

IDŹ DO

PRZYKŁADOWY ROZDZIAŁ



SPIS TREŚCI

KATALOG KSIĄŻEK

KATALOG ONLINE

ZAMÓW DRUKOWANY KATALOG

TWÓJ KOSZYK

DODAJ DO KOSZYKA

CENNIK I INFORMACJE

ZAMÓW INFORMACJE
O NOWOŚCIACH

ZAMÓW CENNIK

CZYTEL尼亚

FRAGMENTY KSIĄŻEK ONLINE

Sieci komputerowe – księga eksperta

Autor: Mark Sportack

Tłumaczenie: Zbigniew Gała

ISBN: 83-7197-076-5

Tytuł oryginału: [Networkig Essentials Unleashed](#)

Format: B5, stron: 616

oprawa twarda



Książka przybliży podstawowe założenia sieci komputerowych, które powinny być znane współczesnemu informatykowi. Krok po kroku wprowadzi Cię w problematykę sieci, pozwalając na poznanie ich architektury i zrozumienie zasad działania. Każdy rozdział zawiera wyczerpujące informacje na temat różnych mechanizmów sieciowych. Jest to również cenny podręcznik, w wielu przypadkach wystarczający do przygotowania się do egzaminów z zakresu sieci komputerowych. Podstawową zaletą książki jest fakt, że daje możliwość ukształtowania własnego punktu widzenia wobec ogromnej ilości coraz to nowszych rozwiązań w dziedzinie techniki komputerowej.



Spis treści

O autorach	14
Wprowadzenie.....	17
CZĘŚĆ I Podstawy sieci	
Rozdział 1. ABC sieci.....	21
Ewolucja sieci	21
Organizacje ustanawiające standardy	24
ANSI	24
IEEE	25
ISO	25
IEC	25
IAB.....	26
Model referencyjny OSI.....	26
Warstwa 1: warstwa fizyczna	28
Warstwa 2: warstwa łączy danych	29
Warstwa 3: warstwa sieci.....	30
Warstwa 4: warstwa transportu.....	30
Warstwa 5: warstwa sesji.....	30
Warstwa 6: warstwa prezentacji	31
Warstwa 7: warstwa aplikacji	31
Zastosowania modelu.....	31
Podstawy sieci	34
Sprzętowe elementy składowe	34
Programowe elementy składowe	37
Składanie elementów w sieć	38
Podsumowanie	42
Rozdział 2. Typy i topologie sieci LAN	43
Urządzenia przyłączane do sieci LAN	43
Typy serwerów	44
Typy sieci	48
Sieci równorzędne (każdy-z-każdym)	48
Sieci oparte na serwerach (klient-serwer).....	51
Sieci mieszane.....	54
Topologie sieci lokalnych	54
Topologia magistrali	55
Topologia pierścienia.....	56
Topologia gwiazdy.....	58
Topologia przełączana	59

Topologie złożone	61
Łącuchy	61
Hierarchie	62
Obszary funkcjonalne sieci LAN	65
Przylączenie stacji	65
Przylączenie serwera	65
Przylączenie do sieci WAN	66
Przylączenie do szkieletu	68
Podsumowanie	72
Rozdział 3. Warstwa fizyczna	75
Warstwa 1: warstwa fizyczna	75
Funkcje warstwy fizycznej	76
Znaczenie odległości	81
Tłumienie	81
Nośniki transmisji fizycznej	84
Kabel koncentryczny	85
Skrętka dwużyłowa	86
Kabel światłowodowy	91
Podsumowanie	95
Rozdział 4. Niezupelnie-fizyczna warstwa fizyczna	97
Spektrum elektromagnetyczne	97
Charakterystyki spektrum	99
Spektrum a szerokość pasma	100
Co to oznacza?	101
Bezprzewodowe sieci LAN	102
Bezprzewodowe łączenie stacji	102
Bezprzewodowe łączenie komputerów w sieci każdy-z-każdym	103
Bezprzewodowe łączenie koncentratorów	104
Bezprzewodowe mostkowanie	104
Technologie transmisji	105
Częstotliwość radiowa szerokiego spektrum	106
Jednopasmowa częstotliwość radiowa	110
Podczerwień	111
Laser	112
Standard IEEE 802.11	114
Dostęp do nośnika	114
Warstwy fizyczne	115
Podsumowanie	116
Rozdział 5. Warstwa łącza danych	117
Warstwa 2 modelu OSI	117
Ramki	118
Składniki typowej ramki	119
Ewolucja struktur ramek firmowych	120
Ramka sieci PARC Ethernet firmy Xerox	120
Ramka sieci DIX Ethernet	121
Projekt IEEE 802	123
Sterowanie łączem logicznym w standardzie IEEE 802.2	124
Protokół dostępu do podsieci (protokół SNAP) standardu IEEE 802.2	126
Ramka sieci Ethernet standardu IEEE 802.3	127
Sieci Token Ring standardu IEEE 802.5	131
Architektura FDDI	133

Zasady sterowania dostępem do nośnika	136
Dostęp do nośnika na zasadzie rywalizacji	136
Dostęp do nośnika na zasadzie priorytetu żądań	137
Dostęp do nośnika na zasadzie pierścienia	138
Wybór technologii LAN	139
Sieć Ethernet 802.3	140
Sieć Token Ring 802.5	140
Sieć FDDI	140
Sieć VG-AnyLAN 802.12	141
Podsumowanie	141
Rozdział 6. Mechanizmy dostępu do nośnika.....	143
Dostęp do nośnika	143
Dostęp do nośnika na zasadzie rywalizacji	144
Dostęp do nośnika na zasadzie pierścienia	149
Dostęp do nośnika na zasadzie priorytetu żądań	152
Dostęp do nośnika w komutowanych sieciach LAN	153
Podsumowanie	156
CZĘŚĆ II Tworzenie sieci lokalnych	
Rozdział 7. Ethernet.....	159
Różne rodzaje sieci Ethernet	159
Obsługiwany sprzęt	162
Funkcje warstwowe	164
Funkcje warstwy łącza danych	164
Funkcje warstwy fizycznej	166
Interfejsy międzyośnikowe warstwy fizycznej	168
10Base2	169
10Base5	169
10BaseT	169
10BaseFL	173
10BaseFOIRL	173
Mieszanie typów nośników	174
Ramka Ethernetu IEEE 802.3	175
Prognozowanie opóźnień	179
Szacowanie opóźnień propagacji	179
Prognozowanie opóźnień Ethernetu	180
Podsumowanie	181
Rozdział 8. Szybsze sieci Ethernet.....	183
Fast Ethernet	183
Nośniki Fast Ethernetu	185
100BaseTX	185
100BaseFX	186
100BaseT4	186
Schematy sygnalizacyjne	187
Maksymalna średnica sieci	188
Podsumowanie sieci Fast Ethernet	189
Gigabit Ethernet	189
Interfejsy fizyczne	189
Co jeszcze nowego?	193
Zbyt dobre, aby mogło być prawdziwe?	194
Podsumowanie	195

Rozdział 9. Token Ring.....	197
Przegląd.....	197
Standaryzacja sieci Token Ring.....	198
Struktura ramki Token Ring.....	199
Ramka Token.....	199
Ramka danych.....	201
Sekwencja wypełniania.....	204
Funkcjonowanie sieci Token Ring.....	204
Sprzęt.....	205
Topologia.....	207
Dynamiczna przynależność do pierścienia.....	207
Monitor aktywny.....	210
Co dalej z Token Ringiem?.....	212
Przełączanie a dedykowane sieci Token Ring.....	212
Zwiększanie szybkości transmisji.....	213
Będzie działać?.....	214
Podsumowanie.....	215
Zalety Token Ringu.....	215
Ograniczenia Token Ringu.....	216
Rozdział 10. FDDI.....	217
FDDI.....	217
Składniki funkcjonalne.....	218
Tworzenie sieci FDDI.....	221
Typy portów i metody przyłączania.....	221
Prawidłowe połączenia portów.....	223
Topologie i implementacje.....	224
Rozmiar sieci.....	230
Ramki FDDI.....	231
Ramka danych.....	231
Ramka danych LLC.....	233
Ramka danych LLC SNAP.....	234
Ramka Token.....	235
Ramki SMT.....	236
Mechanika sieci FDDI.....	236
Inicjalizacja stacji.....	236
Inicjalizacja pierścienia.....	238
Podsumowanie.....	238
Rozdział 11. ATM.....	239
Podstawy sieci ATM.....	240
Połączenia wirtualne.....	240
Typy połączeń.....	241
Szybkości przesyłania danych.....	242
Topologia.....	243
Interfejsy ATM.....	243
Model ATM.....	244
Warstwa fizyczna.....	245
Warstwa adaptacji ATM.....	247
Warstwa ATM.....	252
Komórka.....	253
Emulacja sieci LAN.....	256
Podsumowanie.....	259

Rozdział 12. Protokoły sieciowe	261
Stosy protokołów	261
Protokół Internetu, wersja 4 (Ipv4)	263
Analiza TCP/IP	264
Protokół Internetu, wersja 6 (IPv6).....	270
Struktury adresów <i>unicast</i> IPv6.....	272
Struktury zastępczych adresów <i>unicast</i> IPv6.....	273
Struktury adresów <i>anycast</i> IPv6	274
Struktury adresów <i>multicast</i> IPv6.....	274
Wnioski dotyczące IPv6	275
Wymiana IPX/SPX Novell	275
Analiza IPX/SPX	275
Warstwy łącza danych i dostępu do nośnika	279
Adresowanie IPX	280
Wnioski dotyczące IPX/SPX	280
Pakiet protokołów AppleTalk firmy Apple.....	281
Analiza AppleTalk	281
NetBEUI.....	286
Wnioski dotyczące NetBEUI.....	287
Podsumowanie	288

CZĘŚĆ III Tworzenie sieci rozległych

Rozdział 13. Sieci WAN	291
Funkcjonowanie technologii WAN.....	291
Korzystanie z urządzeń transmisji	292
Urządzenia komutowania obwodów	292
Urządzenia komutowania pakietów	295
Urządzenia komutowania komórek	297
Wybór sprzętu komunikacyjnego	298
Sprzęt własny klienta (CPE).....	299
Urządzenia pośredniczące (<i>Premises Edge Vehicles</i>).....	301
Adresowanie międzysieciowe	301
Zapewnianie adresowania unikatowego	301
Współdziałanie międzysieciowe z wykorzystaniem różnych protokołów	302
Korzystanie z protokołów trasowania	304
Trasowanie na podstawie wektora odległości.....	304
Trasowanie na podstawie stanu łącza	305
Trasowanie hybrydowe.....	305
Trasowanie statyczne.....	306
Wybór protokołu	306
Topologie WAN.....	307
Topologia każdy-z-każdym	307
Topologia pierścienia.....	309
Topologia gwiazdy.....	310
Topologia oczek pełnych	312
Topologia oczek częściowych	313
Topologia dwuwarstwowa	313
Topologia trójwarstwowa	315
Topologie hybrydowe	316
Projektowanie własnych sieci WAN.....	318
Kryteria oceny wydajności sieci WAN.....	318
Koszt sieci WAN	322
Podsumowanie	323

Rozdział 14. Linie dzierżawione	325
Przegląd linii dzierżawionych	325
Techniki multipleksowania	326
Cienie i blaski linii dzierżawionych	327
Topologia linii dzierżawionych	330
Standardy sygnałów cyfrowych	332
Hierarchia ANSI sygnału cyfrowego	333
Systemy nośników SONET	335
System T-Carrier	336
Usługi T-Carrier	337
Kodowanie sygnału	338
Formaty ramek	339
Podsumowanie	341
Rozdział 15. Urządzenia transmisji w sieciach z komutacją obwodów	343
Sieci Switched 56	343
Najczęstsze zastosowania sieci Switched 56	344
Technologie Switched 56	344
Sieci Frame Relay	345
Frame Relay a linie dzierżawione	346
Rozszerzone Frame Relay	348
Stałe a komutowane kanały wirtualne	349
Format podstawowej ramki Frame Relay	350
Projektowanie sieci Frame Relay	351
UNI a NNI	351
Przekraczanie szybkości przesyłania informacji	352
Sterowanie przepływem w sieci Frame Relay	353
Przesyłanie głosu za pomocą Frame Relay	354
Sieci prywatne, publiczne i hybrydowe (mieszane)	355
Współdziałanie międzysieciowe przy zastosowaniu ATM	359
ATM	359
Historia ATM	360
ATM – sedno sprawy	362
Identyfikatory ścieżki wirtualnej (VPI), a identyfikatory kanału wirtualnego (VCI)	365
Połączenia ATM	366
Jakość usług	366
Sygnalizowanie	367
Zamawianie obwodów ATM	367
Współdziałanie przy użyciu emulacji LAN	368
Migrowanie do sieci ATM	368
Podsumowanie	369
Rozdział 16. Urządzenia transmisji w sieciach z komutacją pakietów	371
Sieci X.25	371
Historia X.25	372
Zalety i wady sieci X.25	373
Najczęstsze zastosowania	373
Porównanie z modelem OSI	373
Różne typy sieci	378
Specyfikacje X.25 (RFC 1356)	378
Migrowanie z sieci X.25	379
Podsumowanie	380

