totiž často porušují RFC 952, které v názvech dovoluje pouze písmena, číslice a pomlčky, a bez použití této volby by uzly s Windows s názvem obsahujícím nedovolené názvy nebyly dostupné, atd.

Pokud chceme použít více voleb, musíme je umístit na jeden řádek, například

```
options attempts:1 ipv6 no-check-names
```

Záznamy typu `domain` a `search` by teoreticky neměly být oba najednou použity (dá se říct, že `domain` je speciálním případem `search`, ve kterém zadáváme jen jednu doménu). Prakticky se vyskytovat mohou, ale při načítání konfigurace se bere v úvahu jen poslední uvedený z těchto záznamů.

V Linuxu se běžně používá DNS server *BIND* běžící jako démon `named`, jeho popis je však nad rámec těchto skript. Postup konfigurace včetně ukázkových konfiguračních souborů najdeme například v knize [10] ze seznamu literatury.

B.5.3 Testování DNS

 Podobně jako ve Windows, i v Linuxu máme k dispozici příkaz `nslookup`. Způsob využívání je podobný jako ve Windows:

```
nslookup www.google.com    takto zjistíme IP adresy zadaného serveru
```

```
nslookup 209.85.148.99     zjistíme jméno příslušející zadané IP adrese
```

Do interaktivního režimu se dostaneme jedním ze dvou způsobů:

- jednoduše zadáním příkazu bez parametrů,
- příkaz s jediným parametrem – názvem DNS serveru, před který napíšeme pomlčku, aby byl odlišitelný od názvu, který chceme přeložit.

Název DNS serveru zadáváme tehdy, když výchozí server neposkytuje autorizované odpovědi a nevíme jistě, zda v jeho cache jsou správné záznamy. Pokud už jsme v interaktivním režimu, nastavíme používaný DNS server příkazem `server adresa` (zadáme adresu DNS serveru).

Vnitřní příkazy v podstatě odpovídají tomu, co jsme si ukazovali u příkazu pro Windows, tedy například, překlady podobné jako v neinteraktivním režimu, zobrazení parametrů, určení typu dotazu na DNS server, zobrazení hostitelů (uzlů) v doméně, určení typu zobrazovaných hostitelů (výchozí je typ A) atd. Z interaktivního režimu se dostaneme příkazem `exit`.

 Podobný účel má příkaz `dig` (zkratka z Domain Information Gropher). Tento nástroj je u administrátorů oblíbenější než `nslookup`, zvláště pro účely testování nastavení DNS serverů, už proto, že jeho výstup je obsáhlejší a odpovídá údajům v DNS paketu.

```
dig www.root.cz
```

získáme vyčerpávající odpověď se všemi potřebnými informacemi, kromě žádané IP adresy (resp. více adres) například adresu DNS serveru, který odpověděl, jak dlouho trvala komunikace, ke každé adrese informaci o třídě záznamu (kde je platný), typ záznamu, atd.

`dig -x 91.213.160.118` reverzní dotaz, chceme doménový název k zadané IP adrese
`dig www.root.cz +short` krátká odpověď ve stylu `nslookup` (tj. vypíše se pouze IP adresa)
`dig google.com ANY` chceme všechny typy záznamů týkající se daného serveru (výchozí jsou záznamy typu A, tj. IPv4 adresy)
`dig seznam.cz NS +short` vypíšou se doménová jména DNS serverů v zadané doméně, jejich IP adresy můžeme zjistit následným dotazem podle vypsaných doménových jmen



Příklad

Vyzkoušíme práci s příkazem `dig`:

```
sarka@notebook:~$ dig www.root.cz
```

```
; <<>> DiG 9.7.1-P2 <<>> www.root.cz
;; global options: +cmd
;; Got answer:
;; ->>HEADER<<- opcode: QUERY, status: NOERROR, id: 55084
;; flags: qr rd ra; QUERY: 1, ANSWER: 2, AUTHORITY: 2, ADDITIONAL: 2

;; QUESTION SECTION:
;www.root.cz.      IN A

;; ANSWER SECTION:
www.root.cz.      5 IN CNAME root.cz.
root.cz.          5 IN A~91.213.160.118

;; AUTHORITY SECTION:
root.cz.          5 IN NS ns.iinfo.cz.
root.cz.          5 IN NS ns6.adminit.cz.

;; ADDITIONAL SECTION:
ns6.adminit.cz.   5 IN A~81.91.85.230
ns6.adminit.cz.   5 IN AAAA 2001:1568:b:149::1

;; Query time: 18 msec
;; SERVER: 192.168.184.2#53(192.168.184.2)
;; WHEN: Mon Jul 11 07:09:03 2011
;; MSG SIZE rcvd: 152
```

Ve výpisu jsou tyto položky:

- v záhlaví výpisu jsou informace o tom, jakým způsobem byl příkaz volán a kterou máme verzi, dále část s početními údaji (čísla odpovídají počtům záznamů uvedených v následujících sekcích), příznaky apod.,
- sekce „Question section“ obsahující údaje dotazu (zde doménová adresa, třída „INternet“, typ záznamu „A“)
- sekce „Answer section“ obsahuje údaje z odpovědi (v záhlaví máme uvedeny dva záznamy v odpovědi, proto zde budou dva řádky obsahující doménový název, TTL, třídu, typ záznamu a hodnotu záznamu),
- v sekci „Authority section“ máme seznam doménových adres DNS serverů, které o dotazovaném názvu poskytují autoritativní odpověď (pokud získaným informacím nedůvěřujeme, můžeme se zeptat některého z těchto serverů),
- v následující sekci najdeme IP adresy DNS serverů uvedených v předchozí sekci,

- výpis končí statistickými informacemi.

Ted' podobně zjistíme informace o hostiteli `mail.google.com`, ale zeptáme se jiného DNS serveru. Název DNS serveru musíme uvést symbolem „zavináče“, aby bylo zřejmé, že se nejedná o název, který má být přeložen, za zavináčem můžeme zadat doménovou nebo IP adresu. Zde se (shodou okolností) dotážeme DNS serveru poskytovaného Googlem:

```
sarka@notebook:~$ dig @8.8.4.4 mail.google.com

; <<>> DiG 9.7.1-P2 <<>> @8.8.4.4 mail.google.com
; (1 server found)
;; global options: +cmd
;; Got answer:
;; ->>HEADER<<- opcode: QUERY, status: NOERROR, id: 20135
;; flags: qr rd ra; QUERY: 1, ANSWER: 5, AUTHORITY: 0, ADDITIONAL: 0

;; QUESTION SECTION:
;mail.google.com. IN A

;; ANSWER SECTION:
mail.google.com. 86399 IN CNAME googlemail.l.google.com.
googlemail.l.google.com. 299 IN A~74.125.232.246
googlemail.l.google.com. 299 IN A~74.125.232.248
googlemail.l.google.com. 299 IN A~74.125.232.247
googlemail.l.google.com. 299 IN A~74.125.232.245

;; Query time: 56 msec
;; SERVER: 8.8.4.4#53(8.8.4.4)
;; WHEN: Mon Jul 11 07:21:14 2011
;; MSG SIZE rcvd: 124
```

Je zřejmé, že odpověď jsme sice dostali, ale vyžádali jsme si ji od serveru, který není pro daný záznam autoritativní (údaje bere ze své cache nebo se dotazuje jinde).

Dále vyzkoušíme zkrácený výpis (použijeme parametr `+short`), chceme DNS servery v doméně `seznam.cz`.

```
sarka@notebook:~$ dig seznam.cz NS +short
```

```
ns.seznam.cz
ms.seznam.cz
```



Příklad

Příkaz `dig` můžeme také použít pro zjištění celé trasy našeho dotazu:

```
sarka@notebook:~$ dig seznam.cz +trace
```

```
; <<>> DiG 9.7.1-P2 <<>> seznam.cz +trace
;; global options: +cmd
. 5 IN NS i.root-servers.net.
. 5 IN NS j.root-servers.net.
. 5 IN NS k.root-servers.net.
. 5 IN NS l.root-servers.net.
. 5 IN NS m.root-servers.net.
. 5 IN NS a.root-servers.net.
. 5 IN NS b.root-servers.net.
. 5 IN NS c.root-servers.net.
. 5 IN NS d.root-servers.net.
```

```
. 5 IN NS e.root-servers.net.
. 5 IN NS f.root-servers.net.
. 5 IN NS g.root-servers.net.
. 5 IN NS h.root-servers.net.
;; Received 272 bytes from 192.168.184.2#53(192.168.184.2) in 12 ms

cz. 172800 IN NS f.ns.nic.cz.
cz. 172800 IN NS d.ns.nic.cz.
cz. 172800 IN NS a.ns.nic.cz.
cz. 172800 IN NS c.ns.nic.cz.
cz. 172800 IN NS b.ns.nic.cz.
;; Received 334 bytes from 192.5.5.241#53(f.root-servers.net) in 10 ms

seznam.cz. 18000 IN NS ms.seznam.cz.
seznam.cz. 18000 IN NS ns.seznam.cz.
;; Received 93 bytes from 194.0.14.1#53(c.ns.nic.cz) in 43 ms

seznam.cz. 300 IN A~77.75.72.3
;; Received 43 bytes from 77.75.77.77#53(ms.seznam.cz) in 9 ms
```



Pokud nezadáme DNS server, kterému má být zaslán dotaz, příkaz `dig` použije DNS servery uvedené v souboru `/etc/resolv.conf`.

