

Zastosowania sieci komputerowych

- współdzielenie zasobów, np. plików, drukarek
- komunikacja, np. poczta email, telefonia komórkowa, internetowa
- przypadek specjalny: Internet
 - wiele różnych nowych zastosowań: handel, usługi, reklama, gry on-line, *video on demand*, zdalne nauczanie, zdalne głosowanie, itd.
- zastosowania mobilne
 - podobne jak w poprzednich grupach, + dodatkowe, np. nawigacja
- sieci przemysłowe
 - zastosowania przemysłowe
 - * w szczególności sieci bezprzewodowe
- sieci do zastosowań specjalnych
 - specjalne zastosowania komercyjne, np. systemy alarmowe
 - bezprzewodowe sieci czujników
 - zastosowania wojskowe

Klasy wielkości sieci komputerowych

- **sieci osobiste** PAN *Personal Area Network*

Zasięg to często biurko lub pokój, ewentualnie dom i/lub ogród, należące do jednej osoby. Wiele typowych PAN to sieci bezprzewodowe, np. zrealizowane np. z wykorzystaniem technologii Bluetooth. Często podobną sieć zrealizowaną w technologii przewodowej określa się jako LAN.

- **sieci lokalne** LAN *Local Area Network*

Zasięg do 1000 metrów, należące do jednej organizacji, z jednolitym zarządzaniem.

- **sieci metropolitalne** MAN *Metropolitan Area Network*

Zasięg około 10 kilometrów, zwykle nie należą do jednej organizacji i często nie mają wspólnego zarządzania.

- **sieci rozległe** WAN *Wide Area Network*

Zasięg może obejmować cały kraj, kontynent, lub więcej. Często, mianem WAN określa się nie rzeczywiste sieci komputerowe, ale **intersieci**, czyli sieci, które łączą różne sieci.

Systemy rozproszone

Systemem rozproszonym (*distributed system*) określa się sieciowy system komputerowy, który prezentuje użytkownikowi jednolity interfejs funkcji systemu, częściowo ukrywając jego sieciowy charakter.

Na przykład, sieć WWW prezentuje się jako pojedynczy dokument o strukturze drzewa. Możemy przechodzić między podstronami tego dokumentu, nie mając świadomości (albo nawet możliwości stwierdzenia), że poszczególne elementy są obsługiwane przez różne komputery, pod kontrolą różnych systemów operacyjnych, i zlokalizowane w różnych miejscach geograficznych.

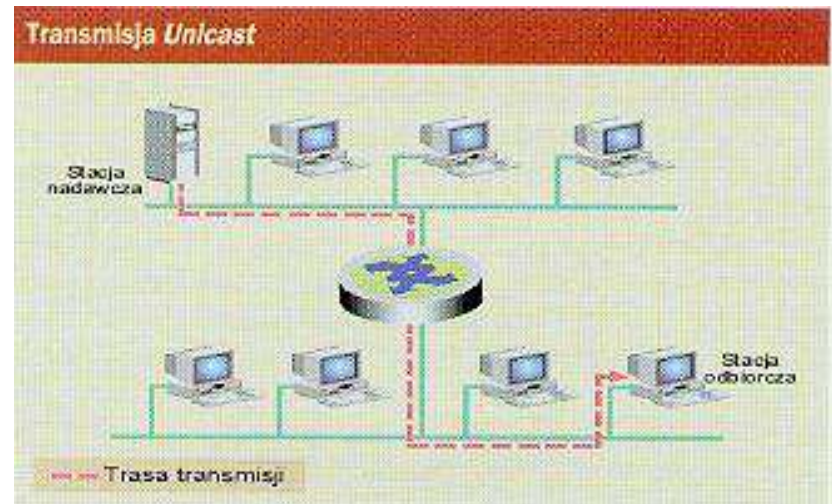
Innym przykładem może być system rezerwacji połączeń PKP lub lotniczych. Znajduje on połączenia realizowane przez różne składy/samoloty, należące do różnych firm, a informacje o tych połączeniach i dostępności miejsc mogą znajdować się w różnych bazach danych.

W odróżnieniu, większość systemów operacyjnych prezentuje interfejs **sieciowego systemu operacyjnego** zmuszającego użytkownika do świadomej nawigacji pomiędzy elementami sieci komputerowej.

Typy transmisji sieciowych

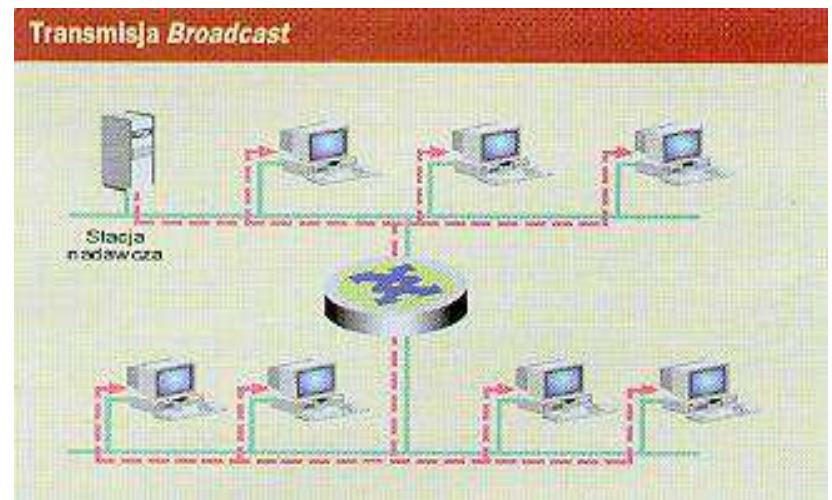
Komunikacja **jednokanałowa**
(*point-to-point, unicast*)

np. połączenie telefoniczne drutowe (aparat końcowy do centrali)