Rozdział 17. Modemy i technologie Dial-Up.....	381
Sposób działania modemu.....	381
Bity i body	383
Typy modulacji modemów	385
Asynchronicznie i synchronicznie	387
Standardowe interfejsy modemów	388
Standardy ITU-T (CCITT) modemów	391
Modemy a Microsoft Networking.....	393
Podsumowanie	395
Rozdział 18. Usługi dostępu zdalnego (RAS).....	397
Historia korzystania z sieci o dostępie zdalnym	397
Lata siedemdziesiąte	398
Lata osiemdziesiąte.....	399
Szaleństwo lat dziewięćdziesiątych	399
Ustawianie połączeń zdalnych	400
Ewolucja standardów protokołów	401
Zestaw poleceń AT	401
Protokoły połączeń zdalnych	403
Ustawianie sesji.....	403
Protokoły dostępu sieci TCP/IP	403
Usługi transportu zdalnego	406
W jaki sposób obecnie łączą się użytkownicy usług dostępu zdalnego	406
Możliwości dostępu zdalnego Windows NT.....	414
Korzystanie z usług dostępu zdalnego jako bramy/routera sieci LAN.....	414
Korzystanie z usług dostępu zdalnego w celu umożliwienia dostępu do Internetu przy użyciu modemów	417
Możliwości dostępu zdalnego Novell NetWare Connect	419
Możliwości dostępu zdalnego systemów Banyan.....	419
Bezpieczeństwo dostępu zdalnego	420
Hasła	421
Dialery	422
Systemy „callback” połączeń zwrotnych.....	422
Podsumowanie	423
Rozdział 19. Sieci Intranet oraz Ekstranet	425
Sieci Intranet	426
Co takiego piszczy w sieci WWW?.....	426
A co śwista w sieci Intranet?	428
Sieci Ekstranet.....	430
Problemy z protokołami otwartymi	430
Problemy z protokołami bezpołączeniowymi.....	431
Problemy z protokołami otwartymi oraz bezpieczeństwem sieci ekstranetowych.....	434
Zasady ochrony sieci Ekstranet	435
Czy aby nie tracę czasu?	437
Wirtualne sieci prywatne.....	438
Wirtualne sieci prywatne dostarczane przez firmy telekomunikacyjne	439
Tunelowanie.....	440
Podsumowanie	441

Część IV Korzystanie z sieci	443
Rozdział 20. Sieciowe systemy operacyjne	445
Historia sieciowych systemów operacyjnych	445
Firma Novell dominuje rynek	446
Wchodzą nowi gracze	446
Uwaga – Microsoft przejmuje pałeczkę	447
Sytuacja obecna	447
Tradycyjne usługi sieciowych systemów operacyjnych	448
Systemy sieciowe Banyan	450
Usługi i aplikacje systemu VINES	450
Standardy obsługiwane przez VINES	452
Mocne i słabe strony VINES	452
Novell NetWare	453
Właściwości NetWare	454
Standardy obsługiwane przez NetWare	454
Mocne i słabe strony NetWare	458
Microsoft Windows NT	458
Właściwości Windows NT	460
Standardy obsługiwane przez Windows NT	462
Bezpieczeństwo Windows NT	462
Mocne i słabe strony Windows NT	463
Podsumowanie	463
Rozdział 21. Administrowanie siecią	465
Administrowanie siecią – coż to oznacza?	465
Zarządzanie kontami sieciowymi	466
Konta użytkowników	466
Konta grup	470
Logowanie wielokrotne	473
Zarządzanie zasobami	473
Zasoby sprzętowe	474
Wydzielone obszary dysku	474
Pliki i katalogi	474
Instalowanie/aktualizowanie oprogramowania	475
Drukowanie w sieci	476
Narzędzia zarządzania	477
Narzędzia zarządzania Microsoftu	477
„Zero administracji”	480
Konsola Zarządzania Microsoftu	480
Podsumowanie	481
Rozdział 22. Zarządzanie siecią	483
Wydajność sieci	483
Warstwa fizyczna	483
Natężenie ruchu	485
Problemy rozróżniania adresów	488
Współdziałanie międzysieciowe	488
Narzędzia i techniki	488
Ping	489
Traceroute	491
Monitor wydajności Windows NT	492
Analizatory sieci	492

Rozwiązanie problemów sprzętowych.....	493
Podsumowanie	496
Rozdział 23. Bezpieczeństwo danych	497
Planowanie w celu zwiększenia bezpieczeństwa sieci oraz danych	497
Poziomy bezpieczeństwa	498
Założenia bezpieczeństwa.....	499
Grupy robocze, domeny i zaufanie	501
Modele czterech domen	504
Konfigurowanie bezpieczeństwa w Windows 95	505
Udostępnianie chronione hasłem	506
Konfigurowanie bezpieczeństwa w Windows NT.....	508
Zgodność z klasyfikacją C2	512
Inspekcja	512
Bezdyskowe stacje robocze	513
Szyfrowanie	514
Ochrona antywirusowa	515
Podsumowanie	515
Rozdział 24. Integralność danych	517
Ochrona systemu operacyjnego	519
Procedury instalacji.....	519
Techniki konserwacji.....	522
Ochrona sprzętu	530
Systemy „UPS” zasilania nieprzerwanego	530
Czynniki środowiskowe.....	534
Bezpieczeństwo fizyczne	535
Nadmiarowość sprzętu.....	535
Ochrona danych	536
Tworzenie kopii zapasowych danych	536
Zapassowe przestrzenie składowania na dysku.....	542
Wdrażanie planu zapewnienia integralności danych	544
Krótki list na temat integralności danych.....	544
Podsumowanie	546
Rozdział 25. Zapobieganie problemom	547
Proaktywne operacje kontroli sieci	547
Zastosowania proaktywnych operacji kontroli sieci.....	550
Testowanie, baselining oraz monitorowanie sieci	553
Doskonalenie istniejących operacji proaktywnej kontroli sieci.....	554
Proaktywne operacje obsługi katastrof sieci	555
Zastosowanie proaktywnych operacji obsługi katastrof sieci.....	556
Testowanie czynności i strategii usuwania skutków katastrof	559
Doskonalenie istniejących operacji obsługi katastrof sieci	559
Podsumowanie	560
Rozdział 26. Rozwiązanie problemów	561
Logiczne wyodrębnianie błędów	561
Określanie priorytetu	562
Kompletowanie stosownej informacji	563
Określanie prawdopodobnych przyczyn problemu	565
Sprawdzanie rozwiązań	565
Badanie i ocena wyników	566
Wyniki oraz przebieg przenoszenia	567

Częste problemy sieciowe	567
Nośnik fizyczny	567
Karta sieciowa.....	568
Parametry konfiguracji karty sieciowej	569
Niezgodność protokołów sieciowych	570
Przeciążenie sieci	571
Sztormy transmisji	572
Problemy zasilania	572
Problemy serwera	573
Narzędzia gromadzenia informacji	574
Cyfrowe mierniki napięcia.....	574
Reflektometry czasowe.....	574
Oscyloskopy.....	575
Zaawansowane urządzenia kontroli kabli	575
Analizatory protokołów	575
Monitory sieci	576
Monitory wydajności	576
Przydatne zasoby.....	577
Serwis techniczny producenta.....	577
Internetowe grupy dyskusyjne oraz listy adresowe	577
Miejsca pobierania danych z sieci	577
Magazyny i czasopisma techniczne	578
Listy zgodnych z Windows NT urządzeń i programów	578
Sieć informacji technicznej Microsoft	578
Sieciowa baza wiedzy Microsoftu	578
Zestaw Resource Kit serwera Windows NT	579
Podsumowanie	579

Dodatki

Słowniczek	583
Skorowidz	605

Rozdział 12.

Protokoły sieciowe

Mark A. Sportack

Termin „protokoły sieciowe” odnosi się przede wszystkim do protokołów warstwy 3 modelu OSI. Protokoły zapewniają adresowanie, dzięki któremu dane mogą być dostarczane na nieokreślone odległości, poza domenę sieci lokalnej nadawcy. Przeważnie protokoły warstwy 3 wykorzystują do transportu danych strukturę znaną jako pakiet.

Choć protokoły warstwy 3 dostarczają mechanizmów niezbędnych do wysyłania pakietów, nie są na tyle wyszukane, aby mieć pewność, że pakiety zostały rzeczywiście odebrane i to we właściwym porządku. Zadania te pozostawiono protokołom transportowym warstwy 4. Protokoły te przyjmują dane z wyższych warstw i osadzają je w segmentach, które przekazują warstwie 3.

W tym rozdziale opisane są funkcje i wzajemne oddziaływania między stosami protokołów warstwy 3 i 4, a następnie badane są zawiłości niektórych najpopularniejszych protokołów sieciowych.

Stosy protokołów

Stos protokołów to komplet powiązanych protokołów komunikacyjnych, oferujących użytkownikowi mechanizmy i usługi potrzebne do komunikacji z innymi maszynami włączonymi do sieci. Z perspektywy użytkownika stos protokołów jest tym, co czyni sieć zdolną do użycia.

W poprzednich rozdziałach omówiono pierwszą i drugą warstwę stosu protokołów (tj. warstwę fizyczną i warstwę łącza danych). Są one mocno zintegrowane i powiązane ze sobą. Warstwę fizyczną narzuca wybrana architektura warstwy łącza danych, jak Ethernet, Token Ring itd.

W obecnej epoce sieci i systemów otwartych wybór określonej architektury LAN nie ogranicza możliwości wyboru protokołów wyższego poziomu. Stos protokołów powinien oferować mechanizmy sprzęgające z istniejącymi, znormalizowanymi środkami dostępu do sieci dla protokołów warstwy łącza danych.

Podobnie jak to było w przypadku warstw pierwszej i drugiej, warstwa 3 modelu referencyjnego OSI jest ściśle powiązana z warstwą 4. Warstwa 3 to warstwa sieci, zaś warstwa 4 jest warstwą transportu. Są one przedstawione na rysunku 12.1. Razem zapewniają one mechanizmy umożliwiające przesyłanie informacji między urządzeniami nadawcy i odbiorcy, z wykorzystaniem sieci komunikacyjnej sięgającej poza domenę warstwy 2. Zapewniają też inne funkcje, takie jak zmienianie porządku pakietów otrzymanych w niewłaściwej kolejności lub ponowną transmisję pakietów, które nie dotarły do odbiorcy lub dotarły uszkodzone.

Protokoły te, aby faktycznie przesłać dane, muszą wykorzystywać architekturę LAN warstwy 1 i 2, muszą też mieć pewne środki sprzęgające je z tymi warstwami. Używają do tego techniki opakowywania. Protokół warstwy 4 umieszcza w segmentach dane otrzymane od protokołów wyższej warstwy. Segmenty te są przekazywane do odpowiedniego protokołu warstwy 3. Protokół warstwy 3 bezzwłocznie opakowuje segment strukturą pakietu z adresami nadawcy i odbiorcy, po czym przekazuje pakiet protokołowi warstwy 2. Warstwa 2 umieszcza ten pakiet danych warstwy 3 w ramce, opatrując go przy tym adresowaniem, przeznaczonym dla urządzeń warstwy 3, takich jak routery i przełączniki IP. Umieszczenie pakietu warstwy 3 (IP) w ramce warstwy 2 (Ethernet) przedstawia rysunek 12.2. Strukturę pakietu zalicza się do części pola danych ramki, choć w rzeczywistości pakiet to nie dane, tylko struktura innej warstwy.

Rysunek 12.1.

Warstwy sieci i transportu w modelu referencyjnym OSI.

	Nazwa warstwy modelu referencyjnego OSI	Numer warstwy OSI
	Aplikacji	7
	Prezentacji	6
	Sesji	5
Stos protokołów sieci	Transportu	4
	Sieci	3
	Łącza danych	2
	Fizyczna	1

Rysunek 12.2.

Umieszczenie pakietu IP w ramce Ethernet.

7-okteta Preambula	1-oktetowy Ogranicznik początku ramki	6-oktetowy Adres odbiorcy	6-oktetowy Adres nadawcy	2-oktetowe pole Długość	Min 20-oktetowy nagłówek IP	Min 20-oktetowy nagłówek TCP	Pole Dane o zmiennej długości od 48 do 1482 oktetów	4-oktetowa Sekwencja kontrolna ramki
--------------------	---------------------------------------	---------------------------	--------------------------	-------------------------	-----------------------------	------------------------------	---	--------------------------------------

Ramki służą do transportowania i adresowania danych dla międzysieciowego urządzenia warstwy 3, znajdującego się na krawędzi warstwy 2 (dziedziny ramek). Urządzenie – zwykle jest to router – przyjmuje obramowany pakiet, usuwa ramkę i czyta informację adresową przeznaczoną dla warstwy 3. Informacja ta służy do ustalenia następnego „skoku” na drodze do miejsca przeznaczenia pakietu. Pakiet jest następnie kierowany do kolejnego punktu. Ostatni router przed miejscem przeznaczenia pakietu musi ponownie umieścić pakiet w strukturze ramki warstwy 2, zgodnej z architekturą sieci LAN w miejscu przeznaczenia.