B.5.4 Zjistění informací o doméně

 Pokud chceme zjistit informace o určité doméně, použijeme příkaz `whois`. Tento příkaz nám vypíše veškeré dostupné informace o doméně – kontakt na vlastníka domény, doba, na kterou je doména registrována apod. Příkaz se hodí například tehdy, když chceme registrovat vlastní doménu a chceme si ověřit, zda už není někým registrována, a nebo pokud z některé domény přichází malware a chceme upozornit vlastníka, aby „si udělal pořádek“.

`whois koupaliste.cz` zjistíme údaje o zadané doméně

`whois -r koupaliste.cz` výpis bude odlišný, ale základní informace budou shodné (uvedeným přepínačem jsme si vyžádali informace z WHOIS databáze RIPE obsahující především evropské servery)

V manuálové stránce příkazu bychom zjistili přepínače pro další možné WHOIS databáze (pozor, manuálové stránky různých UNIXových systémů se liší v tom, jak moc jsou podrobné při komentování seznamu přepínačů tohoto příkazu, pak je dobré použít manuálové stránky na Internetu).



Další informace:

Další informace o doménách v Linuxu:

- <http://www.madboa.com/geek/dig/>
- http://www.brennan.id.au/08-Domain_Name_System_BIND.html



Úkoly

1. Zjistěte název svého počítače. Dále zjistěte, jaké DNS servery používáte.
2. Zobrazte příkazem `dig` podrobnosti o serveru `seznam.cz`, a to všechny typy DNS záznamů (tj. ANY). Srovnajte s výpisem pomocí příkazu `nslookup`.

3. To, že se ptáme jiného DNS serveru než je ten od našeho poskytovatele, nemusí znamenat, že dostaneme komplexnější informace. Vyzkoušejte dotaz na doménu `root.cz` postupně bez zadání DNS serveru a pak se zadáním některého DNS serveru (například `ns.seznam.cz` nebo `8.8.4.4`). Pokud příkaz zkoušíte v síti některé vysoké školy, kde se učí Linux, tak bude výstup od výchozího serveru komplexnější.
4. Zjistěte údaje o několika doménách druhé úrovně podle vlastního výběru (například `seznam.cz`, `slu.cz`, `apod.`).



B.6 Testování a statistiky

B.6.1 Základní nástroje

 Příkaz `ping` slouží ke stejným účelům jako ve Windows:

`ping www.google.com` v pravidelných intervalech odesílá cílovému zařízení ICMP pakety, po přerušení klávesou `[Ctrl+C]` ukončí zasílání a vypíše souhrnnou statistiku (kolik paketů bylo odesláno, přijato, podíl ztracených, dále odezvu cílového počítače – nejrychlejší, nejpomelejší, střední hodnota), cílový počítač můžeme zadat také IP adresou

`ping -c 4 www.google.com` chová se stejně jako ve Windows, tedy zadáme počet zasílaných ICMP paketů (bez tohoto parametru by byly zasílány tak dlouho, dokud bychom nestiskli `[Ctrl+C]`)

`ping -c 1 -R www.google.com` parametr `-R` (pozor, velké písmeno) znamená, že se vypíší všechny routery na cestě podobně jako například u `traceroute` (ale zde získáme méně podrobný výpis)

V manuálové stránce bychom zjistili další možné parametry.

 Existuje také příkaz `arping`, který má podobný účel jako `ping`, ale používá jiný typ paketů (ARP Request) a je určen především pro testování v lokální síti.

`arping -f -c 4 -I eth0 209.85.143.99` zadanému cíli (poslední parametr) přes zadané rozhraní odešle maximálně čtyři ARP pakety, parametr `-f` určuje, že se nemusejí deslat všechny 4, ale pouze tolik, než vzdálený systém odpoví

 Zatímco ve Windows máme příkaz `tracert`, v UNIXových systémech včetně Linuxu se setkáme s příkazem `traceroute`.

`traceroute www.google.com` vypíše trasu k zadanému zařízení (zadáváme doménovou nebo IP adresu), použije UDP pakety (pouze metoda s využitím UDP paketů je použitelná kterýmkoliv uživatelem)

`traceroute -I 209.85.143.99` podobně, ale použije zprávy ICMP Echo Request (přepínač je velké „i“), ekvivalentem je příkaz `tracert` (s ním se setkáme i u Windows)

`traceroute -6 209.85.143.99` vynucení použití IPv6 (také `traceroute6`), k vynucení IPv4 slouží přepínač `-4`

Tento příkaz má ještě několik dalších přepínačů (viz manuálové stránky) včetně u těchto příkazů běžného `-n` zakazujícího mapování IP adres na doménové názvy, určení maximální hodnoty TTL, atd.

 Podobný je příkaz `tracpath` (v některých distribucích ho však nenajdeme). Slouží především ke zjištění hodnot MTU na cestě:

`tracert 209.85.143.99` vypíše údaje o cestě k zadanému uzlu sítě (ke každému úseku zjistíme hodnotu odečítaného TTL na tomto úseku, časový údaj a hodnotu MTU (Path MTU))

`tracert -n 209.85.143.99` místo doménových se vypisují IP adresy

 Samozřejmě nesmí chybět příkaz `netstat`. Parametry jsou obdobné těm ve Windows (odkud se tento příkaz asi do Windows dostal), tedy si zde uvedeme jen několik dalších motivačních příkladů:

Příklad

Zobrazíme seznam otevřených portů, postupně: výpis numerických IP adres (n), výpis UDP portů (u), TCP portů (t), všechny údaje (a, naslouchající i nenaslouchající), vypsání PID procesů, které port používají (p). Výpis příkazu následuje.

```
sarka@notebook:/home/sarka# netstat -nutap
```

```
Aktivní Internetová spojení (servery a~navázaná spojení)
Proto Přích-F Odch-F Místní Adresa      Vzdálená Adresa
Stav      PID/Program name
tcp        0      0 127.0.0.1:2208      0.0.0.0:*
LISTEN     2725/hpiod
tcp        0      0 127.0.0.1:49702     0.0.0.0:*
LISTEN     2728/python
tcp        0      0 0.0.0.0:113         0.0.0.0:*
LISTEN     3038/inetd
tcp        0      0 127.0.0.1:631       0.0.0.0:*
LISTEN     2850/cupsd
tcp        0      0 127.0.0.1:25        0.0.0.0:*
LISTEN     3018/exim4
tcp        0      0 10.0.0.2:54255      74.125.232.216:443
SPOJENO    3517/firefox-bin
tcp        0      0 10.0.0.2:54260      74.125.232.216:443
SPOJENO    3517/firefox-bin
tcp        0      0 10.0.0.2:54265      74.125.232.216:443
SPOJENO    3517/firefox-bin
tcp        0      0 10.0.0.2:39412      74.125.232.213:443
SPOJENO    3517/firefox-bin
tcp        0      0 10.0.0.2:32819      74.125.232.194:80
SPOJENO    3517/firefox-bin
tcp        0      0 127.0.0.1:47622     127.0.0.1:111      TIME_WAIT   -
tcp        1      0 10.0.0.2:55360      91.189.90.132:80
CLOSE_WAIT 3597/python
udp        0      0 0.0.0.0:32768       0.0.0.0:*
          2929/avahi-daemon:
udp        0      0 0.0.0.0:68          0.0.0.0:*
          3152/dhclient
udp        0      0 0.0.0.0:714         0.0.0.0:*
          3082/rpc.statd
udp        0      0 0.0.0.0:5353        0.0.0.0:*
          2929/avahi-daemon:
udp        0      0 0.0.0.0:631         0.0.0.0:*
          2850/cupsd
```

Vidíme, že (aktivně) komunikuje především Firefox, další démony spíše naslouchají. Najdeme tam například nechvalně známý *Avahi* (provádí tzv. Zero-conf automatické nastavení parametrů sítě, což je odborníky hodně kritizováno), *cupsd* (tiskový subsystém), *hpiod* (démon pro podporu tisku kompatibilní s HP zařízeními), *exim4* (síťová podpora pro mail klienty, v současné době se spíše doporučuje použít

postfix), *inetd*,¹² *dhclient* (DHCP klient), atd. (výpis byl záměrně mírně zkrácen).



Příklad

V Linuxu si pomocí příkazu **netstat** můžeme prohlédnout tabulku síťových rozhraní tak, jak ji vidí jádro:

```
sarka@notebook:/home/sarka# netstat -i
```

Tabulka rozhraní v jádru

Iface	MTU	Met	RX-OK	RX-ERR	RX-DRP	RX-OVR	TX-OK	TX-ERR	TX-DRP	TX-OVR	Flg
eth0	1500	0	2218	0	0	0	1919	0	0	0	BMRU
lo	16436	0	48	0	0	0	48	0	0	0	LRU

Co je to MTU, už víme. Následuje metrika, údaje pro příchozí (RX) a odchozí (TX) pakety, kdy se vypisuje počet úspěšně přijatých/odeslaných, chybných, zahozených. Poslední sloupec obsahuje jednopísmenné příznaky:

- B: přijímá broadcast pakety,
- M: zapnut promiskuitní režim,
- R: rozhraní je používáno (běží, running)
- U: rozhraní je aktivní (up)
- L: loopback



Příkaz **netstat-nat** je zdánlivě podobný příkazu **netstat**, ale není tentýž:

netstat-nat výpis adres překládaných mechanismem NAT (pozor, před pomlčkou není mezera, opravdu se jedná o celý příkaz bez jakýchkoliv parametrů)

netstat-nat -n před druhou pomlčkou už mezera je, význam příkazu je stejný jako u předchozího, ale místo názvů počítačů se budou vypisovat IP adresy (tj. bez DNS překladu)

B.6.2 Pokročilé testování

Pro pokročilé testování strojů a celé sítě existuje mnoho velmi užitečných nástrojů. V tomto předmětu je bohužel nestačíme probrat. Namátkou:

- tcpdump
- nmap
- ettercap

B.7 Firewall v Linuxu



V linuxových jádrech (od verze 2.4) najdeme firewall *NetFilter*, což je obousměrný stavový firewall (provádí funkce SPI).¹³ *NetFilter* dokáže filtrovat pakety podle údajů v hlavičkách protokolů TCP, UDP,

¹² *Inetd* má podobnou funkci jako ve Windows *svchost.exe*, může hostit demony. Na rozdíl od *svchost* je zde pro demony „pohostinství“ dobrovolné, mohou běžet samostatně. *Inetd* jednoduše naslouchá na příslušných portech a když zjistí pokus o komunikaci s démonem, jehož hostí, jeho kód spustí ve svém vlastním procesu, není třeba spouštět démona navíc (tj. místo více různých démonů běží jen jeden, *inetd*, a příslušné služby se spouštějí až na vyžádání).