Komunikacja typu **rozgłaszania**
(*broadcast*)

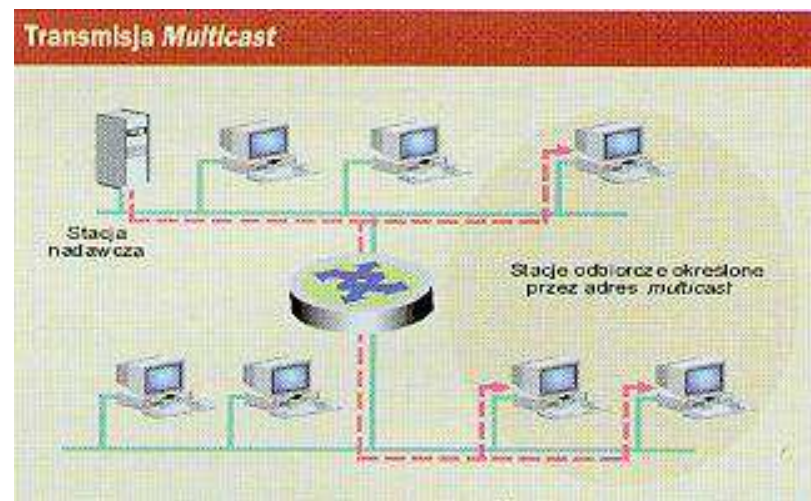
np. transmisja radiowa, satelitarna, itp.



Uwaga: transmisja typu *broadcast* może być adresowana do wybranych odbiorców, pomimo iż może używać pasma dostępnego dla wszystkich.

Typy transmisji sieciowych — multicast

Podobny do *broadcastu* **multicast** oznacza jednoczesne nadawanie do grupy odbiorców. Należy odróżnić transmisję komunikatu *multicast* do określonej grupy (pojedynczy strumień danych) od transmisji jednego komunikatu wiele razy do grupy (zduplikowany strumień danych).



Technologie *multicast* nie są nowym wynalazkiem, ale dotychczas były mało popularne i rozwijane ze względu na komplikacje w niezbędnych technologiach i standardach. (Np. każdy router, nawet w małej sieci domowej lub osiedlowej, powinien być przystosowany do odebrania transmisji *multicast*, zbadania adresu grupy *multicast*, i zadecydowania, czy transmisję należy przekazywać do wnętrza sieci, i konkretnie do których jej części.

W kontekście nabierających popularności szerokopasmowych transmisji wideo te technologie jednak stają się coraz ważniejsze. Zamiast dublować strumień wideo, jak również wysyłać do wszystkich sieci (na świecie) lepiej kierować go do zdefiniowanej grupy.

Typy komunikacji — komunikacja połączeniowa

Wygodnie jest rozważać dwa zasadniczo różne modele komunikacji: połączeniowy i bezpołączeniowy.

Komunikacja połączeniowa jest oparta na zbudowaniu połączenia między stronami, połączenie inicjuje jedna strona, ale utrzymują je obie. Dopóki połączenie istnieje, każda ze stron może nadawać w dowolnym momencie, a druga strona odbiera tę transmisję. Po rozłączeniu, dalsza komunikacja jest niemożliwa do czasu ponownego nawiązania połączenia.

Dobłą analogią komunikacji połączeniowej jest rozmowa przez telefon. Strona inicjująca połączenie musi znać numer telefonu (adres) strony przyjmującej. Strona przyjmująca może nie mieć świadomości numeru dzwoniącego. (Czasami technologia sieci pozwala odbiorcy poznać ten numer, tzw. Caller-ID, ale nie jest to potrzebne do komunikacji.) Jednak po nawiązaniu połączenia żadna ze stron nie musi już pamiętać numeru telefonu drugiej strony.

Typowo w komunikacji połączeniowej strumień danych dochodzi w tej samej postaci w jakiej został nadany (nie ma zamiany kolejności), aczkolwiek przy zawodnym medium jest możliwe przekłamanie, albo utrata części danych.