Warstwa 3 zapewnia tylko międzysieciowy transport danych. Warstwa 4 (transportu) wzbogaca mechanizmy sieciowe warstwy 3 o gwarancje niezawodności i integralności końcowej (czyli na całej długości połączenia). Warstwa transportu może zagwarantować wolne od błędów dostarczenie pakietów i odpowiednie ich uszeregowanie, jak też zapewnić odpowiednią jakość usług. Przykładem mechanizmu warstwy 4 jest protokół TCP (ang. *Transmission Control Protocol*). TCP niemal zawsze występuje razem ze swoim odpowiednikiem z warstwy 3, protokołem internetowym IP, jako „TCP/IP”.

Wykorzystywanie przez aplikacje warstw 3 i 4 do przesyłania danych do innych komputerów/aplikacji sugeruje, że komputery nadawcy i odbiorcy nie są przyłączone do tej samej sieci lokalnej, niezależnie od tego, jaka odległość je dzieli. Dwie różne sieci muszą być ze sobą połączone, by obsłużyć żądaną transmisję. Dlatego mechanizmy komunikacyjne warstwy 2 są niewystarczające i muszą być rozszerzone o adresowanie warstwy 3.

Choć warstwy 3 i 4 istnieją właśnie w tym celu (tzn. w celu łączenia ze sobą różne sieci), to aplikacje mogą przysyłać do siebie dane, wykorzystując protokoły tych warstw, nawet jeśli są przyłączone do tej samej sieci i podsieci LAN. Na przykład, jeśli komputery nadawcy i odbiorcy są przyłączone do tej samej sieci lokalnej, mogą się ze sobą komunikować, wykorzystując tylko ramki i protokoły warstwy 2. Niektóre aplikacje mogą jednak wymagać wspomagania swojej komunikacji pewnymi właściwościami protokołów wyższej warstwy.

Istnieją dwa rodzaje protokołów sieciowych działających w warstwie 3: protokoły trasowane i protokoły trasujące. Protokoły trasowane to te, które umieszczają dane oraz informacje użytkownika w pakietach i są odpowiedzialne za przesłanie pakietów do odbiorcy. Protokoły trasujące stosowane są pomiędzy routerami i określają dostępne trasy, komunikują o nich i przeprowadzają nimi pakiety protokołów trasowanych. Protokoły trasujące są dokładniej omówione w rozdziale 13 pt. „Sieci WAN”. W niniejszym rozdziale koncentrujemy się na najpopularniejszych protokołach trasowanych.

Protokół Internetu, wersja 4 (Ipv4)

Protokół Internetu (IP)¹ został opracowany około 20 lat temu dla Departamentu Obrony USA (ang. *Department of Defense*). Departament Obrony szukał sposobu połączenia różnych rodzajów posiadanych komputerów i sieci je obsługujących w jedną, wspólną sieć. Osiągnięto to za pomocą warstwowego protokołu, który odizolował aplikacje od sprzętu sieciowego. Protokół ten używa modelu nieco różniącego się od modelu referencyjnego OSI. Jest on znany jako model TCP/IP.

¹ Należy w tym miejscu przypomnieć, iż rozpowszechnione obecnie tłumaczenie angielskiego terminu *Internet Protocol* (którego akronimem jest właśnie IP) jako „protokół internetowy” jest konsekwencją rozpowszechnienia się zastosowań Internetu – wierniejszym wydaje się bowiem używany pierwotnie odpowiednik „protokół międzysieciowy”. Sam mianowicie wyraz „Internet”, nim stał się (bodaj najpopularniejszą) nazwą własną, oznaczał po prostu współpracę pomiędzy sieciami – ujętą, trzeba przyznać, w charakterystyczną dla anglosasów lakoniczną formę słowną (przyp. red.)

W odróżnieniu od modelu OSI, model TCP/IP bardziej koncentruje się na zapewnianiu przyłączalności niż na sztywnym przywiązaniu do warstw funkcjonalnych. Uznając znaczenie hierarchicznego uporządkowania funkcji, jednocześnie zostawia projektantom protokołu sporą elastyczność odnośnie implementacji. W konsekwencji model OSI znacznie lepiej wyjaśnia mechanikę komunikacji między komputerami, ale TCP/IP stał się protokołem współdziałania międzysieciowego preferowanym przez rynek.

Elastyczność modelu referencyjnego TCP/IP, w porównaniu z modelem OSI, przedstawia rysunek 12.3.

Rysunek 12.3.

Porównanie modeli referencyjnych OSI i TCP/IP.

Nazwa warstwy modelu referencyjnego OSI	Numer warstwy OSI	Nazwa równoważnej warstwy TCP/IP
Aplikacji	7	Procesu/aplikacji
Prezentacji	6	
Sesji	5	
Transportu	4	Hosta z hostem
Sieci	3	Internetu
Łączy danych	2	Dostępu do sieci
Fizyczna	1	

Model referencyjny TCP/IP, opracowany długo po tym, jak powstały protokoły, które opisuje, oferuje dużo większą elastyczność niż model OSI, gdyż wyraża raczej hierarchiczne uporządkowanie funkcji, a nie ich ścisłą strukturę warstwową.

Analiza TCP/IP

Stos protokołów TCP/IP zawiera cztery warstwy funkcjonalne: dostępu do sieci, Internetu, warstwę host-z-hostem i warstwę procesu/aplikacji. Te cztery warstwy luźno nawiązują do siedmiu warstw modelu referencyjnego OSI, nie tracąc na funkcjonalności.

Warstwa procesu/aplikacji

Warstwa aplikacji dostarcza protokoły zdalnego dostępu i współdzielenia zasobów. Znane aplikacje, jak Telnet, FTP, SMTP, HTTP i wiele innych, znajdują się i działają w tej warstwie i są uzależnione od funkcjonalności niższych warstw.

Warstwa „host-z-hostem”

Warstwa „host-z-hostem” protokołu IP luźno nawiązuje do warstw sesji i transportu modelu OSI. Obejmuje dwa protokoły: protokół sterowania transmisją (ang. *TCP – Transmission Control Protocol*) i protokół datagramów użytkownika (ang. *UDP – User Datagram Protocol*). Obecnie, w celu dostosowania do coraz bardziej zorientowanego

na transakcje charakteru Internetu, definiowany jest trzeci protokół. Protokół ten nosi próbną nazwę protokołu sterowania transmisją i transakcją (ang. *T/TCP – Transaction Transmission Control Protocol*).

Protokół TCP zapewnia połączeniową transmisję danych pomiędzy dwoma lub więcej hostami, może obsługiwać wiele strumieni danych, umożliwia sterowanie strumieniem danych, kontrolę błędów, a nawet ponowne porządkowanie pakietów, otrzymanych w niewłaściwej kolejności.

Nagłówek protokołu TCP ma długość co najmniej 20 oktetów i zawiera następujące pola:

- ♦ Port Źródłowy TCP: 16-bitowe pole portu źródłowego przechowuje numer portu, który inicjuje sesję komunikacyjną. Port źródłowy i adres źródłowy IP funkcjonują jako adres zwrotny pakietu.
- ♦ Port Docelowy TCP: 16-bitowe pole portu docelowego jest adresem portu, dla którego przeznaczona jest transmisja. Port ten zawiera adres interfejsu aplikacji w komputerze odbiorcy, do której przesyłany jest pakiet danych.
- ♦ Numer Sekwencji TCP: 32-bitowy numer sekwencji jest wykorzystywany przez komputer odbierający do zrekonstruowania rozproszonych, rozbitych, podzielonych danych i przywrócenia im pierwotnej postaci. W sieci dynamicznie trasowanej może się zdarzyć, że niektóre pakiety pójdą innymi trasami i dotrą w niewłaściwej kolejności. Pole numeru sekwencji kompensuje tę niekonsekwencję w dostarczaniu.
- ♦ Numer Potwierdzenia TCP: TCP używa 32-bitowego potwierdzenia (ACK) pierwszego oktetu danych zawartego w następnym oczekiwanym segmencie. TCP może obliczyć ten numer, zwiększając numer ostatniego otrzymanego oktetu o liczbę oktetów w każdym segmencie TCP. Numer używany do identyfikowania każdego ACK jest numerem sekwencji potwierdzanego pakietu.
- ♦ Wyrównanie Danych: 4-bitowe pole przechowujące rozmiar nagłówka TCP, którego miarą jest 32-bitowa struktura danych, znana jako „słowo”.
- ♦ Zarezerwowane: 6-bitowe pole, zawsze ustawione na zero. Pole to jest zarezerwowane dla jeszcze nie wyspecyfikowanego, przyszłego zastosowania.
- ♦ Flagi: 6-bitowe pole flagi zawiera sześć 1-bitowych flag, które umożliwiają realizację funkcji sterowania, takich jak pole pilność, potwierdzenie pola znaczącego, pchanie, zerowanie połączenia, synchronizacja numerów sekwencyjnych i zakończenie wysyłania danych.
- ♦ Rozmiar Okna: 16-bitowe pole używane przez komputer docelowy w celu poinformowania komputera źródłowego o tym, ile danych jest gotów przyjąć w jednym segmencie TCP.
- ♦ Suma Kontrolna (16-bitów): Nagłówek TCP zawiera również pole kontroli błędów, znane jako „suma kontrolna”. Komputer źródłowy oblicza wartość tego pola na podstawie zawartości segmentu. Komputer docelowy przeprowadza identyczne obliczenie. Jeśli zawartość pozostała nienaruszona, wynik obydwu obliczeń będzie identyczny, co świadczy o prawidłowości danych.
- ♦ Wypełnienie: do tego pola dodawane są zera, tak aby długość nagłówka TCP była zawsze wielokrotnością 32 bitów.

Protokół datagramów użytkownika (UDP) jest innym protokołem IP warstwy host-z-hostem (odpowiadającej warstwie transportu modelu OSI). Protokół UDP zapewnia proste i mające niewielki narzut transmisje danych, tzw. „datagramy”. Prostota datagramów czyni UDP protokołem nieodpowiednim dla niektórych aplikacji, za to doskonałym dla aplikacji bardziej wyszukanych, które same mogą zapewnić funkcjonalność połączeniową.

Protokół UDP może być wykorzystywany przy wymianianiu takich danych, jak nadane nazwy NetBIOS, komunikaty systemowe itd., gdyż wymiany te nie wymagają sterowania strumieniem danych, potwierdzeń, ponownego uporządkowywania ani innych funkcji dostarczanych przez protokół TCP.

Nagłówek protokołu UDP ma następującą strukturę:

- ♦ Numer portu źródłowego UDP: Port źródłowy jest numerem połączenia w komputerze źródłowym. Port źródłowy i adres źródłowy IP funkcjonują jako adres zwrotny pakietu.
- ♦ Numer portu docelowego UDP: Port docelowy jest numerem połączenia w komputerze docelowym. Port docelowy UDP jest wykorzystywany do przekazywania pakietu odpowiedniej aplikacji, po tym jak pakiet dotrze do komputera docelowego.
- ♦ Suma kontrolna UDP: Suma kontrolna jest polem kontroli błędów, którego wartość jest obliczana na podstawie zawartości segmentu. Komputer docelowy wykonuje taką samą funkcję matematyczną jak komputer źródłowy. Niezgodność dwóch obliczonych wartości wskazuje na wystąpienie błędu podczas transmisji pakietu.
- ♦ Długość komunikatu UDP: Pole długości komunikatu informuje komputer docelowy o jego rozmiarze. Daje to komputerowi docelowemu kolejny mechanizm, wykorzystywany do sprawdzania poprawności wiadomości.

Główną różnicą funkcjonalną pomiędzy TCP a UDP jest niezawodność. Protokół TCP charakteryzuje się wysoką niezawodnością, natomiast UDP jest prostym mechanizmem dostarczania datagramów. Ta fundamentalna różnica skutkuje ogromnie zróżnicowaniem zastosowań tych dwóch protokołów warstwy host-z-hostem.

Warstwa Internetu

Warstwa Internetu protokołu IPv4 obejmuje wszystkie protokoły i procedury potrzebne do przesyłania danych pomiędzy hostami w wielu sieciach. Pakiety przenoszące dane muszą być trasowalne. Odpowiada za to protokół Internetu (IP).

Nagłówek protokołu IP ma następujący rozmiar i strukturę:

- ♦ Wersja: Pierwsze cztery bity nagłówka IP identyfikują wersję operacyjną protokołu IP, np. wersję 4.
- ♦ Długość Nagłówka Internetu: Następne cztery bity nagłówka zawierają jego długość, wyrażoną w wielokrotnościach liczby 32.