¹³ Stavové firewally dokážou pracovat samozřejmě nejen se stavovými, ale i s nestavovými protokoly (jako je například UDP), přestože se v některých publikacích můžeme dočíst, že ne.

IP, ICMP a také dalších uvedených v souboru `/etc/protocols`, je použitelný pro mnoho činností souvisejících se zpracováním paketů, včetně mechanismu NAT.

Jedná se o modul jádra, ke kterému se nepřístupuje přímo, ale přes obslužný program `iptables`.¹⁴ Takže pozor – firewall je NetFilter, obslužný program běžící v uživatelském režimu je `iptables`.

B.7.1 Princip firewallu

 Základem je *tabulka* (table), přes tabulky prochází několik řetězců *chain* (čti: [čejn]), které již obsahují pravidla uplatňovaná na pakety (tato pravidla jsou zřetěžená a na paket se uplatňují postupně, dokud se nenajde „pasující“ pravidlo, proto se tomu říká chain).

Tabulka obvykle určuje, co konkrétně se má v dané části řetězce dít (například jen filtrování, nebo překlad adres či změna některých dalších údajů v záhlaví procházejícího paketu), chain obvykle určuje, pro jaký typ paketů se to má dít (například pro příchozí pakety, odchozí, průchozí u routeru, připravené na směrování, zpracované směrováním, apod.).

Příklad

Pokud jsme na routeru se zapnutým přeposíláním (forwarding), je příchozí paket nejdřív zařazen do chainu *PREROUTING* (tj. pro činnosti, které se mají provést ještě před směrováním), v něm projde určenými tabulkami (například tabulkou *nat*, pokud se provádí překlad adres).

Pak proběhne směrování, které určí, jestli se tento paket má přeposílat (tedy je přerazen do chainu *FORWARD* pro přeposílané pakety) nebo je určen přímo pro toto zařízení (půjde do chainu *INPUT* pro příchozí pakety), a v rámci daného chainu projde stanovenými tabulkami. Přeposílaný paket pak ještě projde chainem *OUTPUT*, protože se stává odchozím paketem.

Takže chain vlastně určuje „kategorii“ paketu, tabulky pravidla pro zacházení. Pokud přes tuto tabulku procházejí dva pakety takové, že každý je v jiném chainu, bude se s každým zacházet trochu jinak, třebaže typ prováděných operací bude stejný (filtrování je ve více chainech, ale jinak se filtrují příchozí, jinak průchozí a jinak odchozí pakety).

 Standardně existují tyto tabulky (v příkazu `iptables` se před název tabulky dává prepínač `-t`):

- *filter* (výchozí)
- *mangle* (pro transformace – upravujeme záhlaví, hodnoty TTL, pole TOS/QoS, apod.)
- *nat* (pro mechanismus NAT)
- *raw* a *conntrack* jsou pomocné mechanismy označování (marking) paketů pro pozdější zpracování a stavového filtru (SPI), obvykle je třeba jejich funkcionalitu do jádra doplnit (načíst moduly)
- *security* je přidávána bezpečnostními moduly jádra jako je například SELinux.

 Přes každou tabulku tedy vede několik *chainů* (řetězců), a to:

- Tabulka *filter* obsahuje standardně tyto chainy:
 - chain *INPUT* se uplatní na pakety, které do systému přicházejí,
 - chain *OUTPUT* se uplatní na pakety, které ze systému odcházejí,

¹⁴Ve starších jádrech, se kterými se (doufejme) už nesetkáme, se používal firewall *IPChains* se stejnojmenným obslužným programem (`ipchains`).

- chain **FORWARD** je určen pro pakety, které nejsou vytvořeny uvnitř systému a ani pro ně nejsou určeny, předávají se jen mezi rozhraními systému (průchozí pakety), typicky v zařízení, které slouží jako směrovač,

pokud jde paket do chainu **FORWARD**, nepadne už do žádného jiného chainu.

- Tabulka *nat* obsahuje standardně tyto chainy:
 - chain **PREROUTING** pro příchozí pakety, na které se má použít **DNAT** (tj. má se přeložit cílová adresa, a to ještě před vlastním směrováním),
 - chain **POSTROUTING** pro odchozí pakety, na které se má použít **SNAT** (překládá se zdrojová adresa, až po směrování),
 - chain **OUTPUT** pro odchozí pakety, kdy se má modifikace provést ještě před dalším zpracováním.
- Tabulka *mangle* obsahuje standardně chainy pojmenované jako všechny, které jsou uvedeny u předchozích tabulek (**INPUT**, **OUTPUT**, **FORWARD**, **PREROUTING**, **POSTROUTING**) s tím rozdílem, že v nich najdeme pravidla modifikující i jiná pole záhlaví než jen ta adresová.
- Tabulka *security* může následovat po tabulce *filter*, obsahuje chainy **INPUT**, **OUTPUT** a **FORWARD** a dokáže podle potřeby pakety označovat a přesměrovávat k jiným modulům jádra.

V tabulkách si můžeme vytvořit další chainy (uživatelsky definované), nemusíme zůstat jen u standardních, případně nové tabulky. Nové tabulky se však často přidávají jako další moduly jádra.

Na obrázku B.3 je znázorněna dráha paketu mezi jednotlivými tabulkami a chainy. Všimněte si, že preposílané (forwarded) pakety nedorazí do žádného chainu typu **INPUT** ani **OUTPUT**.

U pravidel v chainu záleží na pořadí. Každé pravidlo má své pořadové číslo. Nová pravidla můžeme do chainu vkládat na začátek či na konec chainu, anebo na konkrétní pozici zadanou pořadovým číslem.

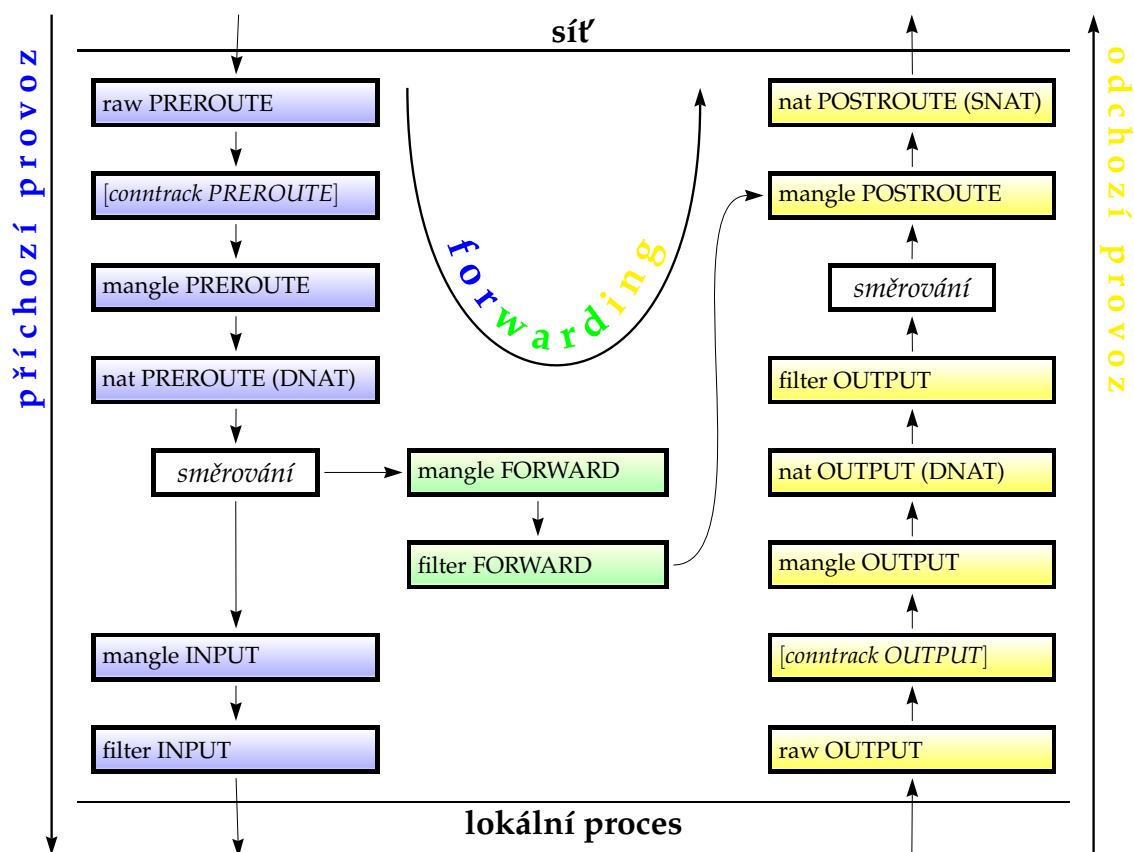
 Program (ve skutečnosti služba, démon) **iptables** umožňuje pracovat s chainy (vytvořit nebo zrušit chain, nastavit výchozí politiku) a s jejich obsahem (přidávat a odstraňovat pravidla v chainu, vypsat seznam pravidel).