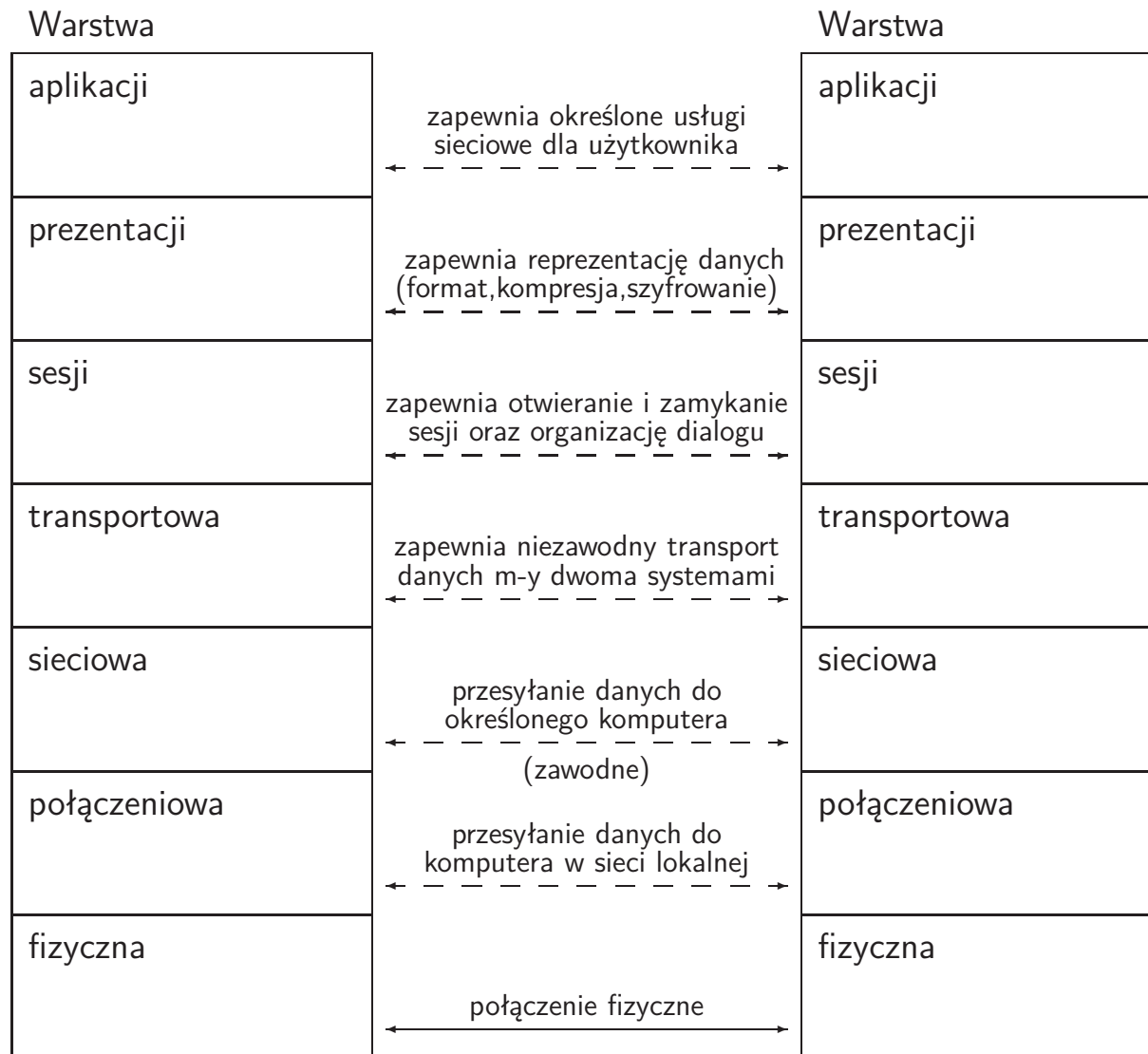
Typy komunikacji — komunikacja bezpołączeniowa

Komunikacja bezpołączeniowa jest oparta na wysyłaniu w pełni adresowanych pakietów danych, z których każdy może być niezależnie doręczony odbiorcy. W każdej chwili możemy wysłać odbiorcy pakiet danych, pod warunkiem, że znamy jego adres.

Podobną analogią komunikacji bezpołączeniowej jest korespondencja listowa. Aby wysłać komuś list trzeba znać jego adres, i żeby odbiorca mógł odpowiedzieć musi on znać adres nadawcy. Przesyłka może być doręczona z adresem nadawcy lub bez tego adresu. (Tradycyjna poczta nie oferuje usługi dostarczenia wraz z listem adresu nadawcy, ale w sieciowej komunikacji bezpołączeniowej takie możliwości zwykle istnieją.)

Typowo w komunikacji bezpołączeniowej możliwa jest zamiana kolejności niektórych pakietów (ich doręczenie w innej kolejności niż były nadane), bo trudno jest kontrolować media komunikacyjne aby tej kolejności przestrzegały. Przekłamanie i gubienie przesyłek jest możliwe podobnie jak w komunikacji połączeniowej.

Warstwowy model sieci ISO-OSI



Każda warstwa realizuje wyłącznie swoje zadania, i nie ma wiedzy ani o zadaniach ani sposobach ich realizacji w warstwach niższych.

Zadania warstw modelu OSI/ISO

- Zadania warstwy fizycznej:
 - zapewnienie dostępu do danych
 - kodowanie strumienia danych
- Zadania warstwy łącza danych
 - dostęp do łącza,
 - formatowanie i transmisja ramek,
 - zapewnienie adresacji.
- Zadania warstwy sieciowej
 - dostarczenie logicznej adresacji
- Zadania warstwy transportowej
 - segmentacje danych w strumień i ponowne ich złożenie w punkcie docelowym
 - zapewnienie niezawodność przesyłu danych
 - zapewnienie parametrów jakości transmisji (QOS - Quality of Service)

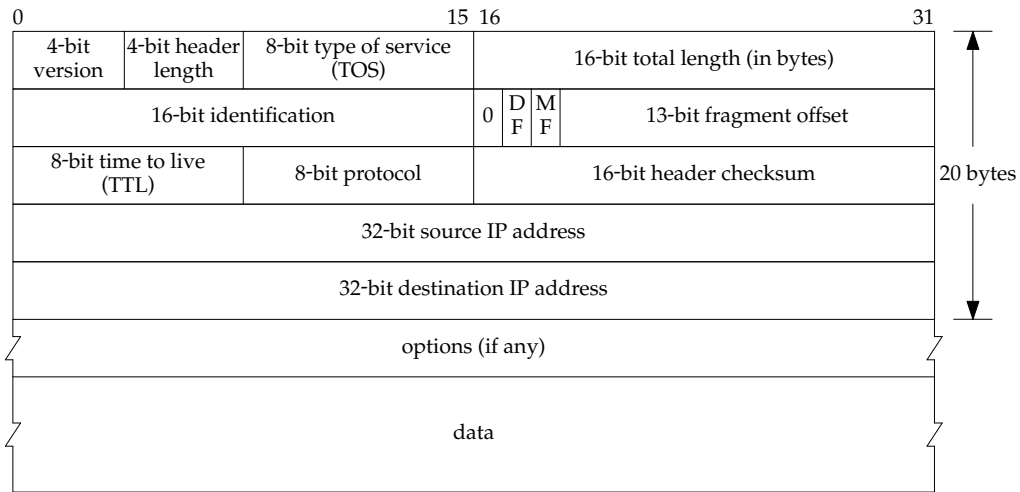
- Zadania warstwy sesji
 - odpowiedzialna za sesje między dwoma procesami na różnych komputerach
 - implementowana jest przez system operacyjny
 - odpowiada za synchronizację danych między komputerami
 - określenie czy stacje mają uprawnienia do komunikacji przez sieć
- Zadania warstwy prezentacji
 - odpowiedzialna za reprezentację danych
 - implementowana przez system operacyjny
 - konwersja między standardami kodowania znaków
- Zadania warstwy aplikacji
 - najbliższej użytkownika
 - przeglądarka WWW, klient poczty elektronicznej, aplikacje konferencyjne, FTP