- ♦ Rodzaj Usługi: Następne 8 bitów to 1-bitowe flagi, które mogą być używane do określania parametrów pierwszeństwa, opóźnienia, przepustowości i niezawodności tego pakietu danych.
- ♦ Długość Całkowita: 16-bitowe pole przechowujące całkowitą długość datagramu IP, mierzoną w oktetach. Prawidłowe wartości mogą mieścić się w przedziale od 576 do 65536 oktetów.
- ♦ Identyfikator: Każdemu pakietowi IP nadaje się unikatowy, 16-bitowy identyfikator.
- ♦ Flagi: Następne pole zawiera trzy 1-bitowe flagi, wskazujące, czy dozwolona jest fragmentacja pakietów i czy jest ona stosowana.
- ♦ Przesunięcie Fragmentu: 8-bitowe pole mierzące przesunięcie fragmentowanej zawartości względem początku całego datagramu. Wartość ta jest mierzona za pomocą 64-bitowych przyrostów.
- ♦ Czas Życia (ang. *TTL* – *Time to Live*): Pakiet IP nie może „włączyć się” w nieskończoność po sieci WAN. Musi mieć ograniczoną liczbę skoków, które może wykonać (patrz niżej). Wartość 8-bitowego pola TTL jest zwiększana o jeden przy każdym skoku, jaki pakiet wykonuje. Gdy osiągnie wartość maksymalną, pakiet jest niszczone.



Pakiety IP są trasowane przez różne sieci za pomocą urządzeń znanych jako routery. Każdy router, przez który przechodzi pakiet, jest liczony jako jeden skok. Ustalenie maksymalnej liczby skoków zapewnia, że pakiety nie będą stale wykonywać pętli w dynamicznie trasowanej sieci.

- ♦ Protokół: 8-bitowe pole identyfikujące protokół, następujący po nagłówku IP, taki jak VINES, TCP, UDP itd.
- ♦ Suma Kontrolna: 16-bitowe pole kontroli błędów. Komputer docelowy lub jakikolwiek inny węzeł bramy w sieci, może powtórzyć działania matematyczne na zawartości pakietu, przeprowadzone wcześniej przez komputer źródłowy. Jeśli dane po drodze nie uległy zmianie, wyniki obydwu obliczeń są identyczne. Pole sumy kontrolnej informuje również komputer docelowy o ilości przychodzących danych.
- ♦ Adres Źródłowy IP: jest adresem IP komputera źródłowego.
- ♦ Adres Docelowy IP: jest adresem IP komputera docelowego.
- ♦ Wypełnienie: do tego pola dodawane są zera, tak aby długość nagłówka TCP była zawsze wielokrotnością 32 bitów.

Te pola nagłówka świadczą o tym, że protokół IPv4 warstwy Internetu jest protokołem bezpołączeniowym – urządzenia kierujące pakietem w sieci mogą samodzielnie ustalać idealną ścieżkę przejścia przez sieć dla każdego pakietu. Nie występują również żadne potwierdzenia, sterowanie strumieniem danych czy też funkcje porządkowania kolejności, właściwe protokołom wyższych warstw, takim jak TCP. Protokół IPv4 pozostawia te funkcje protokołom wyższego poziomu.

Warstwa Internetu musi także obsługiwać inne funkcje zarządzania trasą oprócz formatowania pakietów IP. Musi zapewnić mechanizmy tłumaczące adresy warstwy 2 na adresy warstwy 3 i odwrotnie. Te funkcje zarządzania trasą są dostarczane przez protokoły równorzędne z IP; protokoły trasujące opisane w rozdziale 1 pt. „ABC sieci”. Są to: wewnętrzny protokół bramowy (ang. *IGP* – *Interior Gateway Protocols*), zewnętrzne protokoły bramowe (ang. *EGP* – *Exterior Gateway Protocols*), protokół rozróżniania adresów (ang. *ARP* – *Address Resolution Protocol*), odwrotny protokół rozróżniania adresów (ang. *RARP* – *Reverse Address Resolution Protocol*) i protokół komunikacyjny sterowania Internetem (ang. *ICMP* – *Internet Control Message Protocol*).

Typowe działanie protokołu IPv4

Warstwa aplikacji opatruje pakiet danych nagłówkiem, identyfikując docelowy host i port. Protokół warstwy host-z-hostem (TCP lub UDP, w zależności od aplikacji) dzieli ten blok danych na mniejsze, łatwiej dające sobie radę kawałki. Do każdego kawałka dołączony jest nagłówek. Taką strukturę nazywa się „segmentem TCP”.

Pola nagłówka segmentu są odpowiednio wypełniane, a segment jest przekazywany do warstwy Internetu. Warstwa Internetu dodaje informacje dotyczące adresowania, rodzaju protokołu (TCP lub UDP) i sumy kontrolnej. Jeśli segment był fragmentowany, warstwa Internetu wypełnia również to pole.

Komputer docelowy odwraca właśnie opisane działania. Odbiera pakiety i przekazuje je swojemu protokołowi warstwy host-z-hostem do ponownego złożenia. Jeśli to konieczne, pakiety są ponownie grupowane w segmenty danych, przekazywane odpowiedniej aplikacji.

Schemat adresowania protokołu IP

Protokół IPv4 wykorzystuje 32-bitowy, binarny schemat adresowania, w celu identyfikowania sieci, urządzeń sieciowych i komputerów przyłączonych do sieci. Adresy te, znane jako adresy IP, są ściśle regulowane przez internetowe centrum informacji sieciowej (ang. *InterNIC* – *Internet Network Information Center*). Choć administrator sieci ma możliwość dowolnego wybierania nie zarejestrowanych adresów IP, taka praktyka jest niewybaczalna. Komputery mające takie „podrobione” adresy IP mogą działać prawidłowo tylko w obrębie swej własnej domeny. Próby dostępu do Internetu z pewnością wykażą ograniczenia takiego krótkowzrocznego działania. Skutki mogą być bardzo różne w zależności od wielu rozmaitych czynników, ale na pewno będą to skutki niepożądane.

Każda z pięciu klas adresów IP jest oznaczona literą alfabetu: klasa A, B, C, D i E. Każdy adres składa się z dwóch części: adresu sieci i adresu hosta. Klasy prezentują odmienne uzgodnienia dotyczące liczby obsługiwanych sieci i hostów. Choć są to adresy binarne, zwykle przedstawia się je w tzw. formacie dziesiętnym kropkowym (np. 135.65.121.6), aby ułatwić człowiekowi ich używanie. Kropki rozdzielają cztery oktety adresu.



Notacja dziesiętna kropkowa odnosi się do konwersji adresu binarnego na dziesiętny system liczbowy. Kropka („.”) służy do oddzielania numerów węzła i sieci. Na przykład, 100.99 odnosi się do urządzenia 99 w sieci 100.

- ♦ Adres IP klasy A: Pierwszy bit adresu klasy A jest zawsze ustawiony na „0”. Następne siedem bitów identyfikuje numer sieci. Ostatnie 24 bity (np. trzy liczby dziesiętne oddzielone kropkami) adresu klasy A reprezentują możliwe adresy hostów. Adresy klasy A mogą mieścić się w zakresie od 1.0.0.0 do 126.0.0.0. Każdy adres klasy A może obsłużyć $16777214 (=2^{24}-2)$ unikatowych adresów hostów.
- ♦ Adres IP klasy B: Pierwsze dwa bity adresu klasy B to „10”. Następne 16 bitów identyfikuje numer sieci, zaś ostatnie 16 bitów identyfikuje adresy potencjalnych hostów. Adresy klasy B mogą mieścić się w zakresie od 128.1.0.0 do 191.254.0.0. Każdy adres klasy B może obsłużyć $65534 (=2^{16}-2)$ unikatowych adresów hostów.
- ♦ Adres IP klasy C: Pierwsze trzy bity adresu klasy C to „110”. Następne 21 bitów identyfikuje numer sieci. Ostatni oktet służy do adresowania hostów. Adresy klasy C mogą mieścić się w zakresie od 192.0.1.0 do 223.255.254.0. Każdy adres klasy C może obsłużyć $254 (=2^8-2)$ unikatowe adresy hostów.
- ♦ Adres IP klasy D: Pierwsze cztery bity adresu klasy D to „1110”. Adresy te są wykorzystywane do multicastingu, ale ich zastosowanie jest ograniczone. Adres multicast jest unikatowym adresem sieci, kierującym pakiety do predefiniowanych grup adresów IP. Adresy klasy D mogą pochodzić z zakresu 224.0.0.0 do 239.255.255.254.



Pewna niejasność definicji klasy D adresu IP przyczynia się do potencjalnej rozbieżności pomiędzy jej rozumieniem a stanem faktycznym. Choć IETF zdefiniowało klasy C i D jako oddzielne, różniące się pod względem zakresów liczbowych i zamierzonej funkcjonalności, to wcale nie tak rzadko zdarza się, że zakres adresu klasy D jest utożsamiany z zakresem adresu klasy C. Jest to podejście nieprawidłowe – ale najwidoczniej narzucone przez pewne kursy certyfikacyjne.

- ♦ Adres IP klasy E: Faktycznie – zdefiniowano klasę E adresu IP, ale InterNIC zarezerwował go dla własnych badań. Tak więc żadne adresy klasy E nie zostały dopuszczone do zastosowania w Internecie.

Duże odstępy między tymi klasami adresów marnowały znaczną liczbę potencjalnych adresów. Rozważmy dla przykładu średnich rozmiarów przedsiębiorstwo, które potrzebuje 300 adresów IP. Adres klasy C (254 adresy) jest niewystarczający. Wykorzystanie dwóch adresów klasy C dostarczy więcej adresów niż potrzeba, ale w wyniku tego w ramach przedsiębiorstwa powstaną dwie odrębne domeny. Z kolei zastosowanie adresu klasy B zapewni potrzebne adresy w ramach jednej domeny, ale zmarnuje się w ten sposób $65534 - 300 = 65234$ adresy.

Na szczęście nie będzie to już dłużej stanowić problemu. Został opracowany nowy, międzydomenowy protokół trasujący, znany jako bezklasowe trasowanie międzydomenowe (ang. *CIDR – Classless Interdomain Routing*), umożliwiający wielu mniejszym klasom adresowym działanie w ramach jednej domeny trasowania.

Adresowanie IP wymaga, by każdy komputer miał własny, unikalny adres. Maski podsieci mogą kompensować ogromne odstępstwa między klasami adresowymi, dostosowując długość adresów hosta i/lub sieci. Za pomocą tych dwóch adresów można trasować dowolny datagram IP do miejsca przeznaczenia.

Ponieważ TCP/IP jest w stanie obsługiwać wiele sesji z pojedynczego hosta, musi on zapewnić możliwość adresowania specyficznych programów komunikacyjnych, które mogą działać na każdym z hostów. TCP/IP wykorzystuje do tego numery portów. IETF przypisało kilku najbardziej powszechnym aplikacjom ich własne, dobrze znane numery portów. Numery te są stałe dla każdej aplikacji na określonym hoście. Innym aplikacjom przypisuje się po prostu dostępny numer portu.

Wnioski dotyczące IPv4

Protokół IPv4 ma już prawie dwadzieścia lat. Od jego początków Internet przeszedł kilka znaczących zmian, które zmniejszyły efektywność IP jako protokołu uniwersalnej przyłączalności. Być może najbardziej znaczącą z tych zmian była komercjalizacja Internetu. Przyniosła ona bezprecedensowy wzrost populacji użytkowników Internetu. To z kolei stworzyło zapotrzebowanie na większą liczbę adresów, a także potrzebę obsługi przez warstwę Internetu nowych rodzajów usług. Ograniczenia IPv4 stały się bodźcem dla opracowania zupełnie nowej wersji protokołu. Jest ona nazywana IP, wersja 6 (IPv6), ale powszechnie używa się również nazwy *Następna generacja protokołu Internetu* (ang. *IPng – next generation of Internet Protocol*).

Protokół Internetu, wersja 6 (IPv6)

Protokół IPv6 ma być prostą, kompatybilną „w przód” nowelizacją istniejącej wersji protokołu IP. Intencją przyświecającą tej nowelizacji jest wyeliminowanie wszystkich słabości ujawniających się obecnie w protokole IPv4, w tym zbyt małej liczby dostępnych adresów IP, niemożności obsługi ruchu o wysokich wymaganiach czasowych i braku bezpieczeństwa w warstwie sieci.



Protokół IPv6 był pierwotnie określany jako „IP: następna generacja” lub „Ipng” – co przydawało mu nieco tajemniczości z pogranicza science fiction. Podczas opracowywania specyfikacji protokół ten otrzymał oficjalną nazwę „IP wersja 6” (IPv6).

Dodatkowym bodźcem dla opracowania i rozwoju nowego protokołu IP stało się trasowanie, które w ramach protokołu IPv4 jest skrupowane jego 32-bitową architekturą adresową, dwupoziomą hierarchią adresowania i klasami adresowymi. Dwupozioma hierarchia

adresowania „host.domena” po prostu nie pozwala konstruować wydajnych hierarchii adresowych, które mogłyby być agregowane w routerach na skalę odpowiadającą dzisiejszym wymaganiom globalnego Internetu.