V různých distribucích se však můžeme setkat i s dalšími obslužnými programy, například **uwf** (Uncomplicated FireWall) s grafickým rozhraním **gufw**, **shorewall**, **sanewall**, **apf-firewall** a další. Často jde spíše o nástavby **iptables**.

B.7.2 Základní vnitřní příkazy a parametry pro filtrování

 Ve vnitřních příkazech se používají tzv. *targety* určující akci, která se má s určeným paketem provést. Můžeme použít tyto targety:

- **ACCEPT** akceptuj, paket má povoleno projít, obvykle v tabulce *filter*,
- **DROP** zahod' bez odezvy (tiché zahození), taktéž v tabulce *filter*,
- **REJECT** odmítni (s odezvou), odešle ICMP zprávu o odmítnutí (přesněji – o nedostupnosti adresy) uzlu, od kterého paket přišel, pro tabulku *filter*,
- **DNAT** a **SNAT** patří k tabulce *nat*, účel je zřejmý – akcí je překlad cílové nebo zdrojové adresy na jinou IP adresu, používáme při překladu statických adres,
- **MASQUERADE** také patří k tabulce *nat* chainu **POSTROUTING**, u odchozích paketů překládá celý rozsah vnitřních (zdrojových) adres na jednu vnější (slouží ke skrývání lokálních adres), na rozdíl od **SNAT** bývá adres více a častěji se používá pro **PAT** (překlad portů), používáme pouze tehdy, když potřebujeme překládat dynamické soukromé adresy,



Obrázek B.3: Dráha paketu mezi výchozími tabulkami a chainy v Netfilteru

- TOS a TTL patří k tabulce *mangle*, mění pole TOS (type of service) a TTL v záhlaví paketu,
- MARK a CONNMARK se v tabulce *mangle* používají pro přidávání značek do záhlaví paketů,
- LOG umožňuje logovat informace ze záhlaví paketu, používá se obvykle s programem *syslog*,
- RETURN pokud jsme během vyhodnocování pravidel odskočili do jiného chainu, tento target nás vrátí zpět,

atd. Jednotlivé targety obvykle patří ke konkrétní tabulce, tedy když přidáme do jádra další modul s tabulkou, přidají se i další použitelné targety. Za targety se považují i uživatelsky definované chainy.

Příkaz `iptables` se používá s nejméně jedním parametrem (obvykle více), což bývá určení tabulky, vnitřní příkaz nebo parametr. Obvykle (ne však vždy) zadáváme tabulku, se kterou chceme v příkazu pracovat, k tomu použijeme přepínač `-t`, například `-t nat`. Zde se podíváme na vnitřní příkazy a parametry, které se používají při základním filtrování, v dalších podsekcích projdeme další (překlady adres, stavový firewall, značení paketů).

 Obvykle tedy používáme některý vnitřní příkaz (pozor, velká písmena) představující:

`-P` nebo `-policy` určuje výchozí zásadu (politiku) pro chain, ve kterém chceme pracovat (pokud na paket nebude pasovat žádné pravidlo chainu, použije se tato politika), za přepínačem následuje název chainu a dále *target* určující akci, která se má provést, výchozí politiku lze určit pouze pro předdefinované chainy, nikoliv pro uživatelsky definované; například

```
iptables -P INPUT -j DROP,
```

`-L` nebo `-list` vypisuje seznam všech pravidel v zadaném chainu, například

`iptables -L` (výpis bývá dlouhý dokonce i na desktopu, všechny tabulky)

`iptables -t filter -L -n`

`iptables -t filter -L -nv`

`iptables -t filter -L INPUT -n`

první příkaz vypíše všechna pravidla tabulky *filter*, nechává IP adresy (nepřekládá na doménové, díky poslednímu parametru), druhý příkaz udělá totéž, ale ve verbose („upovídáném“) módu, tudíž se dozvíme více, třetí příkaz vypíše pouze pravidla v chainu INPUT,

- A nebo **-append** připojí nové pravidlo na konec chainu, následuje název chainu a specifikace pravidla
- I nebo **-insert** připojí nové pravidlo na začátek chainu (pokud nezadáme žádný index) nebo na zadaný index (číslo pravidla můžeme napsat za název chainu), dto., první index je 1,
- R nebo **-replace** přepíše existující pravidlo na daném indexu,
- D nebo **-delete** odstraní pravidlo z chainu, následuje název chainu a buď číslo pravidla nebo jeho specifikace,
- F nebo **-flush** následované názvem chainu tento chain vyprázdní (vymaže všechna jeho pravidla)
- Z nebo **-zero** vynuluje všechny čítače v tabulce, obvykle se používá zároveň s příkazem -L, abychom viděli stav čítačů před jejich vynulováním, tedy například
`iptables -t filter -LZ`
- N nebo **-new-chain**, -X nebo **-delete-chain** pro vytvoření nebo odstranění chainu, následuje vždy název chainu,
- E nebo **-rename-chain** pro přejmenování chainu, následuje původní a nový název chainu.

 *Parametry* slouží k upřesnění příkazu (určují, podle čeho se pakety mají filtrovat), na rozdíl od příkazů jsou pro ně určena malá písmena, používáme především:

- p nebo **-protocol** určí protokol, například `-p tcp`, negace s vykřičníkem, např. `-p ! udp`,
- dport nebo **-sport** dovoluje k protokolu podle předchozí odrážky přidat konkrétní číslo portu (zdrojového nebo cílového, také se dá označit názvem příslušného aplikačního protokolu), například
`iptables -t filter -A OUTPUT -p tcp -dport 80 -j DROP`
`iptables -t filter -A OUTPUT -p tcp -dport http -j DROP`
 (obojí: zablokovali jsme odchozí tcp spojení přes http port)
`iptables -t filter -A FORWARD -p tcp -dport imap -j ACCEPT`
`iptables -t filter -A FORWARD -p tcp -dport pop3 -j ACCEPT`
 (povolili jsme přeposílání paketů se stahovanou poštou, aplikační protokol IMAP a POP3)
- icmp-type můžeme použít filtrování pro určitý typ ICMP zprávy, například
`iptables -t filter -A FORWARD -p icmp -icmp-type 8 -j DROP`
 zahazujeme všechny zprávy ICMP Echo Request (tím zablokuje odpovědi na ping)
- syn v TCP záhlaví je nastaven příznak SYN (samotný bez příznaku ACK), což znamená, že zřejmě jde o paket, který navazuje nové spojení (ne nutně, také je důležité otestovat stav spojení, viz dále o SPI)
- s nebo **-source** nebo **-src** a **-d** nebo **-destination** nebo **-dst** je zdrojová a cílová adresa, u sítě píšeme i délku prefixu, například `-s 192.89.120.0/24`, opět můžeme použít vykřičník pro negování (tj. všechny adresy kromě této),

- i nebo -in-interface zadává vstupní rozhraní, používá se v chainech INPUT, FORWARD a PRE-ROUTING, například -i eth0, můžeme také zadávat negaci, například -i ! eth4 znamená pakety ze všech síťových rozhraní kromě eth4,
- o nebo -out-interface podobně pro výstupní rozhraní, na které je paket směřován (v chainech OUTPUT, FORWARD a POSTROUTING), například

```
iptables -t filter -A OUTPUT -o eth1 -j ACCEPT
```

 veškerý ochozí provoz (neodpovídající jiným pravidlům) odcházející přes port eth1 bude povolen

 Další parametry slouží k upřesnění činnosti, která se má s paketem provést:

- j nebo -jump slouží k zadání provedení targetu (viz výše), který za tímto přepínačem následuje, znamená tedy odskok (jump) k provedení targetu, kterým může být některá akce (ACCEPT, DROP, SNAT apod.), a *nebo* název uživatelsky definovaného chainu, do kterého takto „odskočíme“ s tím, že se po procházení uživatelského chainu můžeme vrátit zpět (tj. odskok se zapamatováním adresy), například

```
iptables -t filter -A INPUT -p tcp -j tcppakety
```

 (pokud jde o TCP paket, přesuneme se do chainu *tcppakety*; pokud v tom chainu nenajdeme žádné vyhovující pravidlo, vrátíme se zpět do chainu INPUT a pokračujeme následujícími pravidly, uživatelský chain *tcppakety* byl předem vytvořen příkazem -N)
- g nebo -goto je pro uživatelsky definované chainy podobné jako -j, s tím rozdílem, že se před odskokem jinam nezapamatuje adresa momentálního chainu (přesněji – pokud už byla předtím některá adresa zapamatována pomocí -j, nepřepíše se jinou adresou, tudíž při případném návratu během vyhodnocování jednoho paketu se nevrátíme do chainu, ze kterého odcházíme, ale bude přeskočen), například

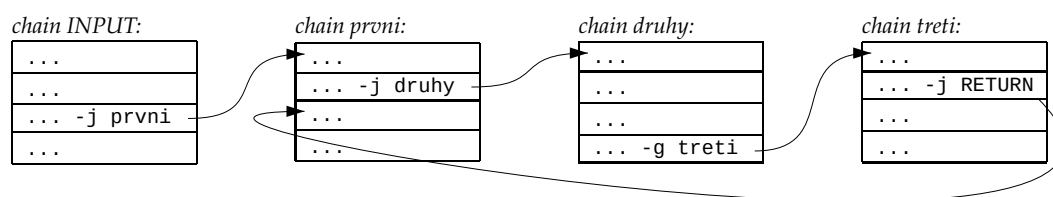
```
iptables -t filter -A INPUT -p tcp -g tcppakety
```

 podobně jako předchozí, ale pokud v uživatelském chainu nenalezne vhodné pravidlo, už se do INPUT nevrací (ovšem v *tcppakety* může být také parametr -j nebo -g směřující vyhodnocování paketu jinam).