Model ISO a protokoły internetowe

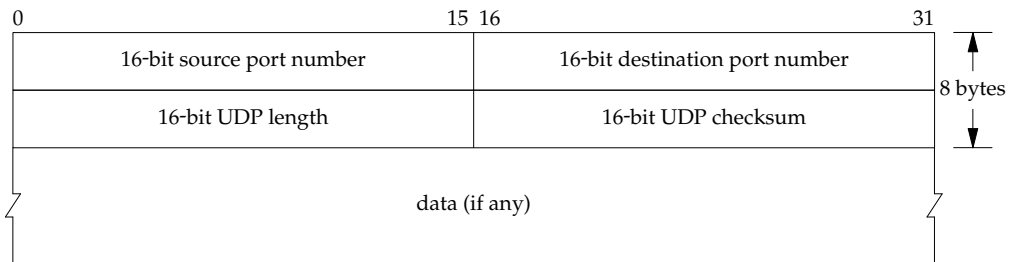
Warstwy modelu ISO	Protokoły internetowe	Funkcja
aplikacji	aplikacji	
prezentacji		
sesji		interface gniazdek
transportowa	TCP, UDP, ...	dostarczanie danych w trybie połączeniowym lub bezpołączeniowym pod określony adres (komputer+port)
sieciowa	IP, ICMP IPv6, ICMPv6	znajdowanie ścieżek sieciowych, przekazywanie pakietów do właściwego adresu lokalnego, funkcje kontrolne
połączeniowa	systemowy driver	
fizyczna	karta sieciowa inny sprzęt sieciowy	

Kapsułkowanie: nagłówki pakietów TCP, UDP i IP

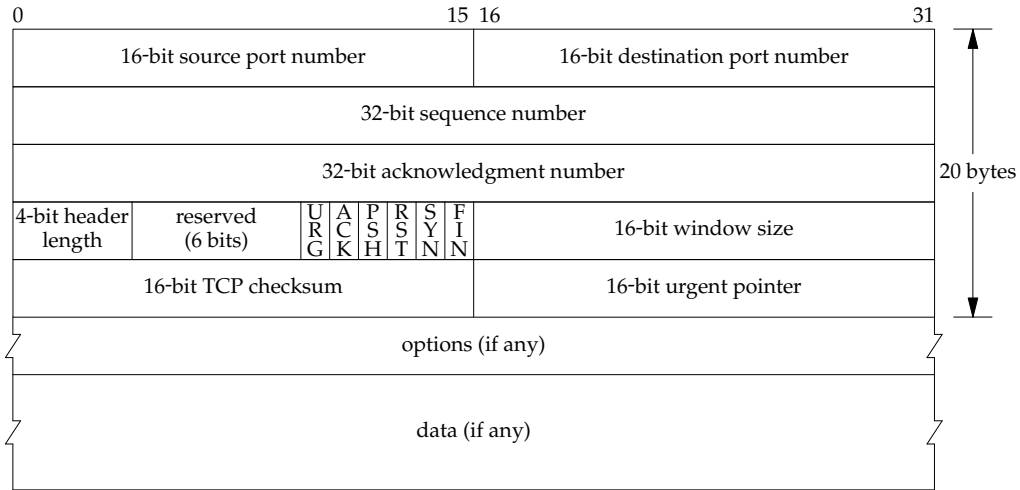
IP Header



UDP Header



TCP Header



Adresowanie w Internecie

32-bitowe (4-oktetowe) adresy IP wersji 4 (IPv4), stosowane od 1.1.1983

- klasa A: pierwszy oktet jest adresem sieci (a pozostałe trzy adresem komputera), z czego pierwszy bit jest zerem; możliwe 128 takich sieci (przedział 0-127) i 16 milionów komputerów w każdej
- klasa B: pierwsze dwa oktety są adresem sieci (a dwa adresem komputera), z czego pierwsze dwa bity są 10; możliwe 16 tysięcy adresów sieci (przedział 128-191) po 65 tysięcy komputerów
- klasa C: pierwsze trzy oktety są adresem sieci (a jeden adresem komputera), z czego pierwsze trzy bity są 110; możliwe dwa miliony takich sieci (przedział 192-223) i 256 komputerów w każdej
- klasa D: cały adres jest jednym adresem, ale cztery pierwsze bity muszą być 1110; adresy te stosuje się do komunikacji multicast
- adresy *broadcast*: same jedyńki lub same zera w adresie komputera

Te system sprawdzał się dobrze w latach 1980-tych w sieci zbudowanej pod nazwą ARPANET¹, a potem stopniowo coraz częściej nazywanej Internetem.

¹ARPA (obecnie DARPA — *Defense Advanced Research Project Agency*) to amerykańska agencja rozdzielająca granty badawcze dla potrzeb Departamentu Obrony U.S.A., która pierwotnie sponsorowała projekt i budowę tej sieci.