Następna generacja protokołu IP – IPv6 – rozwiązuje wszystkie wymienione problemy. Będzie oferować znacznie rozszerzony schemat adresowania, aby nadążyć za stałą ekspansją Internetu, a także zwiększoną zdolność agregowania tras na wielką skalę.

IPv6 będzie także obsługiwać wiele innych właściwości, takich jak: transmisje audio i/lub wideo w czasie rzeczywistym, mobilność hostów, bezpieczeństwo końcowe (czyli na całej długości połączenia) dzięki mechanizmom warstwy Internetu – kodowaniu i identyfikacji, a także autokonfiguracja i autorekonfiguracja. Oczekuje się, że usługi te będą odpowiednią zachętą dla migracji, gdy tylko staną się dostępne produkty zgodne z IPv6. Wiele z tych rozwiązań wciąż wymaga dodatkowej standaryzacji, dlatego też przedwcześnie byłoby ich obszerne omawianie.

Jedynym jednakże aspektem protokołu IPv6, który wymaga szerszego omówienia, jest adresowanie. 32-bitowa długość adresu w protokole IPv4 teoretycznie umożliwiała zaadresowanie około 4 miliardów ($2^{32}-1$) urządzeń. Niewydajne podsieciowe techniki maskowania i inne rozrzutne praktyki roztrwoniły niestety ów zasób.

Protokół IPv6 wykorzystuje adresy 128-bitowe i teoretycznie jest w stanie zwiększyć przestrzeń adresową protokołu o czynnik 2^{96} – co daje astronomiczną liczbę

340.282.366.920.938.463.463.374.607.431.768.211.456

potencjalnych adresów. Obecnie zajęte jest około 15% tej przestrzeni adresowej – reszta jest zarezerwowana dla – bliżej nie określonych – przyszłych zastosowań.

W rzeczywistości przypisanie i trasowanie adresów wymaga utworzenia ich hierarchii. Hierarchie mogą zmniejszyć liczbę potencjalnych adresów, ale za to zwiększają wydajność protokołów trasujących zgodnych z IPv6. Jedną z praktycznych implikacji długości adresu IPv6 jest to, że usługa nazwy domeny (ang. *DNS – Domain Name Service*), stanowiąca w wersji IPv4 jedynie wygodny luksus, staje się absolutną koniecznością.



Usługa nazwy domeny jest narzędziem sieciowym, odpowiedzialnym za tłumaczenie (wygodnych dla użytkowników) mnemonicznych nazw hostów na numeryczne adresy IP.

Równie znacząca, jak zwiększona potencjalna przestrzeń adresowa, jest jeszcze większa elastyczność, na jaką pozwalają nowe struktury adresowe IPv6. Protokół ten uwalnia się od adresowania bazującego na klasach. Zamiast tego rozpoznaje on trzy rodzaje adresów typu *unicast*, adres klasy D zastępuje nowym formatem adresu *multicast* oraz wprowadza nowy rodzaj adresu; przed przejściem do dalszej części wykładu nieodzowne staje się wyjaśnienie szczegółów tych koncepcji.

Struktury adresów *unicast* IPv6

Adresowanie *unicast* zapewnia przyłączalność od jednego urządzenia końcowego do drugiego. Protokół IPv6 obsługuje kilka odmian adresów *unicast*.

Adres dostawcy usług internetowych (ISP)

Podczas gdy protokół IPv4 z góry przyjął grupy użytkowników wymagających przyłączalności, IPv6 dostarcza format adresu *unicast*, specjalnie przeznaczony dla dostawców usług internetowych, w celu przyłączania indywidualnych użytkowników do Internetu. Te oparte na dostawcach adresy *unicast* oferują unikatowe adresy dla indywidualnych użytkowników lub małych grup, uzyskujących dostęp do Internetu za pośrednictwem dostawcy usług internetowych. Architektura adresu zapewnia wydajną agregację tras w środowisku użytkowników indywidualnych.

Format adresu *unicast* ISP jest następujący:

- ◆ 3-bitowa flaga adresu *unicast* ISP, zawsze ustawiana na „010”
- ◆ Pole ID rejestru, o długości „n” bitów
- ◆ Pole ID dostawcy, o długości „m” bitów
- ◆ Pole ID abonenta, o długości „o” bitów
- ◆ Pole ID podsieci, o długości „p” bitów
- ◆ Pole ID interfejsu, o długości $128-3-(n+m+o+p)$ bitów

Litery n,m,o,p oznaczają zmienne długości pól. Długość pola ID interfejsu stanowi różnicę długości adresu (128) i łącznej długości pól poprzedzających, wraz z trójbitową flagą.

Przykładem adresu tego typu może być 010:0:0:0:x, gdzie „x” może być dowolną liczbą. Ponieważ większość nowej przestrzeni adresowej dopiero musi zostać przypisana, adresy te będą zawierać mnóstwo zer. Dlatego grupy zer mogą być zapisywane skrótem w postaci podwójnego dwukropka (::) – skróconą formą adresu 010:0:0:0:x jest więc 010::x.

Inne rodzaje adresów *unicast* są przeznaczone do użytku lokalnego. Adresy użytku lokalnego mogą być przypisane do urządzeń sieciowych w samodzielnym Intranecie lub do urządzeń w Intranecie, którym potrzebny jest dostęp do Internetu.

Adres użytku lokalnego dla łącza

Adres użytku lokalnego dla łącza jest przeznaczony dla pojedynczego łącza, do celów takich jak konfiguracja auto-adresu, wykrywanie sąsiadów, a także w przypadku braku routerów. Adresy lokalne dla łącza mają następujący format:

- ◆ 10-bitowa flaga adresu lokalnego, zawsze ustawiana na „1111111011”
- ◆ Zarezerwowane, nienazwane pole, mające długość „n” bitów, ale ustawiane domyślnie na wartość „0”
- ◆ Pole ID interfejsu o długości $118-n$ bitów

ID interfejsu może być adresem MAC karty sieciowej Ethernetu. Adresy MAC, będące teoretycznie adresami unikalnymi, mogą być skojarzone z przedrostkami standardowego adresu IP w celu utworzenia unikalnych adresów dla mobilnych lub zastępczych użytkowników. Przykładem adresu użytku lokalnego dla łącza z adresem MAC mógłby być: 111111011:0:adres_mac.

Adres użytku lokalnego dla miejsca

Adresy lokalne dla miejsca są przeznaczone do stosowania w pojedynczym miejscu. Mogą być używane w miejscach lub organizacjach, które nie są przyłączone do globalnego Internetu. Nie muszą żądać czy też „kraść” przedrostka adresu z przestrzeni adresowej globalnego Internetu. Zamiast tego mogą używać adresów protokołu IPv6 lokalnych dla miejsca. Gdy organizacja łączy się z globalnym Internetem, może utworzyć unikatowe adresy globalne, zastępując przedrostek lokalny dla miejsca przedrostkiem abonenta, zawierającym identyfikatory rejestru, dostawcy i abonenta.

Adresy lokalne dla miejsca mają następujący format:

- ♦ 10-bitowa flaga użytku lokalnego, zawsze ustawiana na „111111011”
- ♦ Zarezerwowane, nienazwane pole, mające długość „n” bitów, ale ustawiane domyślnie na wartość „0”
- ♦ Pole ID podsieci o długości „m” bitów
- ♦ Pole ID interfejsu o długości $118 - (n+m)$ bitów

Przykładem adresu lokalnego dla miejsca jest: 111111011:podsieć:interfejs.

Struktury zastępczych adresów *unicast* IPv6

Dwa specjalne adresy unicast protokołu IPv6 zostały określone jako mechanizmy przejściowe, umożliwiające hostom i routerom dynamiczne trasowanie pakietów IPv6 przez infrastrukturę sieci protokołu IPv4 i na odwrót.

Adres *unicast* IPv6 zgodny z IPv4

Pierwszy typ adresu unicast nosi nazwę „adres IPv6 zgodny z IPv4”. Ten zastępczy adres unicast może być przypisywany węzłom IPv6, a jego ostatnie 32 bity zawierają adres IPv4. Adresy takie mają następujący format:

80 bitów	16 bitów	32 bity
000...0000	00...00	adres IPv4

Adres *unicast* IPv6 wzorowany na IPv4

Drugi, podobny typ adresu IPv6, również zawierający adres IPv4 w ostatnich 32 bitach, jest znany jako „adres IPv6 wzorowany na IPv4”. Adres ten jest tworzony przez router

o podwójnym protokole i umożliwia węzłom pracującym wyłącznie z protokołem IPv4 tunelowanie przez infrastrukturę sieci z protokołem IPv6. Jedyną różnicą między adresami unicast IPv6 wzorowanymi na IPv4 a adresami unicast IPv6 zgodnymi z IPv4 jest taka, że adresy wzorowane na IPv4 to adresy tymczasowe. Są one automatycznie tworzone przez routery o podwójnym protokole i nie mogą być przypisane do żadnego węzła. Format takiego adresu wygląda następująco:

80 bitów	16 bitów	32 bity
000...0000	FF...FF	adres IPv4

Obydwa adresy unicast, zarówno wzorowany na IPv4, jak i zgodny z IPv4, mają zasadnicze znaczenie dla tunelowania. Tunelowanie umożliwia przesyłanie pakietów przez niedostępny w inny sposób rejon sieci dzięki umieszczaniu pakietów w obramowaniu akceptowalnym na zewnątrz.

Struktury adresów *anycast* IPv6

Adres anycast, wprowadzony w protokole IPv6, jest pojedynczą wartością przypisaną do więcej niż jednego interfejsu. Zwykle interfejsy te należą do różnych urządzeń. Pakiet wysłany pod adres anycast jest trasowany tylko do jednego urządzenia. Jest on wysyłany do najbliższego – według zdefiniowanej przez protokoły trasujące miary odległości – interfejsu o tym adresie. Na przykład, strona WWW (World Wide Web) może być powielona na kilku serwerach. Dzięki przypisaniu tym serwerom adresu anycast żądania połączenia z tą stroną WWW są automatycznie trasowane do tylko jednego serwera – najbliższego względem użytkownika.



W środowisku trasowanym „najbliższy” interfejs może nie być tym, który jest najbliższy w sensie fizycznego położenia. Routery wykorzystują przy obliczaniu tras zaskakująco szeroki zestaw metryk. Określanie najkrótszej trasy jest uzależnione od aktualnie używanego protokołu trasującego oraz od jego metryk.

Adresy anycast są tworzone (pobierane) z przestrzeni adresów unicast i mogą przybrać formę dowolnego typu adresu unicast. Tworzy się je, przypisując po prostu ten sam adres unicast więcej niż jednemu interfejsowi.

Struktury adresów *multicast* IPv6

Protokół IPv4 obsługiwał multicasting, ale wymagało to stosowania niejasnego adresowania klasy D. Protokół IPv6 rezygnuje z adresów klasy D na korzyść nowego formatu adresu, udostępniającego tryliony możliwych kodów grup multicast. Każdy kod grupy identyfikuje dwóch lub więcej odbiorców pakietu. Zakres pojedynczego adresu multicast jest elastyczny. Każdy adres może być ograniczony do pojedynczego systemu, do określonego miejsca, powiązany z danym łączem sieciowym lub rozpowszechniany globalnie.

Należy zauważyć, że nadawanie adresów IP również zostało wyeliminowane i zastąpione nowym multicastingowym formatem adresu.

Wnioski dotyczące IPv6

Pomimo potencjalnych korzyści związanych z protokołem IPv6, migracja z IPv4 nie jest wolna od ryzyka. Rozszerzenie długości adresu z 32 do 128 bitów automatycznie ogranicza współoperacyjność protokołów IPv4 i IPv6. Węzły „tylko-IPv4” nie mogą współdziałać z węzłami „tylko-IPv6”, ponieważ architektury adresowe nie są kompatybilne w przód. To ryzyko biznesowe, w połączeniu z nieustanną ewolucją protokołu IPv4, może stanowić przeszkodę dla rynkowej akceptacji protokołu IPv6.

Wymiana IPX/SPX Novell

Zestaw protokołów firmy Novell bierze nazwę od swoich dwóch głównych protokołów: międzysieciowej wymiany pakietów (ang. *IPX-Internet Packet Exchange*) i sekwencyjnej wymiany pakietów (ang. *SPX-Sequenced Packet Exchange*). Ten firmowy stos protokołów został oparty na protokole systemów sieciowych firmy Xerox (ang. *XNS – Xerox’s Network System*), wykorzystywanym w pierwszej generacji sieci Ethernet. Wymiana IPX/SPX zyskała na znaczeniu we wczesnych latach 80. jako integralna część systemu Novell Netware. Netware stał się *faktycznym* standardem sieciowego systemu operacyjnego (ang. *NOS – Network Operating System*) dla sieci lokalnych pierwszej generacji. Novell uzupełnił swój system zestawem aplikacji biznesowych i klienckich narzędzi łączności.

Protokół IPX w dużym stopniu przypomina IP. Jest bezpołączeniowym protokołem datagramowym, który nie wymaga ani nie zapewnia potwierdzenia każdego transmitowanego pakietu. Protokół IPX polega na SPX w taki sam sposób, w jaki protokół IP polega na TCP w zakresie porządkowania kolejności i innych usług połączeniowych warstwy 4. Rysunek 12.4 przedstawia stos protokołów IPX/SPX w porównaniu z modelem referencyjnym OSI.