Parametr -j tedy lze použít pro odskok do jiného (obvykle uživatelsky definovaného) chainu. Zpět se vrátíme buď tehdy, když následný chain neobsahuje vhodné pravidlo, a nebo lze návrat vynutit targetem RETURN (příkaz umístíme na konec vlastního chainu):

```
iptables -A vlastni_chain -j RETURN
```

(zde by bylo vhodné přidat ještě nějakou filtrovací podmínku, při jejímž splnění dojde k návratu, jinak pravidlo nemá moc smysl).



Obrázek B.4: Rozdíl mezi parametry -j a -g ve firewallu

 Zbývá ještě několik pomocných přepínačů používaných u příkazu -L:

`-v` nebo `-verbose` zapíná verbose („ukecaný“) mód, chceme delší výpis s více informacemi,
`-n` nebo `-numeric` použijeme, když chceme ve výpisu IP adresy místo doménových, výpis je pak celkově rychlejší a snadněji se hromadně zpracovává,

další bychom našli v manuálové stránce příkazu: `man 8 iptables`

Výše uvedený výčet příkazů a parametrů není zdaleka úplný, podrobnější informaci bychom našli opět v manuálové stránce příkazu, případně na odkazech na konci této sekce.



Poznámka:

Měli bychom si uvědomit, že zpracovávání paketu v tabulce končí v okamžiku, kdy NetFilter najde první vyhovující pravidlo s cílovou akcí akceptující nebo zahazující (například ACCEPT nebo DROP). Pokud je v tomto pravidle akce DROP nebo jiná „zahazovací“, paket bude okamžitě zahozen, i kdyby třeba hned následující pravidlo tento paket povolovalo (k dalšímu pravidlu se nedostaneme).



Příklad

Pár ukázek práce s `iptables`:

```
/sbin/iptables -t filter -L    výpis pravidel v tabulce filter
/sbin/iptables -L -v -n -line-numbers    podrobný výpis (statistika) všech tabulek, s IP adresami (n),
                                         v chainech jsou řádky číslovány
/sbin/iptables -P INPUT -j DROP    stanovíme výchozí politiku pro chain INPUT, ve které řekneme, že
                                         vše příchozí se má zahodit (samozřejmě musíme kromě jiného přidat před toto pravidlo také pravidla,
                                         která upřesní, co se zahazovat nemá)
/sbin/iptables -P OUTPUT -j ACCEPT    všechno odchozí povolíme, propustíme dál (opět bychom měli
                                         uvažovat, jestli nepřidáme „výjimky“)
/sbin/iptables -A INPUT -i lo -j ACCEPT    povolíme to, co probíhá mezi lokálními procesy (komunikují
                                         přes sockety na loopbacku)
/sbin/iptables -A INPUT -p tcp -syn -j DROP    budou zahozeny všechny příchozí pakety, které mají
                                         v TCP záhlaví nastaven pouze SYN bit (to je vždy první paket spojení, tedy znemožníme navazování
                                         spojení zvnějšku)
/sbin/iptables -A INPUT -p icmp -icmp-type echo-request -j ACCEPT    povolili jsme pakety s ICMP
                                         zprávou č. 8 (Echo Request) z vnějšku, můžeme použít číselné i slovní označení typu zprávy
/sbin/iptables -A INPUT -p icmp -icmp-type time-exceeded -j ACCEPT    povolili jsme příchozí ICMP
                                         pakety Time Exceeded (vypršení času při navazování spojení)
/sbin/iptables -A INPUT -p icmp -icmp-type destination-unreachable -j ACCEPT
                                         povolili jsme příchozí ICMP pakety hlásící nedostupnost cíle
```



Úkoly

1. Nastavte výchozí politiku pro odchozí provoz na akceptování.

2. Sestavte sadu pravidel, která z příchozího provozu zakážou všechno kromě provozu protokolu HTTP, POP a IMAP.
3. Na koncové stanici chceme povolit u příchozích paketů pouze ty, které jdou na TCP port 80 nebo TCP port 8080, všechny ostatní příchozí zakážeme. Jak budou vypadat pravidla pro příslušný chain?
4. Sestavte pravidla, která povolí ICMP pakety (jakékoliv, všechny typy ICMP zpráv) v příchozím, odchozím i průchozím provozu.



B.7.3 SPI

 Firewall v Linuxu s jádrem 2.4 a novějším podporuje také SPI (Stateful Packet Inspection). Aby bylo možné pracovat se stavy spojení, je třeba mít v jádře zaveden modul pro sledování stavu spojení (connection tracking). Dříve se používal modul `ip_conntrack` a další s ním související, nyní se spíše setkáme s bezpečnějším a funkčně širším `nf_conntrack` a moduly s ním souvisejícími. `nf_conntrack` plně podporuje IPv4 a IPv6 včetně jejich kombinování.



Příklad

Jak zjistíme, jestli je modul v jádře zaveden (výpis bývá dlouhý, přefiltrujeme si ho příkazem `grep`), výstup v SUSE Linuxu:

```
sarka@notebook: /home/sarka# /sbin/lsmmod | grep conntr
```

```
nf_conntrack_ipv6      20196      4
nf_conntrack_netbios_ns  2152      0
nf_conntrack_ipv4      10480     4
nf_conntrack           67400     5 nf_conntrack_ipv6,xt_NOTRACK,xt_state,
nf_conntrack_netbios_ns,nf_conntrack_ipv4
ip6v6                  242608    25 ip6t_REJECT,nf_conntrack_ipv6,ip6table_mangle
```

V prvním sloupci je název modulu, v druhém jeho velikost, následuje počet uživatelů tohoto modulu a případné závislosti (které zavedené moduly dotýčný modul vyžadují). Pokud výpis vypadal podobně jako ten výše, pak je vše OK.

Pokud potřebné moduly v jádře nejsou, tak je třeba je do jádra zavést. U každého potřebného modulu nejdřív ověříme, zda je dostupný:

```
/sbin/modeprobe -l *conntr*
```

(případně lze vypsát vše a pak to profiltrovat `grepem`). S dostupností obvykle nebývá problém. Zbývá modul dostat do jádra:

```
/sbin/modeprobe nf_conntrack
```

Pro další moduly je postup podobný.



Podrobný návod včetně ukávek využití modulů: <http://www.abclinuxu.cz/blog/pb/2007/2/nf-conntrack>



 Pokud máme modul v jádře, můžeme ve firewallu filtrovat pakety podle *stavu spojení* (můžeme je považovat za stavy paketů ve vztahu ke spojení, ke kterému patří). Rozlišují se tyto stavy:

- NEW: klient přes firewall navazuje nové spojení (tj. první paket spojení), může být také u jedno-směrných spojení

- ESTABLISHED: paket je součástí již navázaného spojení
- RELATED: paket sice navazuje nové spojení, ale zároveň je ve vztahu k některému existujícímu (typicky u FTP, kdy mohou být navázána dvě různá spojení na dvou portech)
- INVALID: paket nelze přiřadit k žádnému existujícímu ani navazovanému spojení

Pak v příkazech programu `iptables` můžeme používat filtrování podle stavů. Existují také virtuální stavy SNAT a DNAT pro pakety, kde byla přeložena zdrojová resp. cílová adresa.

Conntrack ve skutečnosti není samostatná tabulka, funkcionality se používá v existujících tabulkách a chainech (na obrázku B.3 na straně 283 vidíme, kde konkrétně v datovém toku se vyskytuje).



Příklad

Příklady pravidel pro příchozí spojení:

```
iptables -A INPUT -p tcp ! -syn -m state --state NEW
    -j LOG --log-prefix "Nenastaven SYN:" protokolují se veškeré pakety navazující spojení zvnějšku,
    které nemají nastaven (samotný) příznak SYN, toto pravidlo je obvykle následováno tímto:
iptables -A INPUT -p tcp ! -syn -m state --state NEW -j DROP   zajímáme se o pakety zapouzdřující
    TCP segment (protokol TCP), které nemají nastaven (samotný) příznak SYN (je tam negace)
    a zároveň stav spojení je nové (tj. navazuje se nové spojení a přitom není nastaven SYN bit),
    výslednou akcí je tiché zahození paketu
iptables -A INPUT -m state --state NEW -j DROP   zdánlivě totéž, co předtím, ale pozor – pokud zvenku
    přijde paket navazující nové spojení a nebude mít (chybně) nastaven příznak SYN, tak projde (tudíž
    pokud chceme blokovat příchozí navazování spojení, uvedeme obě pravidla)
iptables -A INPUT -m state --state RELATED, ESTABLISHED -j ACCEPT   necháme projít dovnitř pakety,
    které patří do existujícího spojení nebo s některým existujícím souvisejí
```



Pro práci s příchozími a odchozími spojeními a modulem `nf_conntrack` je možné použít také program `conntrack` z balíčku `conntrack-tools` (bývá běžně v repozitářích, můžeme doinstalovat).