Konsekwencje stosowania klasowego systemu adresowania

Zauważmy, że zdefiniowane klasy adresowe nie mają żadnego związku z geografią ani topologią sieci. Sąsiadujące adresy z danej puli adresowej mogą być przyznane jednostkom znajdującym się w różnych częściach świata, jak również w różnych fizycznych sieciach. Powoduje to, że globalne routery sieciowe, kierujące ruch pomiędzy różnymi sieciami, rozlokowanymi w różnych częściach świata, i różnie połączonymi, są niezwykle obciążone, ponieważ muszą rozróżniać wszystkie istniejące adresy klasowe.

Ponadto, przydzielanie adresów IP według schematu czterech klas jest nieefektywne.

Pomimo iż istnieje ponad 4 miliardy liczb 32-bitowych, muszą one być przydzielane w dość rozrzutny sposób. Jeśli duża organizacja potrzebuje przestrzeni adresowej na więcej niż 65 tysięcy adresów (być może z uwzględnieniem przyszłego rozwoju organizacji), to nie może używać sieci adresowej klasy B, tylko musi używać klasy A. Jednak takich sieci jest zaledwie 128, i część z nich została rozdysponowana już w początkowych latach Internetu. Dodatkowo, spośród wszystkich chętnych na sieci adresowe klasy A, mało kto rzeczywiście był w stanie wykorzystać całą 16-milionową pulę adresową takiej sieci. Podobne zjawisko występuje na granicy liczności klas B i C.

Wyczerpywanie adresów IP

W pewnym momencie oszacowano, że pomimo wyczerpania przestrzeni adresowej IPv4, wykorzystywanych było jedynie 14% wszystkich adresów.

Z drugiej strony, rewolucyjny rozwój Internetu od początku lat 1990-tych przekroczył wszelkie najśmielsze wyobrażenia twórców opartej na klasach struktury adresowania IP.

Kombinacja tych zjawisk spowodowała, że już w połowie lat 1990-tych zaczęło brakować adresów IP. Gdy stało się jasne, że żadne próby oszczędnego gospodarowania tymi adresami, w tym odbierania niektórym organizacjom większych bloków adresowych, i przydzielania im mniejszych, nie są w stanie rozwiązać tego problemu, podjęto prace nad wypracowaniem rozwiązania.

Rozwiązanie problemów klasowego systemu adresowania IPv4

Poprawnym i docelowym rozwiązaniem problemu jest wprowadzenie nowego — znacznie pojemniejszego — systemu adresowania, docelowo nazwanego IPv6. Ponieważ jednak szybkie i skuteczne wprowadzenie tego systemu byłoby niezwykle kosztowne i uciążliwe (wymagałoby wymiany całego istniejącego oprogramowania, jak również sprzętu sieciowego operującego adresami internetowymi) prace nad nim przeciągały się.

W związku z tym wdrażano różne rozwiązania pośrednie, rozszerzające możliwości wykorzystania starego systemu adresowania, i umożliwiające pracę i dalszy rozwój technologii w warunkach braku adresów IPv4.

Spośród wielu takich inicjatyw przedstawione tu będą skutecznie wdrożone, i stosowane obecnie: adresowanie bezklasowe CIDR, oraz translacja adresów NAT.

Adresowanie bezklasowe CIDR

Wprowadzony w 1993 system adresowania bezklasowego CIDR (*classless inter-domain routing*) odrzuca sztywny podział adresu na adres sieci i komputera wyznaczony przez klasę. Obecnie dowolny adres IPv4 składa się z *prefiksu* (pierwszej części) dowolnej długości, stanowiącego adres sieci, i reszty, stanowiącej adres komputera.

W systemie CIDR nie da się poznać części adresu stanowiącej adres sieci po klasie adresu. Zatem dla wskazania długości prefiksu stosuje się notację $x.y.z.t/p$. Na przykład, 156.17.9.0/25 oznacza adres sieci, w którym 25 bitów stanowi adres sieci, a pozostałych 7 adres komputera. Oznacza to, że może być 128 adresów w tej sieci, z których pierwszy (same zera w części adresu komputera) jest adresem sieci, a ostatni (same jedynki w części adresu komputera) jest adresem rozgłaszania (*broadcast*), co pozostawia 126 rzeczywistych adresów.

Adresowanie bezklasowe pozwala właścicielom bloków adresów efektywniej nimi gospodarować, co powoduje **mniejsze marnowanie adresów** w poszczególnych blokach. Jednak co równie ważne, pozwala ono administratorom sieci na **definiowanie zagregowanych ścieżek routingu**. Na przykład, pomimo iż istnieje wiele małych sieci o adresach zaczynających się na 156.17.x.x, to dla globalnego routera mogą one być reprezentowane jedną ścieżką 156.17/16.

Sieci prywatne

W schemacie adresowania IPv4 trzy zakresy adresów zostały zarezerwowane jako „prywatne”. Przeznaczone były do wykorzystania w sieciach LAN: domowych, biurowych, i firmowych, nie połączonych z Internetem. Korzystanie z tych adresów nie wymaga żadnych zezwoleń, uzgodnień, ani rejestracji:

- 10.0.0.0 – 10.255.255.255
- 172.16.0.0 – 172.31.255.255
- 192.168.0.0 – 192.168.255.255