Protokoły IPX i SPX Novella są funkcjonalnym ekwiwalentem warstw modelu OSI, odpowiednio warstw 3 i 4. Pełny zestaw protokołów IPX/SPX, składający się z czterech warstw, funkcjonalnie odpowiada innym warstwom modelu OSI.

Analiza IPX/SPX

Stos protokołów IPX/SPX obejmuje cztery warstwy funkcjonalne: dostępu do nośnika, łącza danych, Internetu i aplikacji. Te cztery warstwy luźno nawiązują do siedmiu warstw modelu referencyjnego OSI, nie tracąc nic na funkcjonalności.

Rysunek 12.4.

Porównanie modelu
OSI z modelem
IPX/SPX.

Nazwa warstwy modelu referencyjnego OSI	Numer warstwy OSI	Opis równoważnej warstwy IPX/SPX				
Aplikacji	7	R I P	S A P	N C P	N L S P	Rozmaite protokoły
Prezentacji	6					
Sesji	5					
Transportu	4	SPX				
Sieci	3	Międzysieciowa wymiana pakietów				
Łącza danych	2	Interfejs otwartego łącza danych				
Fizyczna	1	Dostęp do nośnika				

Warstwa aplikacji

Warstwa aplikacji Novella obejmuje trzy warstwy – aplikacji, prezentacji i sesji – modelu OSI, choć niektóre z jej protokołów aplikacyjnych rozciągają ten stos w dół, aż do warstwy sieci. Głównym protokołem warstwy aplikacji w tym stosie jest protokół rdzenia NetWare (ang. *NCP – NetWare Core Protocol*). Protokół NCP można bezpośrednio sprzęgać zarówno z protokołem SPX, jak i IPX. Jest wykorzystywany do drukowania, współdzielenia plików, poczty elektronicznej i dostępu do katalogów.

Innymi protokołami warstwy aplikacji są między innymi: protokół informacyjny trasowania (ang. *RIP – Routing Information Protocol*), firmowy protokół ogłoszeniowy usługi (ang. *SAP – Service Advertisement Protocol*) i protokół obsługi łącza systemu Netware (ang. *NLSP – Netware Link Services Protocol*).

Protokół RIP jest domyślnym protokołem trasującym systemu NetWare. Jest to protokół trasowania wektora odległości wykorzystujący tylko dwie metryki: kwanty (ang. *ticks*) i skoki (ang. *hops*). Kwant jest miarą czasu, zaś liczba skoków, jak już wyjaśniono wcześniej w tym rozdziale, jest licznikiem routerów, które manipulowały trasowanym pakietem. Na tych dwóch metrykach opiera się wybór ścieżki trasowania protokołu IPX. Podstawową metryką są kwanty – skoki rozstrzygają tylko w przypadku, gdy dwie ścieżki (lub więcej) mają taką samą wartość znaków kontrolnych.

RIP jest bardzo prostym protokołem trasującym. Oprócz ograniczonej liczby metryk wektora odległości, cechuje się też wysokim poziomem narzutu sieciowego. Narzut ten powstaje, ponieważ aktualizacje tabeli trasującej RIP są nadawane co 60 sekund. W wielkich lub mocno obciążonych sieciach taka szybkość aktualizacji może mieć szkodliwe działanie.

SAP jest unikatowym protokołem firmowym, który Novell udanie zastosował do polepszenia związku klienta z serwerem. Serwery wykorzystują protokół SAP do automatycznego wysyłania w sieć informacji o udostępnianych przez nie usługach natychmiast po tym, jak uaktywnią się w sieci. Oprócz tego okresowo nadają informacje SAP, aby dostarczać klientom i innym serwerom informacje o swoim statusie i usługach.

Transmisje SAP generowane przez serwer informują o statusie i usługach tego serwera. Transmisje zawierają nazwę i typ serwera, jego status operacyjny, a także numery sieci, węzła i gniazda. Routery mogą przechowywać informacje z transmisji SAP i rozprowadzać je do innych segmentów sieci. Klienci także mogą inicjować zgłoszenie SAP, gdy potrzebują określonej usługi. Ich żądanie jest rozsyłane po całym segmencie sieci. Hosty mogą wtedy odpowiedzieć i dostarczyć klientowi informacje SAP wystarczające do określenia, czy usługa jest dostępna w rozsądnej odległości.

Niestety, SAP jest dojrzałym protokołem, który coraz gorzej funkcjonuje we współczesnych sieciach. Tak jak w przypadku protokołu RIP, ogłoszenia o usługach są nadawane co 60 sekund. Przy dzisiejszych ogromnych, jednorodnych, komutowanych sieciach LAN, taka częstość nadawania może być problematyczna.

Najnowszym protokołem warstwy aplikacji jest protokół obsługi łącza systemu Netware (NLSP). Jest to protokół trasowania w zależności od stanu łącza, którym Novell zamierza zastąpić starzejące się protokoły RIP i SAP. Protokół NLSP aktualizuje trasy tylko wtedy, gdy zaszły jakieś zmiany.

Protokoły warstwy Internetu

Warstwa Internetu wymiany IPX/SPX luźno nawiązuje do warstw sieci i transportu modelu referencyjnego OSI. IPX jest w przeważającej części protokołem warstwy 3 (sieci), choć może też być bezpośrednio sprzęgany z warstwą aplikacji. SPX jest wyłącznie protokołem warstwy 4 (transportu) i nie może być bezpośrednio sprzęgnięty z interfejsem ODI warstwy łącza danych. Musi przekazywać dane poprzez protokół IPX sprzęgnięty z ODI. IPX i SPX funkcjonują jako protokoły podwarstw we wspólnej warstwie Internetu.

SPX jest protokołem połączeniowym i może być wykorzystywany do przesyłania danych między klientem serwerem, dwoma serwerami czy nawet dwoma klientami. Tak jak w przypadku protokołu TCP, protokół SPX zapewnia niezawodność transmisjom IPX, zarządzając (administrując) połączeniem i udostępniając sterowanie strumieniem danych, kontrolę błędów i porządkowanie kolejności pakietów.

Nagłówek SPX ma następujący rozmiar i strukturę:

- ♦ Sterowanie połączeniem: Pierwszy oktet (8 bitów) nagłówka SPX zawiera cztery 2-bitowe flagi, sterujące dwukierunkowym przepływem danych przez połączenie SPX.
- ♦ Typ strumienia danych: Następnych osiem bitów nagłówka definiuje typ strumienia danych.
- ♦ Identyfikacja połączenia źródłowego: 16-bitowe pole identyfikacji połączenia źródłowego identyfikuje proces odpowiedzialny za inicjowanie połączenia.
- ♦ Identyfikacja połączenia docelowego: 16-bitowe pole identyfikacji połączenia docelowego służy do identyfikowania procesu, który zaakceptował żądanie (zgłoszenie) połączenia SPX.

- ◆ Numer sekwencji: 16-bitowe pole numeru sekwencji dostarcza protokołowi SPX hosta docelowego informację o liczbie wysłanych pakietów. To sekwencyjne numerowanie może być wykorzystywane do zmiany kolejności odebranych pakietów, gdyby przybyły w niewłaściwej kolejności.
- ◆ Numer potwierdzenia: 16-bitowe pole numeru potwierdzenia wskazuje następny oczekiwany segment.
- ◆ Liczba alokacji: 16-bitowe pole liczby alokacji jest wykorzystywane do śledzenia liczby pakietów wysłanych, ale nie potwierdzonych przez odbiorcę.
- ◆ Dane: Ostatnie pole nagłówka SPX zawiera dane. W jednym pakiecie SPX można przesłać do 534 oktetów danych.

Protokołem warstwy sieci dla sieci Novell jest IPX. Protokół ten zapewnia bezpołączeniowe usługi dostarczania datagramów. Przygotowuje pakiety protokołu SPX (lub pakiety innych protokołów) do dostarczenia przez wiele sieci, dołączając do nich nagłówki IPX. W ten sposób powstaje struktura zwana datagramem IPX. Nagłówek tego datagramu zawiera wszystkie informacje niezbędne do skierowania pakietów do miejsca przeznaczenia, niezależnie od tego, gdzie mogłoby się ono znajdować.

Długość nagłówka IPX wynosi 11 oktetów. Jego struktura obejmuje następujące pola:

- ◆ Suma Kontrolna: Nagłówek IPX zaczyna się od 16-bitowego pola dziedziczenia, które istnieje tylko po to, aby zapewnić kompatybilność wsteczną z protokołem XNS. Protokół XNS wykorzystywał to pole do kontrolowania błędów, ale IPX domyślnie ustawia to pole na „FFFFH”, a wykrywanie (i korekcję) błędów transmisji pozostawia protokołom wyższego poziomu.
- ◆ Długość Pakietu: 16-bitowe pole określające długość datagramu IPX, wliczając nagłówki i dane. Pole to jest sprawdzane w celu weryfikacji integralności pakietu.
- ◆ Sterowanie Transportem: 8-bitowe pole wykorzystywane przez routery podczas przesyłania datagramu. Przed wysłaniem IPX ustawia to pole na „0”. Każdy router, który odbiera i przesyła dalej datagram, zwiększa wartość pola o jeden.
- ◆ Typ Pakietu: 8-bitowe pole identyfikujące typ pakietu zawartego w datagramie IPX. Pole to umożliwia hostowi docelowemu przekazanie zawartości do następnej, odpowiedniej warstwy protokołów. Typy mogą obejmować RIP, NCP, SPX, błąd itd.
- ◆ Numer Sieci Docelowej: 32-bitowe pole określające numer sieci, w której znajduje się węzeł docelowy.
- ◆ Węzeł Docelowy: 48-bitowe pole zawierające numer węzła, w którym znajduje się komputer docelowy.
- ◆ Numer Gniazda Docelowego: Ponieważ IPX umożliwia wiele jednoczesnych połączeń z jednym systemem, istotne jest określenie numeru gniazda procesu lub programu odbierającego pakiety. Informacji takiej dostarcza to 16-bitowe pole.
- ◆ Numer Sieci Źródłowej: 32-bitowe pole określające numer sieci, w której znajduje się węzeł źródłowy.

- ♦ Adres Węzła Źródłowego: 48-bitowe pole zawierające numer węzła, w którym znajduje się komputer źródłowy.
- ♦ Numer Gniazda Źródłowego: 16-bitowe pole, określające numer gniazda procesu lub programu wysyłającego pakiety.

Typowe działanie protokołów IPX/SPX

Protokół SPX tworzy i utrzymuje połączeniowy strumień bitów między dwoma przyłączonymi do sieci urządzeniami. Protokół przyjmuje duże bloki danych z protokołów wyższych warstw i dzieli je na łatwiejsze w kierowaniu kawałki, nie przekraczające długości 534 oktetów. Do danych dołączany jest nagłówek SPX i w ten sposób powstają segmenty danych SPX. Segmenty przekazywane są protokołowi warstwy Internetu, czyli protokołowi IPX. IPX umieszcza segmenty w polu danych swoich pakietów i wypełnia wszystkie pola nagłówka IPX.

Pola nagłówka IPX obejmują adresowanie sieci, długość, sumę kontrolną i inne informacje nagłówkowe. Następnie pakiet przekazywany jest warstwie łącza danych.

Rysunek 12.5 pokazuje umiejscowienie nagłówków IPX i SPX w ramce Ethernet 802.3. Jest to struktura używana do przekazywania danych pomiędzy dwiema podwarstwami warstwy Internetu sieci Novell.

Rysunek 12.5.

Struktura ramki Ethernet 802.3, zawierającej nagłówki IPX/SPX.

7 oktetowa Preambuła	1 oktetowy Ogranicznik początku ramki	6 oktetowy Adres odbiorcy	6 oktetowy Adres nadawcy	2 oktetowe pole Długość	30 oktetowy nagłówek IPX	Nagłówek SPX o zmiennej długości	Pole Dane o zmiennej długości od 46 do 1482 oktetów	4 oktetowa Sekwencja kontrolna ramki
----------------------	---------------------------------------	---------------------------	--------------------------	-------------------------	--------------------------	----------------------------------	---	--------------------------------------

Komputer docelowy odwraca opisane wyżej działania. Odbiera pakiety i przekazuje je własnemu protokołowi SPX do ponownego złożenia. Jeśli to konieczne, pakiety są ponownie grupowane w segmenty danych, przekazywane odpowiedniej aplikacji.

Warstwy łącza danych i dostępu do nośnika

W systemie Netware odpowiednikami warstw fizycznej i łącza danych OSI są warstwy dostępu do nośnika i łącza danych. Warstwa łącza danych jest bezpośrednio kompatybilna ze standardem interfejsu otwartego łącza danych (ODI). Podobnie warstwa dostępu do nośnika jest bezpośrednio kompatybilna ze wszystkimi popularnymi, znormalizowanymi protokołami dostępu do nośnika.