Seznam aktivních spojení, která jsou povolena mechanismem `conntrack`, najdeme takto:

```
cat /proc/net/nf_conntrack
```



Příklad

Teď se podíváme na mezilehlé zařízení (třeba firewall), na kterém chceme nastavit port pro DMZ. Předpokládejme, že `eth0` je port do lokální sítě, `eth1` bude WAN rozhraní a `eth2` povede do DMZ:

```
iptables -A FORWARD -i eth0 -o eth2 -m state --state NEW, ESTABLISHED, RELATED -j ACCEPT   povo-
    líme pakety patřící do existujícího spojení, vztažené k existujícímu i navazující spojení (provoz
    z LAN do DMZ), tento směr prakticky není omezován (až na INVALID pakety), výstup příkazu
iptables -L -v bude:
```

```
Chain FORWARD (policy DROP)
target    prot opt in     out    source        destination
...
ACCEPT    all  -  eth0  eth2  anywhere      anywhere      state ESTABLISHED,RELATED
```

```
iptables -A FORWARD -i eth1 -o eth2 -m state --state NEW, ESTABLISHED,  
RELATED -j ACCEPT
```

 podobně pro provoz vedoucí z vnější sítě do DMZ, tady je potřeba mít povoleno navazování spojení, protože v DMZ jsou často webové a další servery, se kterými je potřeba navazovat spojení i zvnějšku

```
iptables -A FORWARD -i eth2 -o eth0 -m state --state ESTABLISHED,  
RELATED -j ACCEPT
```

 zařízení umístěná v DMZ jsou obecně považována za méně důvěryhodná, tj. navazovaná spojení nepustíme do LAN (pouze pakety patřící do navázaných nebo vztahených spojení)

```
iptables -A FORWARD -i eth2 -o eth1 -m state --state ESTABLISHED,  
RELATED -j ACCEPT
```

 podobně pro provoz z DMZ do vnější sítě, servery obvykle nemají co navazovat nová spojení

Pak samozřejmě potřebujeme zbývající pravidla pro řízení provozu na portech eth0 a eth1. U chainu FORWARD předpokládáme jako výchozí zásadu „zahazovat vše“ (DROP), tedy co jsme explicitně v pravidlech nepovolili nebo neurčili k zahazení s ICMP odezvou, bude zahazeno bez odezvy.



Poznámka:

Podpora conntrack je důležitá i pro funkčnost SNAT/DNAT/MASQUERADE, protože je třeba překládat podle jedné konkrétní dvojice adres v rámci jednoho spojení (v rámci dalšího spojení se může používat jiná soukromá adresa).



Úkoly

1. Zjistěte, zda máte v jádře načtené alespoň některé moduly pro SPI. Pokud ano, zjistěte, zda se sledování stavů spojení používá v aktuálních pravidlech firewallu (zjistíte z výpisu pravidel). Pokud máte dostatečná přístupová oprávnění, vypište obsah souboru, ve kterém jsou uchovávána momentálně navázaná spojení.
2. Sestavte pravidlo, které bude zakazovat příchozí spojení (na běžné pracovní stanici).
3. Sestavte pravidlo, které bude akceptovat pakety z již navázaných spojení.



B.7.4 NAT a maškaráda



Překlad adres se provádí vždy co nejbližší síťovému rozhraní, tj. těsně před odesláním paketu do sítě (všechny ostatní tabulky se procházejí předtím), resp. okamžitě po příchodu paketu ze sítě (a až potom se provádějí další tabulky).

Kromě targetů ACCEPT, DROP a jiných můžeme využívat také targety pro překlad adres:

- *SNAT* (Source NAT) – v odchozím provozu překládá statickou adresu odesílatele na jinou (statickou) adresu viditelnou ve vnější síti
- *DNAT* (Destination NAT) – v příchozím provozu překládá statickou adresu příjemce na jinou, obvykle soukromou statickou adresu viditelnou pouze ve vnitřní síti, používá se také pro přeposílání příchozích paketů na proces, který má tyto pakety dále zpracovat

- *MASQUERADE* (maškaráda) – je to obdoba SNAT, ale používáme tehdy, když překládáme *dynamické* adresy (v odchozím provozu překládá soukromé adresy z vnitřní sítě na jednu adresu viditelnou vně naší sítě, obvykle adresu toho zařízení, na kterém je toto pravidlo)

 Nejdřív bychom měli nastavit výchozí politiky:

`iptables -t nat -P OUTPUT -j ACCEPT` pokud něco odchází přes tabulku nat, necháme to jít (to je obvyklá výchozí politika pro tabulku nat)

`iptables -t nat -P PREROUTING -j ACCEPT` podobně pro chain PREROUTING, případné zahození necháme na jiných tabulkách, totéž bychom mohli provést pro POSTROUTING

 Nejdřív se podíváme na *překlad statické adresy*. Nemusí být zrovna veřejná, ale neměla by se měnit a měli bychom ji znát, protože tuto adresu napevno zadáváme do příkazu.

`iptables -t nat -A POSTROUTING -o eth2 -j SNAT -to 193.168.210.80` zajistíme SNAT na odchozím provozu odcházejícím přes eth2: zdrojová adresa v záhlaví odcházejícího paketu bude zaměněna za uvedenou (statickou viditelnou venku), do překladové tabulky se uloží údaj o této komunikaci (původní zdrojová plus cílová) pro zpětný překlad.

 Pokud v síti používáme DHCP, použijeme maškarádu. Maškaráda funguje podobně jako SNAT, ale využívá se výhradně u dynamických adres (ostatně, dynamickou adresu bychom stejně nemohli do SNAT pravidla naklepat), údaje pro překlad se uchovávají v jednom ze souborů v souborovém systému `/proc` (pro každé spojení se ukládá dvojice soukromá adresa v LAN plus vnější adresa, se kterou se komunikuje).

`iptables -t nat -A POSTROUTING -o eth2 -j MASQUERADE` zajistili jsme maškarádu (masquerade) na portu eth2, tj. vše, co odchází ven ze sítě na portu eth2, bude mít přeloženu zdrojovou adresu (obvykle soukromou, přidělenou DHCP serverem) na adresu routeru, automaticky je zajištěn záznam s uvedením portu pro zajištění zpětné konverzace (tj. PAT)

`iptables -t nat -A POSTROUTING -o ppp0 -j MASQUERADE` totéž, jen port je jinak pojmenován (to je velice častý název u routerů kombinovaných s DSL modemem)

Pokud má také náš router (firewall) IP adresu viditelnou vně naší sítě také dynamickou (od DHCP serveru našeho ISP), je třeba pozměnit hodnotu v jednom ze souborů virtuálního soub. systému `/proc`:

```
echo 7 > /proc/sys/net/ipv4/ip_dynaddr
```

To je pro případ, že „vnější“ adresa, která je při maškarádě použita pro překlad místních soukromých adres, „vypadne“, je potřeba požádat o novou, která je pravděpodobně jiná. Pokud je dotyčný parametr takto nastaven, pak se v takové situaci zamění stará dynamická adresa za nově získanou ve všech zaznamenaných právě probíhajících spojeních.

 Vrátime se ke statickým adresám. Mechanismus NAT se používá také pro *přeposílání příchozího provozu*:

`iptables -t nat -A PREROUTING -i eth2 -p tcp -dport 80 -j REDIRECT -to-port 3128` pokud přes rozhraní eth2 přijde paket z TCP portu č. 80 (velmi pravděpodobně pro protokol HTTP), přesměruj ho na port 3128; to se používá tehdy, když na portu 3128 naslouchá Squid konfigurovaný jako transparentní proxy, tj. dále prověřuje a zpracovává tento typ paketů

Všimněte si, že filtrování a NAT lze jednoduše kombinovat s překladem portů.

 Mechanismus NAT může sloužit pro *překlad jedné statické adresy na jinou statickou*, pak je potřeba zajistit ručně překlad v obou směrech a u každého směru uvést obě adresy. Výhodou tohoto řešení oproti jednoduchému SNAT a maškarádě je rychlejší zpracování (provoz není zdržován vyhledáváním dynamicky přidělených adres v dočasných souborech). Předpokládejme, že v DMZ máme server se soukromou (uvnitř viditelnou) adresou 10.201.51.2, chceme, aby byl vně sítě viditelný pod adresou 195.201.51.2. Je třeba zadat následující dvě pravidla, jedno pro provoz směrem od serveru ven, druhé pro opačný směr (zpětný překlad):

```
iptables -t nat -A PREROUTING -d 10.201.51.2 -j DNAT -to-destination 195.201.51.2
```

 nastavení DNAT, tedy překladu cílové adresy v paketu (provoz směrem do sítě), zamění soukromou adresu za uvedenou 195.201.51.2 v záhlaví přicházejícího paketu

```
iptables -t nat -A POSTROUTING -s 195.201.51.2 -j SNAT -to-source 10.201.51.2
```

 nastavení mechanismu SNAT, tedy překladu zdrojové adresy (provoz směrem ven)

Poznámka:

Na jednu věc bychom si měli dát pozor: mechanismus NAT funguje tak, že se pravidlo s přesměrováním adresy použije pouze na první paket spojení, a pouze první paket spojení prochází tabulkou *nat*. Další pakety patřící do téhož spojení (tj. pakety, které nejsou „první“), vůbec do tabulky *nat* nepřijdou, proto se jich tato pravidla netýkají.