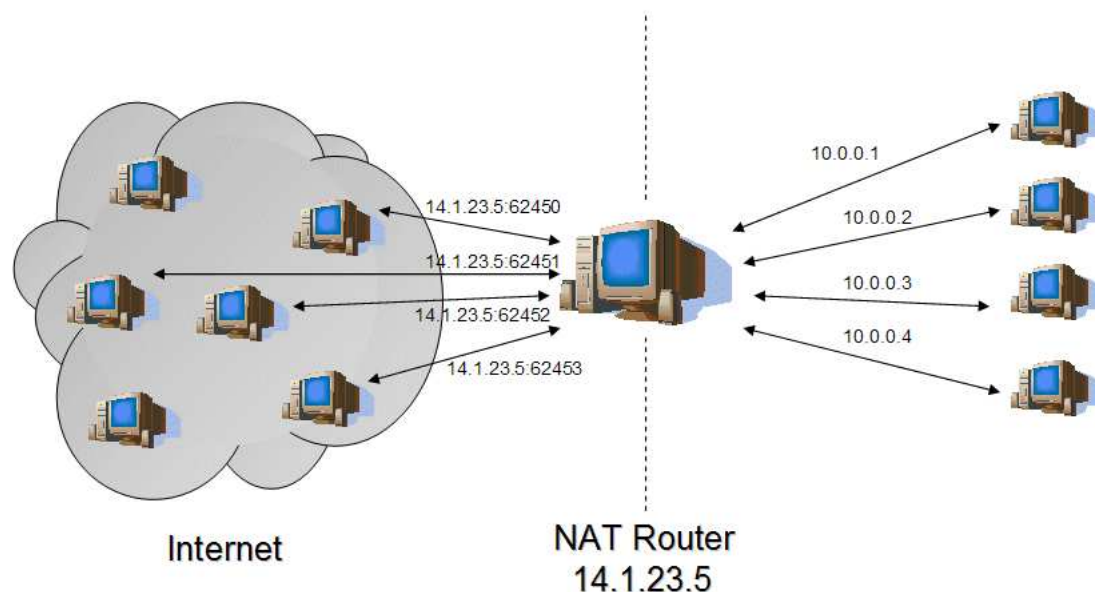
Oznacza to, że w takich sieciach można stosować oprogramowanie przeznaczone do Internetu, takie same urządzenia i konfiguracje, ale sieć nie będzie komunikować się z Internetem, ponieważ te adresy w Internecie nie są powiązane z jakąkolwiek siecią. Przykład: duży bank posiada sieć komputerową, która wykorzystuje oprogramowanie TCP/IP, ale nie jest planowane połączenie tej sieci z Internetem.

Zapewnia to z jednej strony bezpieczeństwo infrastruktury firmy przed możliwymi atakami z Internetu, ale również brak zagrożenia możliwym wyciekiem danych przy okazji nieostrożnych poczynań pracowników.

Co znacznie ważniejsze, z drugiej strony prywatne bloki adresowe okazały się skutecznym remedium problemu wyczerpania zwykłych adresów.

Translacja adresów NAT

Wprowadzony w 1996 NAT (*Network Address Translation*) pozwala sieciom LAN wykorzystującym adresy prywatne na łączenie się z Internetem za pomocą routerów dokonujących zamiany adresów prywatnych w LAN na publiczny adres internetowy routera obsługującego sieć LAN, plus unikalnego numeru portu, identyfikującego połączenie.



Od strony Internetu wygląda to tak, jakby cały ruch z takiej sieci prywatnej pochodził z samego routera. Natomiast router zamienia adresy wychodzące z sieci lokalnej na swój własny adres, a adresy przychodzących odpowiedzi na właściwe adresy prywatne sieci LAN, na podstawie numeru portu.

Praktyczne aspekty stosowania NAT

NAT może wprowadzić na własny użytek małe biuro, gospodarstwo domowe, ale także całkiem duża firma, lub dostarczyciel usługi Internetu dla osiedla albo nawet całego miasta. NAT może być realizowany przez komputer zapewniający łączność sieci LAN z Internetem, może być zrealizowany przez router sprzętowy, a nawet aplikację na telefon komórkowy, tworzący prywatną sieć WiFi i łączący ją z Internetem przez sieć komórkową GSM.

Wprowadzenie NAT odsunęło na długi czas kryzys wynikający z wyczerpania adresów IP, ponieważ większość użytkowników Internetu na świecie korzysta z sieci za pośrednictwem dostawcy usługi Internetu, którzy mogą świadczyć tę usługę z wykorzystaniem NAT i adresów prywatnych w obsługiwanej sieci.

Praktyczne aspekty stosowania NAT (2)

Należy pamiętać, że NAT nie zapewnia pełnej łączności sieci LAN z Internetem, **a jedynie połączenia wychodzące**. Nie jest możliwe umieszczenie pełnofunkcyjnego serwera internetowego w sieci LAN z NAT, nie da się przypisać normalnych adresów IP nabierającym popularności urządzeniom IoT, oraz nie jest to odpowiednie rozwiązanie dla pewnych typów komunikacji, jak np. telefonia internetowa VOIP.

W efekcie użytkownicy Internetu w sieciach z NAT są w gorszym położeniu, ponieważ nie mają dostępu do pełnej funkcjonalności sieci. Oczywiście, dla wielu z nich nie ma to praktycznego znaczenia dopóty, dopóki są w stanie zamówić pizzę, kupić bilety, wykonać przelew, sprawdzić pogodę, mailować, porozmawiać przez Messengera, ...

Praktyczne aspekty stosowania NAT (3)

Poza oczywistą konsekwencją, w postaci braku w pełni routowalnego adresu IP, masowe użycie NAT wprowadza wiele dodatkowych komplikacji technicznych, np.:

- NAT łamie transparentność sieci, która ułatwia wprowadzanie nowych technologii i jest ogólnie uważana za zaletę,
- w szczególności, routery sieciowe obsługujące NAT muszą przechowywać stan wszystkich połączeń, co pochłania zasoby i wprowadza pewne subtelne problemy techniczne (np. jak długo należy przechowywać taką informację o stanie, itp.),
- użytkownik sieci LAN/NAT, który nadużywa dobrych obyczajów, albo łamie prawo, jest w Internecie widoczny z tego samego adresu IP co wszyscy użytkownicy tej sieci, i tym samym naraża ich na: ostracyzm, wykluczenie, albo śledztwo,
- administrator sieci LAN/NAT może być zmuszony do rejestrowania wszystkich indywidualnych połączeń.