Ta niskopoziomowa zgodność z przemysłowymi standardami otwartymi sprawia, że system Netware ze stosem protokołów IPX/SPX może być implementowany niemal wszędzie.

Adresowanie IPX

Adresy IPX mają długość 10 oktetów (80 bitów). Jest to znacznie więcej niż 32 bity adresu IPv4, ale mniej niż 128 bitów adresu IPv6. Każdy adres składa się z dwóch części składowych: numeru sieci o maksymalnej długości 32 bitów oraz 48-bitowego numeru węzła. Numery te są wyrażane w notacji kropkowo-szesnastkowej. Na przykład, 1a2b.0000.3c4d.5e6d mogłoby być prawidłowym adresem IPX, w którym „1a2b” reprezentuje numer sieci, a „0000.3c4d.5e6d” jest numerem węzła.

Adresy IPX mogą być tworzone przez administratora sieci. Jednakże tak utworzone numery po znalezieniu się w sieci mogą spowodować występowanie konfliktów adresów. Wymyślanie numerów sieci obciąża administratora obowiązkiem ich utrzymywania i administrowania nimi. Lepszym rozwiązaniem jest więc pozyskanie zarejestrowanych numerów sieci IPX od firmy Novell.

Jako numer hosta IPX wykorzystuje się zwykle powszechnie przypisywany adres (adres MAC) karty sieciowej (NIC). Ponieważ adresy te są unikatowe, przynajmniej w teorii i w stopniu zależnym od zapewnienia jakości przez producenta, oferują wygodną i unikatową numerację hostów.

Podobnie jak IP, protokół IPX może obsługiwać wiele jednoczesnych sesji. Stwarza to potrzebę identyfikowania określonego procesu lub programu, który bierze udział w danej sesji. Identyfikację osiąga się dzięki stosowaniu 16-bitowego numeru „gniazda” w nagłówku IPX. Numer gniazda jest analogiczny do numeru portu w protokole TCP/IP.

Wnioski dotyczące IPX/SPX

Firma Novell Inc. zaobserwowała, jak pozycja rynkowa będącego jej własnością stosu protokołów IPX/SPX słabnie pod naporem konkurencji. Gdy dostępne stały się stosy protokołów otwartych, takich jak OSI, IP i inne, pozycja IPX/SPX bardzo na tym ucierpiała. Dostępne w handlu pakiety oprogramowania wspomagającego prace biurowe również wpłynęły na sprzedaż produktów firmy Novell. Będące jej własnością, ściśle połączone ze sobą serie produktów zapewniły początkowy sukces, ale stały się ciężarem w warunkach rynku ceniącego otwartość i współpracę.

Novell zademonstrował swoje zaangażowanie w staraniach o odzyskanie utraconej pozycji, czyniąc IPv6 domyślnym protokołem przyszłych wersji systemu Netware. Aby pomyślnie wprowadzić tę zmianę strategii, Novell musi zapewnić kompatybilność między protokołami IPv6 i IPX/SPX. By osiągnąć ten cel, Novell blisko współpracował z Grupą Roboczą ds. Technicznych Internetu podczas projektowania IPv6. Dzięki temu wiele usług IPX stało się integralną częścią IPv6.

Przygotowawszy grunt pod przyszłość, Novell musi teraz umożliwić bezbolesną migrację obecnego stosu protokołów i zestawu aplikacji do nowego środowiska. Co więcej, powinien także dostarczyć produkty i usługi podnoszące wartość wykorzystywania platformy sieci otwartej. Dla firmy Novell wizją na przyszłość jest dostarczenie usługi katalogów sieciowych (ang. *NDS – Network Directory Service*) i powiązanych produktów dla dwóch grup użytkowników: środowiska Internetu i korporacyjnych intranetów.

Usługa NDS oferuje jeden, globalny, logiczny widok na wszystkie usługi i zasoby sieciowe. Umożliwia to użytkownikom dostęp do tych usług i zasobów po wykonaniu pojedynczego logowania, niezależnie od lokalizacji użytkownika czy zasobów.

Pakiet protokołów AppleTalk firmy Apple

Gdy komputery Apple zyskały większą popularność, a ich użytkownicy zaczęli z nich korzystać w sposób coraz bardziej wyszukany, nieunikniona stała się konieczność połączenia ich w sieć. Nie jest niespodzianką, że sieć opracowana przez Apple jest tak przyjazna użytkownikowi jak komputery tej firmy. Z każdym komputerem Apple sprzedawany jest AppleTalk, czyli stos protokołów pracy sieciowej firmy Apple, a także niezbędny sprzęt.

Przyłączenie do sieci jest równie proste jak wetknięcie wtyczki do złącza sieciowego i włączenie zasilania komputera. AppleTalk jest siecią równoprawną dostarczającą proste funkcje jak wspólne korzystanie z plików i drukarek. Inaczej niż w sieciach klient/serwer, funkcjonalności sieci równoprawnej nie ograniczają żadne sztywne definicje. Każdy komputer może działać jednocześnie jako serwer i klient.

AppleTalk został także przyjęty przez wielu innych producentów systemów operacyjnych. Nierzadko spotyka się możliwość obsługi stosu protokołów AppleTalk na komputerach innych niż Apple. Pozwala to klientom wykorzystywać AppleTalk i komputery Apple do tworzenia lub przyłączania się do istniejących sieci klient/serwer, innych niż sieci Apple.

Analiza AppleTalk

Stos protokołów AppleTalk obejmuje pięć warstw funkcjonalnych: dostępu do sieci, datagramową, sieci, informacji o strefach i aplikacji. Stos protokołów AppleTalk dość wiernie naśladuje funkcjonalność warstw transportu i sesji modelu referencyjnego OSI. Warstwy fizyczna i łącza danych zostały rozbite na wiele odrębnych warstw, specyficznych ze względu na ramki. AppleTalk integruje warstwy aplikacji i prezentacji, tworząc pojedynczą warstwę aplikacji. Rysunek 12.6 przedstawia to powiązanie funkcjonalne.

Stos protokołów AppleTalk odwzorowuje funkcjonalność warstw sieci, transportu i sesji modelu referencyjnego OSI, ale pozostałe cztery warstwy zawiera w dwóch.

Warstwa aplikacji sieci AppleTalk

AppleTalk łączy w pojedynczej warstwie aplikacji funkcjonalność warstw aplikacji i prezentacji modelu OSI. Ponieważ AppleTalk jest dosyć prostym stosem protokołów, warstwę tę zajmuje tylko jeden protokół. Jest to protokół dostępu do plików sieci AppleTalk (ang. *AFP – AppleTalk Filing Protocol*). Protokół AFP dostarcza usługi plików sieciowych aplikacjom istniejącym oddzielnie od stosu protokołów, takim jak poczta elektroniczna, kolejki wydruków itd. Każda aplikacja uruchamiana na komputerze Apple musi przejść przez protokół AFP, jeśli chce wysłać informacje do sieci lub je z niej odebrać.

Rysunek 12.6.

Porównanie modelu
referencyjnego OSI
i AppleTalk.

Nazwa warstwy modelu referencyjnego OSI	Numer warstwy OSI	Opis równoważnej warstwy AppleTalk
Aplikacji	7	Aplikacji
Prezentacji	6	
Sesji	5	Sesji
Transportu	4	Transportu
Sieci	3	Datagramowa
Łącza danych	2	Dostępu do sieci
Fizyczna	1	

Warstwa sesji sieci AppleTalk

Warstwa sesji w sieci AppleTalk obejmuje pięć podstawowych protokołów, dostarczających takie usługi, jak pełnodupleksowa transmisja, logiczne rozróżnianie nazw i adresów, dostęp do drukarki, ustalanie kolejności pakietów i inne.

Pierwszym protokołem warstwy sesji jest protokół strumienia danych sieci AppleTalk (ang. *ADSP – AppleTalk Data Stream Protocol*). Protokół ten dostarcza pełnodupleksowe usługi połączeniowe w wysoce niezawodny sposób, poprzez ustanawianie logicznego połączenia (sesji) pomiędzy dwoma komunikującymi się procesami na komputerach klientów. Protokół ADSP również zarządza tym połączeniem, dostarczając usługi sterowania strumieniem danych, zarządzania kolejnością i potwierdzania transmitowanych pakietów. Protokół ADSP wykorzystuje adresy gniazd do ustanowienia logicznego połączenia procesów. Po ustanowieniu tego połączenia dwa systemy mogą wymieniać dane.

Innym protokołem warstwy sesji sieci AppleTalk jest protokół sesji sieci AppleTalk (ang. *ASP – AppleTalk Session Protocol*). Protokół ten zapewnia niezawodne dostarczanie danych, wykorzystując sekwencyjne zarządzanie sesją, a także usługi transportowe protokołu transportu sieci AppleTalk (ang. *ATP – Apple Talk Transport Protocol*), który jest protokołem warstwy transportu.

Protokół trasowania AppleTalk (ang. *AURP – AppleTalk Update-Based Routing Protocol*) jest wykorzystywany w większych sieciach AppleTalk. Protokół ten służy przede wszystkim do zarządzania trasą i wymianą informacji pomiędzy urządzeniami trasującymi, zwłaszcza routerami bramek zewnętrznych.

Warstwa sesji sieci AppleTalk zawiera także protokół dostępu do drukarki (ang. *PAP – Printer Access Protocol*). Choć protokół ten został pierwotnie opracowany dla administrowania dostępem do drukarek sieciowych, może być wykorzystywany w rozmaitych wymianach danych. Zapewnia dwukierunkową sesję między dwoma urządzeniami, uzupełnioną o sterowanie strumieniem danych i zarządzanie kolejnością.

Ostatnim z protokołów warstwy sesji sieci AppleTalk jest protokół informacji o strefach (ang. *ZIP – Zone Information Protocol*). Zapewnia on mechanizm logicznego grupowania

indywidualnych urządzeń sieciowych z wykorzystaniem nazw przyjaznych dla użytkownika. Te grupy logiczne są nazywane strefami. W rozszerzonej sieci komputery mogą być rozrzucone po wielu sieciach, ciągle będąc zgrupowanymi w strefie. Jednak w małych, nie rozszerzonych sieciach, może być zdefiniowana tylko jedna strefa.

Protokół ZIP korzysta z protokołu wiązania nazw (ang. *NBP – Name Binding Protocol*), który jest protokołem warstwy transportu, do tłumaczenia nazw na numery sieci i węzła, a także z protokołu transportu ATP do aktualizowania informacji o strefach.

Pięć wymienionych protokołów warstwy sesji zapewnia klientom AppleTalk logiczne połączenia i transfery danych między komputerami, niezależnie od tego, jak bardzo są od siebie oddalone.

Warstwa transportu sieci AppleTalk

Warstwa transportu sieci AppleTalk oferuje usługi transportowe wszystkim warstwom wyższych poziomów. W warstwie tej istnieją cztery odrębne protokoły. Najczęściej używanym spośród nich jest protokół transportu AppleTalk (ATP).

Protokół ATP zapewnia niezawodny mechanizm dostarczania pakietów między dwoma komputerami. ATP korzysta z pól sekwencji i potwierdzenia, znajdujących się w nagłówku pakietu, aby zapewnić, że pakiety nie zaginą na drodze do miejsca przeznaczenia.

Kolejnym ważnym protokołem warstwy transportu AppleTalk jest protokół wiązania nazw (NBP). Jak wspominałem wcześniej, NBP wykorzystuje protokół ZIP do tłumaczenia nazw przyjaznych dla użytkownika na rzeczywiste adresy. Protokół NBP przeprowadza faktyczną translację nazw stref na adresy sieci i węzłów. Protokół ten obejmuje cztery podstawowe funkcje:

- ♦ Rejestracja nazwy: Funkcja ta rejestruje unikalną nazwę logiczną w bazie rejestrów NBP.
- ♦ Przeglądanie nazw: Funkcja ta jest udostępniana komputerowi, który prosi o adres innego komputera. Prośba jest zgłaszana i zaspokajana w sposób jawny. Jeśli w prośbie podawana jest nazwa obiektu, protokół NBP zmienia tę nazwę w adres numeryczny. NBP zawsze przystępuje do zaspokajania takich prośb, przeglądając numery węzłów lokalnych. Jeśli żaden z nich nie pasuje, protokół NBP rozsyła prośbę do innych, połączonych ze sobą sieci AppleTalk. Jeśli wciąż nie można znaleźć pasującego adresu, czas prośby mija i proszące urządzenie otrzymuje komunikat o błędzie.
- ♦ Potwierdzenie nazwy: Żądania potwierdzenia są używane do weryfikacji związku obiektu z adresem.
- ♦ Usunięcie nazwy: W każdej sieci urządzenia są czasowo wyłączane lub odłączane. Gdy wystąpi taka sytuacja, wysyłane jest żądanie usunięcia nazwy, a tablice „obiekt-nazwa-adresowanie” są uaktualniane automatycznie.

Kolejnym protokołem warstwy transportu jest protokół echa sieci AppleTalk (ang. *AEP – AppleTalk Echo Protocol*). Służy on do określania dostępności systemu i obliczania czasu transmisji i potwierdzenia przyjęcia (ang. *RTT – Round Trip Transmit Time*).