Úkoly

1. Ve výpisu pravidel firewallu zjistěte, zda je nastaven překlad adres.
2. Sestavte pravidlo pro SNAT odchozího provozu na portu wan0, překládá se na statickou adresu 195.49.88.1.
3. Sestavte pravidlo pro maškarádu odchozího provozu na portu wan0.



B.7.5 Označování paketů

 Firewall NetFilter může do záhlaví paketů přidávat *značky* (marks), kterým rozumí jak jeho pravidla, tak i některé další nástroje (včetně příkazu *ip*). Přidávání značek je jednou z funkcí souvisejících s tabulkou *mangle* (protože jde o změnu neadresního údaje v záhlaví) nebo *nat*, značku pak můžeme využívat jinde (typicky v tabulce *filter*).

Rozlišujeme dva druhy značek:

- *NetFilter mark* se nastavuje targetem MARK, je viditelný i mimo firewall, touto značkou se označují pakety, v terminologii *iproute2* se také nazývá *fwmark* (firewall mark),
- *connection mark* se nastavuje targetem CONNMARK, slouží pouze pro vnitřní potřebu firewallu; touto značkou se označuje *spojení* (pozor, ne paket), značku lze přenést i do záhlaví paketů (pak jde o NetFilter značky), nebo naopak.

Značky NetFilter lze do záhlaví paketu uložit pomocí příkazu firewallu `iptables` targetem `MARK`, využívat je lze buď v tabulkách firewallu pro další filtrování či úpravy, a nebo mechanismem `iproute2` (příkazem `ip rule`).

 Ukážeme si několik příkazů s využitím značek:

```
iptables -t mangle -A PREROUTING -i eth2 -j MARK -set-mark 2
```

 pokud paket přišel z rozhraní `eth2`, dostane značku 2 (použijeme target `MARK`, protože označení je důsledkem, nikoliv podmínkou), platí jak pro pakety, které půjdou do chainu `INPUT`, tak i pro pakety mířící do chainu `FORWARD` v jiných tabulkách (na obrázku B.3 můžeme vidět, že označujeme prakticky hned na vstupu firewallu, pokud nepoužíváme tabulku `raw`)

```
iptables -t filter -A INPUT -m mark -mark 2 -p tcp -dport 80 -j ACCEPT
```

 chceme filtrovat pakety s podmínkou „pokud má paket značku 2“ (znamená: pokud má paket značku 2 a zároveň jde na TCP port 80, akceptuj), všimněte si odlišnosti práce s označením oproti předchozímu příkazu (zde je targetem `ACCEPT`)

```
iptables -t nat -A PREROUTING -p tcp -dport 80 -j CONNMARK -set-mark 2
```

 u příchozích spojení mířících na TCP port 80 (`http`) nastavíme značku spojení na 2 (protože je pravidlo v tabulce `nat`, provede se vždy jen pro první paket spojení, proto se často řadí právě do tabulky `nat`)

```
iptables -t mangle -A PREROUTING -j CONNMARK -save-mark
```

 značku spojení můžeme nastavit také takto: předpokládejme, že značka paketu byla nastavena již předtím (v předchozí tabulce), pak hodnota této značky paketu se uloží do značky spojení (tj. `connmark := fwmark`)

```
iptables -t mangle -A POSTROUTING -j CONNMARK -restore-mark
```

 opačný přenos informace, značka spojení se uloží do značky paketu (tj. `fwmark := connmark`)

 Pokud firewall označil paket značkou `fwmark`, lze tuto značku dále zpracovávat během směrování, například takto:

```
ip rule add fwmark 4 table tabVen prio 2021
```

 přidali jsme nové směrovací pravidlo: pakety označené značkou 4 jsou směrovány podle tabulky `tabVen` (na obrázku B.3 na straně 283 si všimněte, kdy ke směrování vzhledem k chainům firewallu dochází, a kdy tedy je třeba paket označit, aby se podle značky dalo směrovat)

Další informace:

Další informace o této problematice (včetně příkladů) najdeme na

- <http://www.linuxexpres.cz/praxe/ipp2p-kladivo-na-stahovace>
- <http://blog.khax.net/2009/12/01/multi-gateway-balancing-with-iptables/>
- <http://nerdboys.com/2006/05/05/conning-the-mark-multiwan-connections-using-iptables-mark-connmark-and-iproute2/>
- <http://home.regit.org/netfilter-en/netfilter-connmark/>
- http://ornellas.apanela.com/dokuwiki/pub:firewall_and_adv_routing



B.7.6 Skripty

Je vhodné si vytvořit jeden nebo více skriptů pro konfiguraci firewallu, a to z několika důvodů: předně potřebujeme, aby veškerá nastavení, která provedeme, byla platná neustále (i po případném restartu), abychom měli zálohu pro případ poruchy a aby bylo možné občas sadu pravidel měnit či někdy vypínat (mazat pravidla a pak podle potřeby obnovit).



Další informace:

Ukázky několika typických konfiguračních skriptů s pravidly `iptables` najdeme například na

- <http://www.faqs.org/docs/iptables/examplecode.html> (to je první skript, na další se dostaneme klepnutím na odkaz *Next* dole).
- <http://www.pmoghadam.com/blog/categories/Scripts/Round-robin%20load%20balancing%20NAT.txt>
- <http://tldp.org/HOWTO/IP-Masquerade-HOWTO/firewall-examples.html>
- <http://tldp.org/HOWTO/pdf/VPN-Masquerade-HOWTO.pdf>
- <http://jk.frozen-doe.net/hack/iptables/>



B.7.7 Uložení změn



Aby byla konfigurace nejen platná, ale také uložena pro načtení po vypnutí a zapnutí počítače, je nutné tuto konfiguraci uložit. Je velmi pravděpodobné, že distribuce, kterou máme, podporuje příkazy `iptables-save` a `iptables-restore`, které nám mohou velmi usnadnit práci s „hromadným“ startováním či ukončováním firewallu.¹⁵

`iptables-save` uloží obsah všech tabulek a chainů (tj. všechna pravidla) na standardní výstup, tj. používáme přesměrování, například `iptables-save > /etc/firerules`

`iptables-save -c -t filter` při použití přepínače `-c` se uloží také stav čítačů paketů a oktetů, přepínač `-t` zase omezí uložení pouze na zadanou tabulku (zde tabulku `filter`)

`iptables-restore` načte pravidla (tabulky, chainy) ze standardního vstupu (tj. opět musíme použít přesměrování nebo třeba rouru s příkazem `cat`, například `iptables-restore < /etc/firerules`)

`iptables-restore -c` načte také hodnoty čítačů

`iptables-restore -n` pokud už předtím byla nějaká pravidla ve firewallu aktivována, nebudou nově načtenými přepsána, nová pravidla se pouze přidají (bez uvedení parametru `-n` by po vykonání příkazu byla původně aktivní pravidla přepsána, odstraněna – flushed)

Pokud tyto příkazy nenajdeme, pak musíme použít postup typický pro danou distribuci. Například v distribucích založených na RedHat (včetně Mandrivy) je to příkaz

```
/sbin/service iptables save
```

V distribucích založených na Debianu v souboru `/etc/default/iptables` nastavíme řádek `enable_iptables_initd=true`

¹⁵Pozor – tím, že ukončíme program `iptables`, rozhodně nevypneme firewall, ten je totiž modulem jádra. Vpodstatě je firewall zapnutý po celou dobu, co je jeho modul zaveden v jádře, obdobou vypnutí by bylo vyprázdnění všech tabulek vnitřním příkazem `-F`.

čímž uložení povolíme, samotné uložení se provede příkazem
`/etc/init.d/iptables save_active`



Poznámka:

Pokud chceme firewall deaktivovat (pojem „vypnout“ není zrovna správný), vyprázdníme jeho chainy:
`iptables -F`
(případně zadáme tabulku a chain). Pokud ale něco takového opravdu chceme udělat, měli bychom si předem provést zálohu pravidel, třeba výše popsáním způsobem.



Není vždy nutné provádět konfiguraci a její uložení v textovém režimu. Existuje dokonce několik různých grafických rozhraní k programu `iptables`. V desktopových distribucích je najdeme celkem běžně, případně je možné si některý z nich stáhnout z repozitáře a nainstalovat.

V UNIXových systémech založených na BSD (například OpenBSD) se často setkáme s firewallem *PacketFilter* (starší byl *IPFilter*), obslužný program je `pfctl`.



Další informace:

Problematika firewallu v Linuxu je velmi rozsáhlá, sekce v těchto skriptech zdaleka nepokrývá všechny možnosti, které NetFilter nabízí. Další informace o `iptables`:

- <http://www.netfilter.org/documentation/index.html>
- <http://security.maruhn.com/iptables-tutorial/x9125.html>
- <http://lartc.org/lartc.html>
- <http://www.frozentux.net/iptables-tutorial/iptables-tutorial.html>
- <http://www.root.cz/serialy/vse-o-iptables/>
- <http://www.webstep.net/media/PDF/iptables.pdf>
- <http://www.faqs.org/docs/iptables/>
- http://ornellas.apanela.com/dokuwiki/pub:firewall_and_adv_routing
- http://www.asia-oss.net/aoss25_slides/TH_Sawangpong_Kernel_Development_and_Embedded_System.pdf



B.8 Vzdálený přístup



Pokud chceme přistupovat k některému uzlu sítě s Linuxem vzdáleně (z jiného počítače), lze použít několik nástrojů. Předně všechny takové uzly určitě podporují příkaz `telnet`. Ovšem to, že tento příkaz podporují, ještě neznamená, že ho můžeme použít. Především protokol *telnet* je na mnoha počítačích už během instalace zakázán, a to z bezpečnostních důvodů. Pokud přesto chceme tímto způsobem na daný počítač přistupovat, musíme *telnet* povolit. Postup je v různých distribucích odlišný, popis najdeme zde:

- Ubuntu: <http://www.cyberciti.biz/faq/ubuntu-linux-enable-telnet-service/>
- RedHat: http://aplawrence.com/Linux/enable_telnet.html

atd. (vždy bychom se měli raději podívat na stránky naší distribuce).