Nowy standard IPv6

System adresowania IPv4 jest wykorzystywany wraz z podstawowym protokołem komunikacyjnym Internetu IP od 1 stycznia 1983. Pozwolił on na poprawną pracę sieci od początku, kiedy istniało zaledwie kilkadziesiąt węzłów, aż do masowego wzrostu w latach 1990-tych, kiedy adresów IP zaczęło brakować.

Ponadto, dowolność w przydzielaniu pojedynczych adresów sieci powodowała przeciążenie do granic fizycznych możliwości globalnych routerów internetowych. Wynikało to z faktu przydzielania różnym rozproszonym geograficznie jednostkom dowolnych adresów sieci, które musiały być routowane na poziomie globalnym.

Jednym z projektów zmiany systemu adresowania była koncepcja IPng (IP nowej generacji), która później rozwinęła się w standard IPv6. Niestety, ta koncepcja powstawała wolno i nie nadążała za potrzebami Internetu. Co gorsza, zatwierdzony ostatecznie w 1998 standard IPv6 nie zapewnia interoperacyjności z IPv4. W czasie gdy powstawał nie było jeszcze oczywiste, że w Internecie takiej drastycznej zmiany nie da się przeprowadzić.

Jednym z elementów standardu IPv6 jest nowy system adresowania.

System adresowania IPv6

Adresy IPv6 mają 128 bitów długości. Teoretycznie daje to ponad 10^{38} różnych adresów, co jest liczbą zdecydowanie przekraczającą wszelkie przyszłe potrzeby, i to w skali intergalaktycznej. Jednak celem tego systemu nie było wygenerowanie dużej liczby używalnych adresów. Raczej, przestrzeń adresowa ma zapewniać duży nadmiar, pozwalając zarówno na uproszczenie przetwarzania adresu przez zewnętrzne routery, jak i przydzielanie pewnym sekcjom adresu określonych funkcji.

Stosuje się heksadecymalną notację tych adresów, oddzielającą dwukropkami co drugi oktet, gdzie zblokowane dwa oktety bywają nazywane **hekstetami**:

np. 2001:0a48:000f:0146:0000:0000:0000:0001

Stosuje się skróconą formę notacji adresów. W każdej sekcji można pominąć początkowe zera, a jedną najdłuższą sekwencję bloków samych zer — nawet obejmującą kilka hekstetów — można pominąć i zastąpić podwójnym dwukropkiem:

np. 2001:a48:f:146::1

Ponadto stosuje się notację dla określenia długości prefiksu sieci, analogiczną do notacji CIDR w IPv4:

np. sieć 2001:a48:f:146::/64

obejmuje przedział adresów od 2001:a48:f:146:0:0:0:0 do 2001:a48:f:146:ffff:ffff:ffff:ffff

System adresowania IPv6 (2)

Istnieje szereg typów adresów, które można rozróżnić po kilku początkowych bitach:

globalny adres *unicast* i *anycast*

globalnie rozpoznawalny i dostępny adres indywidualnego urządzenia

przykład: 2001:0db8:3c4d:0015:0000:0000:1a2f:1a2b

Pierwsza część adresu stanowi tzw. prefiks sieciowy (*network prefix*) składający się z prefiksu kierunkowego (*routing prefix*) i identyfikatora podsieci (*subnet ID*), natomiast druga część stanowi identyfikator interfejsu (*interface ID*) przypisany konkretnemu urządzeniu i jego konkretnemu interfejsowi (karcie sieciowej).

Ideą ogromnego nadmiaru długości adresu IPv6 jest żeby prefiks kierunkowy mógł służyć do lokalizowania sieci, upraszczając rozwiązywanie problemu **routingu**.

Ten typ adresu jest rozpoznawalny po prefiksie 2xxx: lub 3xxx: (pierwsze trzy bity 002).

Nowym rodzajem adresu w standardzie IPv6 jest adres *anycast*, który jest wariantem adresu ***unicast***, z tym samym prefiksem. Te adresy mogą być nadawane wielu urządzeniom jednocześnie, a doręczenie pakietu do tego adresu polega na jego doręczeniu do jednego wybranego (najbliższego?) urządzenia współdzielącego adres.

unikalny adres lokalny

adres urządzenia rozpoznawalny i dostępny lokalnie, ale nie globalnie, a więc jest odpowiednikiem adresu prywatnego standardu IPv4

rozpoznawalny po prefiksie FCxx: lub FDxx: (pierwsze siedem bitów 1111110)

adres lokalnego połączenia

adres do szybkiego automatycznego nadawania indywidualnym urządzeniom, nierozpoznawany globalnie ani lokalnie, odpowiednik adresu 169.254.x.x standardu IPv4

rozpoznawalny po prefiksach: FE8x: FE9x: FEAx: FEBx: (pierwsze 10 bitów: 1111111010)

adres *multicast*

adres do rozsyłania grupowego

rozpoznawalny po prefiksie FFxx: (pierwsze 8 bitów: 11111111)

W standardzie IPv6 nie ma adresów *broadcast*, ich rolę w całości muszą przejąć adresy *multicast*.

Routing (trasowanie?)

Routing jest czynnością określania dokąd należy wysłać pakiet sieciowy mający zawarty w nim adres odbiorcy. Wymaga to lokalizacji sieci komputerowej na podstawie jej adresu. Lokalizację rozumiemy tu w sensie połączeń, to znaczy znalezienie ścieżki połączeń sieciowych (szeregu połączonych routerów), prowadzącej do lokalizowanej sieci.