Ostatnim protokołem warstwy transportu jest protokół utrzymania wyboru trasy (ang. *RTMP – Routing Table Maintenance Protocol*). Ponieważ AppleTalk stosuje w swojej warstwie sieci protokoły trasowane, musi zapewnić zarządzanie (administrowanie) tablicami trasowania. Protokół RTMP dostarcza routerom zawartość dla ich tablic trasowania.

Warstwa datagramowa sieci AppleTalk

Warstwa datagramowa sieci AppleTalk, analogiczna do warstwy 3 (sieci) modelu OSI, zapewnia bezpołączeniowe dostarczanie pakietowanych datagramów. Jest podstawą dla ustanawiania komunikacji i dostarczania danych przez sieć AppleTalk. Warstwa datagramowa jest również odpowiedzialna za zapewnianie dynamicznego adresowania węzłów sieciowych, jak też za rozróżnianie adresów MAC dla sieci IEEE 802.

Podstawowym protokołem tej warstwy jest protokół dostaw datagramów (ang. *DDP – Datagram Delivery Protocol*). Zapewnia on transmisję danych przez wiele sieci w trybie bezpołączeniowym. Dostosowuje swoje nagłówki w zależności od miejsca przeznaczenia przesyłki. Podstawowe elementy pozostają stałe; dodatkowe pola są dodawane w razie potrzeby.

Datagramy, które mają być dostarczone lokalnie (innymi słowy w obrębie tej samej podsieci), wykorzystują tzw. „krótki nagłówek”. Datagramy, które wymagają trasowania do innych podsieci, wykorzystują format „rozszerzonego nagłówka”. Format rozszerzony zawiera adresy sieci i pole licznika skoków.

Nagłówek DDP składa się z następujących pól:

- ♦ Liczba Skoków: Pole to zawiera licznik, zwiększany o jeden po każdym przejściu pakietu przez router. Pole liczby skoków jest wykorzystywane tylko w rozszerzonym nagłówku.
- ♦ Długość Datagramu: Pole zawiera długość datagramu i może służyć do sprawdzenia, czy nie został on uszkodzony podczas transmisji.
- ♦ Suma Kontrolna DDP: Jest to pole opcjonalne. Kiedy jest używane, zapewnia pewniejszą metodę wykrywania błędów niż proste sprawdzanie długości datagramu. Weryfikacja sumy kontrolnej wykrywa nawet niewielkie zmiany zawartości, niezależnie od tego, czy długość datagramu uległa zmianie.
- ♦ Numer Gniazda Źródłowego: To pole identyfikuje proces komunikujący w komputerze, który zainicjował połączenie.
- ♦ Numer Gniazda Docelowego: To pole identyfikuje proces komunikujący w komputerze, który odpowiedział na żądanie (prośbę) połączenia.
- ♦ Adres Źródłowy: Pole zawierające numery sieci i węzła komputera źródłowego. Jest używane tylko w rozszerzonym formacie nagłówka i umożliwia routerom przesyłanie datagramów przez wiele podsieci.
- ♦ Adres Docelowy: Pole zawierające numery sieci i węzła komputera docelowego. Jest używane tylko w rozszerzonym formacie nagłówka i umożliwia routerom przesyłanie datagramów przez wiele podsieci.

- ♦ Typ DDP: Pole identyfikujące zawarty w datagramie protokół wyższej warstwy. Jest wykorzystywane przez warstwę transportu komputera docelowego do identyfikowania odpowiedniego protokołu, do którego powinna być przesłana zawartość.
- ♦ Dane: Pole to zawiera przesyłane dane. Jego rozmiar może wynosić od 0 do 586 oktetów.

Warstwa datagramowa zawiera także protokół używany do przekształcania adresów węzłów w adresy MAC dla komputerów przyłączonych do sieci IEEE 802. Jest to protokół rozróżniania adresów sieci AppleTalk (ang. *AARP – AppleTalk Address Resolution Protocol*). Może być także używany do określania adresu węzła danej stacji. Protokół AARP przechowuje swoje informacje w tablicy odwzorowywania adresów (AMT). Stosownie do dynamicznego przypisywania numerów węzłów, tablica ta jest stale i automatycznie aktualizowana.

Warstwa łączy danych sieci AppleTalk

Warstwa łączy danych sieci AppleTalk odwzorowuje funkcjonalność warstw fizycznej i łączy danych modelu OSI. Funkcjonalność ta jest zintegrowana w podwarstwach specyficznych dla ramek. Na przykład, „EtherTalk” jest protokołem warstwy łączy danych, zapewniającym całkowitą funkcjonalność warstw fizycznej i łączy danych modelu OSI w ramach jednej podwarstwy. Podwarstwa ta umożliwia opakowywanie AppleTalk w strukturze ramki Ethernetu zgodnej z 802.3.

Istnieją podobne podwarstwy AppleTalk dla Token Ringu (znane jako „TokenTalk”) i dla FDDI („FDDITalk”). Protokoły te są nazywane „protokołami dostępu” ze względu na oferowane przez nie usługi dostępu do sieci fizycznej.

EtherTalk używa protokołu dostępu szeregowego, znanego jako „protokół dostępu do łączy EtherTalk” (ang. *ELAP – Ether Talk Link Access Protocol*) do pakowania danych i umieszczania ramek zgodnych z 802.3 w nośniku fizycznym. Taka konwencja nazewnictwa i funkcjonalność protokołu dostępu szeregowego dotyczy również pozostałych protokołów dostępu. Na przykład, TokenTalk korzysta z „protokołu dostępu do łączy TokenTalk” (ang. *TLAP – Token Talk Link Access Protocol*).

Oprócz protokołów dostępu pasujących do standardów przemysłowych, firma Apple oferuje własny protokół sieci lokalnych, należący do warstwy łączy danych. Jest on znany pod nazwą „LocalTalk”. LocalTalk działa z szybkością 230 Kbps, korzystając ze skrętki dwużyłowej. Wykorzystuje, jak można się spodziewać, protokół dostępu do łączy LocalTalk (ang. *LLAP – Local Talk Link Access Protocol*) do składania ramek i umieszczania ich w sieci. Protokół LLAP zawiera również mechanizmy zarządzania dostępem do nośnika, adresowania na poziomie łączy danych, opakowywania danych oraz reprezentacji bitowej dla transmisji ramki.

Schemat adresowania sieci AppleTalk

Schemat adresowania sieci AppleTalk składa się z dwóch części: numeru sieci i numeru węzła.

Numerzy sieci mają zwykle długość 16 bitów, choć w przypadku sieci nie rozszerzonych lub rozszerzonych w małym stopniu może być stosowane numerowanie jednoskładnikowe (8 bitów). Numerzy te muszą być zdefiniowane przez administratora sieci i używane przez AppleTalk do trasowania pakietów między różnymi sieciami. Numer sieci „0” jest zarezerwowany przez protokół do wykorzystania przy pierwszym przyłączaniu nowych węzłów sieci. Numer sieci musi mieć wartość z zakresu od 00000001 do FFFFFFFF.

Numerzy węzłów są liczbami 8-bitowymi – dopuszczalny zakres adresów dla hostów, drukarek, routerów i innych urządzeń wynosi od 1 do 253; numery 0, 254 i 255 są zarezerwowane przez AppleTalk do wykorzystania w rozszerzonych sieciach. Węzły są numerowane dynamicznie przez warstwę łącza danych sieci AppleTalk.

Adresy AppleTalk są wyrażane w notacji kropkowo-dziesiętnej. Jak już wyjaśniono wcześniej w tym rozdziale, adres binarny jest zamieniany na dziesiętny system liczbowy, a kropka (.) służy do oddzielania numerów węzła i sieci. Na przykład, 100.99 odnosi się do urządzenia 99 w sieci 100. Początkowe zera zostały pominięte.

Wnioski dotyczące AppleTalk

AppleTalk jest firmowym stosem protokołów, przeznaczonym specjalnie dla pracujących w sieci komputerów osobistych firmy Apple. Jego przyszłość jest bezpośrednio związana z losami firmy Apple Corporation i kierunkami rozwoju jej technologii. Tak jak w przypadku firmowego stosu protokołów Novella, warstwy fizyczna i łącza danych służą do zapewnienia zgodności z technologiami sieciowymi opartymi na ustanowionych standardach. Jedynym wyjątkiem jest warstwa fizyczna LocalTalk, która może połączyć ze sobą komputery Apple, używając skrętki dwużyłowej przy szybkości do 230 Kbps.

NetBEUI

Ostatnim, zasługującym na omówienie protokołem jest NetBEUI. Ta niewygodna nazwa jest częściowo skrótem, a częściowo akronimem. Oznacza rozszerzony interfejs użytkownika NetBIOS (co z kolei jest skrótem od Podstawowego sieciowego systemu wejścia-wyjścia). Interfejs NetBEUI został opracowany przez IBM i wprowadzony na rynek w 1985 roku. Jest stosunkowo małym, ale wydajnym protokołem komunikacyjnym LAN.

Wizja firmy IBM dotycząca obliczeń rozproszonych zakładała w tamtych czasach segmentację sieci LAN, opartą na potrzebie wspólnej pracy. Poszczególne segmenty obsługiwałyby środowisko powiązane procesami pracy. Dane, do których potrzebny był dostęp, ale znajdujące się poza segmentem, mogły być odnalezione za pomocą pewnego rodzaju bramy aplikacji. Ze względu na takie pochodzenie nie powinno dziwić, że NetBEUI najlepiej nadaje się do małych sieci LAN. Wizja ta wyjaśnia również, dlaczego protokół NetBEUI nie jest trasowalny.

Protokół ten obejmuje warstwy 3 i 4 modelu referencyjnego OSI. Rysunek 12.7 przedstawia odpowiednie porównanie.

Rysunek 12.7.

Porównanie modelu referencyjnego OSI i NetBEUI.

Nazwa warstwy modelu referencyjnego OSI	Numer warstwy OSI	
Aplikacji	7	
Prezentacji	6	
Sesji	5	
Transportu	4	
Sieci	3	NetBEUI
Łącza danych	2	
Fizyczna	1	

Jak dowodzi rysunek 12.7, NetBEUI ustanawia komunikację pomiędzy dwoma komputerami i dostarcza mechanizmy zapewniające niezawodne dostarczenie i odpowiednią kolejność danych.

Ostatnio firma Microsoft wypuściła protokół NetBEUI 3.0. Jest to ważne z kilku powodów. Po pierwsze, wersja 3.0 jest bardziej tolerancyjna dla wolniejszych środków transmisji niż wersje wcześniejsze. Posiada też możliwość w pełni automatycznego dostrajania się. Najbardziej znaczącą zmianą w NetBEUI 3.0 jest wyeliminowanie samego protokołu NetBEUI. W sieciowych systemach operacyjnych firmy Microsoft został on zastąpiony protokołem ramki NetBIOS (ang. *NBF – NetBIOS Frame*). Zarówno NetBEUI, jak i NBF są ściśle związane z NetBIOS. Dlatego NetBEUI 3.0 (NBF) jest całkowicie kompatybilny i może współpracować z wcześniejszymi wersjami Microsoft NetBEUI.



NetBEUI, niezależnie od wersji, jest integralną częścią sieciowych systemów operacyjnych firmy Microsoft. Jeśli podejmiesz próbę uruchomienia systemu Windows NT 3.x (lub wyższego), Windows for Workgroups 3.11 czy nawet LAN Manager 2.x bez zainstalowanego protokołu NetBEUI, komputer nie będzie mógł się komunikować.

Wnioski dotyczące NetBEUI

NetBEUI jest wyłącznie protokołem transportu sieci LAN dla systemów operacyjnych firmy Microsoft. Nie jest trasowalny. Dlatego jego implementacje ograniczają się do domen warstwy 2, w których działają wyłącznie komputery wykorzystujące systemy operacyjne firmy Microsoft. Aczkolwiek staje się to coraz mniejszą przeszkodą, to jednak skutecznie ogranicza dostępne architektury obliczeniowe i aplikacje technologiczne.

Zalety korzystania z protokołu NetBEUI są następujące:

- ♦ Komputery korzystające z systemów operacyjnych lub oprogramowania sieciowego firmy Microsoft mogą się komunikować
- ♦ NetBEUI jest w pełni samodostrajającym się protokołem i najlepiej działa w małych segmentach LAN

- ♦ NetBEUI ma minimalne wymagania odnośnie pamięci
- ♦ NetBEUI zapewnia doskonałą ochronę przed błędami transmisji, a także powrót do normalnego stanu w razie ich wystąpienia

Wadą protokołu NetBEUI jest fakt, że nie może być trasowany i niezbyt dobrze działa w sieciach WAN.

Podsumowanie

Protokoły sieciowe są umiejscowione powyżej warstwy łącza danych. Prawdłowo zaprojektowane i skonstruowane są niezależne od architektur sieci LAN (opisanych w części II pt. „Tworzenie sieci LAN”) i zapewniają całościowe zarządzanie transmisjami w domenach sieci LAN.