Nicméně používání protokolu *telnet* se moc nedoporučuje, místo něho bychom měli raději používat protokol *SSH* (příkaz `ssh` pro zajištění spojení a pak příkaz `scp` pro zabezpečené kopírování s využitím

protokolu SSH), který je taktéž dostupný na prakticky všech uzlech sítě a klienti jsou dostupní i pro Windows. SSH má výhody především v oblasti zabezpečení, šifruje se celá komunikace včetně zasílání přihlašovacích informací.

Nejdřív je potřeba nainstalovat na „vzdálený stroj“ SSH server, pak na stroj, ze kterého budeme na ten vzdálený přistupovat, nainstalujeme SSH klienta. Jako SSH server se v Linuxu běžně používá *OpenSSH Server*, který najdeme v repozitáři obvykle pod názvem `openssh-server` (záleží na konkrétní distribuci, prostě si prohlédneme obsah repozitáře v některém vhodném programu, třeba i v grafickém prostředí). Taký je možné, že už máme tento balíček nainstalovaný.

Co se týče klientů, na strojích s Linuxem už obvykle bývá nainstalován, ve Windows si můžeme stáhnout třeba program *puTTY* nebo *TTSSH*.



Další informace:

Postup konfigurace a používání SSH můžeme najít na těchto odkazech:

- <http://wiki.ubuntu.cz/SSH>
- http://www.xenocafe.com/tutorials/openssh_linux/redhat/openssh_linux_redhat.php
(oproti návodu lze OpenSSH instalovat i jednodušeji z repozitářů)
- <http://www.debianhelp.co.uk/ssh.htm>



Pokud chceme jen stáhnout soubor z některého FTP nebo WWW serveru, nepotřebujeme výše uvedené nástroje. Obvykle nám stačí webový prohlížeč nebo některý grafický klient, ale někdy se hodí i nástroj pro použití v textovém režimu. V tomto směru je k nezaplacení příkaz `wget`.

Tento příkaz tedy slouží ke stahování souborů z webu (podporuje protokoly FTP, HTTP, HTTPS včetně proxy), dokonce zvládá i rekurzivní stahování (například HTML stránky včetně vnitřní struktury, vytvoří se offline verze stránek) a dokáže navázat na dřívější přerušené stahování. Jednoduché příklady použití:

```
wget http://ftp.edUNIX.cz/pub/danix/Evolution2/dvd/DANIX-evol2-dvd-2007-11-04.iso
```

stahujeme DVD Linuxu Danix (pozor, vždy je třeba si zjistit adresu pro nejnovější verzi)

```
wget -c ftp://adresa/dokument.pripona
```

 naváže na předchozí přerušené stahování (příkaz je spouštěn v tomtéž pracovním adresáři jako při předchozím stahování)

Úkol

Zobrazte si manuálovou stránku příkazu `wget` a zjistěte, jak se dá spustit

- ve „spider“ módu, tedy dokument, který má zadán ke stáhnutí, nestáhne, ale jen zkontroluje, zda je na svém místě,
- bezpečně, aby se stahovalo s využitím některého šifrování (SSL nebo TLS),
- s rekurzivním stahováním.

Vyzkoušejte si stáhnutí některého dokumentu s příponou PDF z webu (případně ISO obsaz některé distribuce Linuxu, pokud máte dostatečně rychlé připojení bez limitů).



B.9 Další příkazy

V Linuxu máme k dispozici velmi mnoho dalších příkazů, které souvisejí s adresami na síti, například

- mtr** nástroj pro aktivní diagnostiku sítě (pozor, generuje provoz na síti, podle něho pak vyhodnocuje parametry sítě)
- lft** další diagnostický nástroj
- firewalk** bezpečnostní testovací nástroj, který zasíláním TCP nebo UDP paketů (lze určit parametrem) zjišťuje, které porty jsou na přeposílajících (forwarding) zařízeních v síti otevřené
- brctl** program pro administraci zařízení sloužícího jako ethernetový bridge, jako první parametr se zadává některý z příkazů pro konfiguraci mostu (lze například zobrazovat informace, pracovat s porty či s nastavením protokolu STP, apod.)
- tcpdump** analyzátor paketů (podobně jako Wireshark)
- ipcalc** jednoduchý program, který nám může pomoci orientovat se v IP adresách sítí/podsítí/uzlů („IP kalkulačka“), ukázka použití je v příkladu níže
- host** používá se podobně jako **nslookup** v neinteraktivním režimu, tedy IP adresu přeloží na název a naopak (oproti **nslookup** bere přednostně údaje ze souboru `/etc/hosts`, až když tam údaj nenajde, táže se DNS serveru)

Pro konfiguraci bezdrátové síťové karty lze použít nástroje **iwlist** a **iwconfig**.



Příklad

Ukážeme si práci se zajímavým příkazem **ipcalc**, který podle zadané IP adresy podsítě vypíše masku, inverzní masku, rozsah adres uzlů v zadané podsíti, broadcast adresu a počet možných uzlů v síti. Příkaz chceme spustit na stanici s nainstalovaným Ubuntu. Tato distribuce sice dotyčný nástroj ve výchozí instalaci neobsahuje, ale dokáže poradit, jakým způsobem máme nástroj nainstalovat:

```
sarka@notebook:~$ ipcalc 10.1.0.0/16
```

```
The program 'ipcalc' is currently not installed. You can install it by typing:
sudo apt-get install ipcalc
```

Takže program není nainstalován. Ale byli jsme upozorněni, že se nachází v repozitáři (jsme připojeni k Internetu, máme přístup do repozitářů Ubuntu), máme před očima příkaz, kterým program nainstalujeme. Protože známe administrátorské heslo, není to problém (všimněte si příkazu **sudo**, který v Ubuntu a některých dalších distribucích slouží k získání vyšších přístupových oprávnění).

```
sarka@notebook:~$ sudo apt-get install ipcalc
```

```
[sudo] password for sarka: *****
Reading package lists... Done
Building dependency tree
Reading state information... Done
The following packages were automatically installed and are no longer required:
  python-pylibacl python-brasero rdiff-backup python-pyxattr
Use 'apt-get autoremove' to remove them.
The following NEW packages will be installed:
  ipcalc
0 upgraded, 1 newly installed, 0 to remove and 378 not upgraded.
Need to get 26.6kB of archives.
After this operation, 131kB of additional disk space will be used.
```

```
Get:1 http://archive.ubuntu.com/ubuntu/ maverick/universe ipcalc all 0.41-2 [26.6kB]
Fetched 26.6kB in 0s (72.0kB/s)
Selecting previously deselected package ipcalc.
(Reading database ... 162465 files and directories currently installed.)
Unpacking ipcalc (from .../archives/ipcalc_0.41-2_all.deb) ...
Processing triggers for man-db ...
Setting up ipcalc (0.41-2) ...
```

Instalace je hotova, restartovat netřeba (nejsme přece ve Windows), tedy konečně můžeme spustit příkaz, jako parametr zadáme IP adresu (pod)sítě včetně délky prefixu:

```
sarka@notebook:~$ ipcalc 10.1.0.0/16
```

```
Address: 10.1.0.0          00001010.00000001. 00000000.00000000
Netmask: 255.255.0.0 = 16  11111111.11111111. 00000000.00000000
Wildcard: 0.0.255.255      00000000.00000000. 11111111.11111111
=>
Network: 10.1.0.0/16       00001010.00000001. 00000000.00000000
HostMin: 10.1.0.1         00001010.00000001. 00000000.00000001
HostMax: 10.1.255.254     00001010.00000001. 11111111.11111110
Broadcast: 10.1.255.255   00001010.00000001. 11111111.11111111
Hosts/Net: 65534          Class A, Private Internet
```



Existují různé nástroje pro testování zabezpečení sítě, které nejsou součástí běžných distribucí, ale rozhodně stojí za to si je stáhnout a používat. Například *SATAN*¹⁶ (System Administrators Tool for Analyzing Networks) prochází celou síť a provádí bezpečnostní testy. Je oblíbený jak u administrátorů, tak i u hackerů, i když každý z nich má samozřejmě jinou motivaci.



Další informace:

- Zajímavý tutoriál o sítích v Linuxu najdeme na adrese <http://linux-ip.net/html/part-concepts.html>.
- <http://www.linuxsoft.cz/wifi/>
- <http://www.root.cz/serialy/linux-jako-internetova-gateway/>
- <http://www.root.cz/clanky/luci-webove-rozhrani-pro-openwrt/>
- ČEČÁK, O. *Linux v příkazech – práce s Wi-Fi* [online]. Dostupné na http://www.linuxsoft.cz/article.php?id_article=1351



¹⁶ Ke stažení na <http://www.porcupine.org/satan/>.