Realizacja tej czynności opiera się na powiązaniu prefiksu adresu sieci z fizyczną lokalizacją sieci. Informacje wymieniane między komputerami znajdującymi ścieżki (routerami) pozwalają im na określanie tych ścieżek przez abstrakcję.

Algorytm routera: (1) jeśli adres jest w mojej podsieci to wysyłam pakiet do docelowego odbiorcy; (2) jeśli tak się składa, że mam w pamięci ścieżkę do podsieci odbiorcy pakietu, wraz z bramą lokalną (routerem) stanowiącą początek tej ścieżki, to przekazuję pakiet temu routerowi, oraz (3) w przeciwnym wypadku nie wiem jak doręczyć dany pakiet, więc muszę go skasować!!

Szczególnym przypadkiem warunku (2) jest sytuacja, kiedy komputer posiada tzw. ścieżkę domyślną, określającą bramę obsługującą wszystkie adresy, do których nie jest pamiętana indywidualna ścieżka. Jeśli komputer posiada zdefiniowaną taką ścieżkę to jest w stanie doręczyć wszystkie pakiety.

Routing — tablica ścieżek

Decyzja wyboru ścieżki sieciowej, do której należy wysłać dany pakiet sieciowy, jest podejmowana na podstawie docelowego adresu IP pakietu, i jej wynikiem jest wybór komputera w (jednej z) sieci lokalnej(ych), do której(ych) dany komputer jest podłączony.

Dodatkowym parametrem każdej ścieżki jest jej **metryka** określająca łatwość przesłania pakietu przez tą ścieżkę. Mogą istnieć różne ścieżki do tej samej sieci docelowej, ale z różnymi metrykami. W najprostszym przypadku metryka może określać liczbę segmentów sieci które pakiet będzie musiał pokonać na drodze do sieci docelowej. Pokonanie każdego segmentu wiąże się z przetwarzaniem pakietu w jakimś urządzeniu, a więc im większa ta liczba segmentów tym dłużej będzie trwało przesyłanie do sieci docelowej.

Routing jest czynnością wykonywaną w ramach protokołu IP (warstwy sieciowej, w nomenklaturze ISO).

Routing — tablica ścieżek

Routing realizowany jest w sposób zasadniczo dość prosty: system operacyjny posiada tablicę ścieżek sieciowych, określając powiązania docelowych adresów IP komputerów i całych sieci, z bramami, czyli adresami IP komputerów w sieci lokalnej, czyli takich, do których przesłanie jest bezpośrednie.

Może istnieć wiele ścieżek w tej tablicy, i wybierana jest zawsze najlepiej dopasowana, to znaczy najbardziej szczegółowa ścieżka zgodna z danym adresem docelowym. W braku takiej ścieżki używana jest specjalna ścieżka domyślna, a gdy jej nie ma, pakietu nie da się wysłać do miejsca przeznaczenia, i routing kończy się niepowodzeniem. Pakiet zostaje zwyczajnie skasowany, natomiast do nadawcy może zostać wysłany komunikat informujący go o błędzie w jego tablicy ścieżek.

Routing — tworzenie ścieżek sieciowych

Skąd komputery biorą tablice ścieżek? Skąd komputer właśnie włączony do sieci może znać ścieżkę komputera z określonym adresem położonym np. w Australii? Zwłaszcza, że sieci komputerowe i połączenia między nimi zmieniają się dynamicznie, powstają nowe połączenia, znikają istniejące, zmieniają się metryki pewnych połączeń, występują awarie, itp.

Odpowiedzią na te pytania jest cały szereg dość złożonych procesów i technologii.

W niewielkich sieciach stosuje się **routing statyczny** polegający na ręcznym kopiowaniu informacji o zmianach ścieżek sieciowych do wszystkich komputerów sieci. W większych sieciach jest to niemożliwe, i stosuje się **routing dynamiczny**. Polega on na propagowaniu informacji o zmianach ścieżek sieciowych przez komputery, i automatycznej aktualizacji tablicy ścieżek. Służą do tego specjalne protokoły komunikacji, pozwalające określić kto komu może przesyłać informacje o których ścieżkach, i od kogo można przyjmować wiążące informacje o zmianach.

Na poziomie światowym istnieją routery rdzenia sieci szkieletowej (*backbone core routers*), które łączą pomiędzy sobą główne routery sieci składowych traktowanych jako systemy autonomiczne. Routery rdzeniowe przekazują pakiety od jednego systemu do drugiego, a routingiem w ramach systemu zajmują się ich routery wewnętrzne.

Adresy symboliczne

Dla wygody wprowadzono dualną przestrzeń adresów — adresy symboliczne. Są one zorganizowane w hierarchiczną strukturę tzw. domen adresowych. Struktura ta nie ma ograniczeń; domeny położone „niżej” w hierarchii są własnością różnych organizacji, które same dostarczają informacji o domenach w nich zawartych.

Np.: komputer sequoia.kcir.pwr.edu.pl (IP:156.17.9.3) należy do domeny kcir.pwr.edu.pl, która jest własnością Katedry Cybernetyki i Robotyki Politechniki Wrocławskiej, i która otrzymała prawa do tej domeny od Politechniki Wrocławskiej, właściciela domeny pwr.edu.pl.

Przestrzeń adresów symbolicznych służy tylko wygodzie ludzkiej pamięci, i istnieje odwzorowanie adresów symbolicznych na adresy numeryczne. Do realizacji tego odwzorowania służy specjalny system DNS (*domain name system*) składający się z sieci serwerów wymieniających informacje o tych odwzorowaniach i serwujących te informacje na życzenie.

Translacja nazw symbolicznych — system DNS

- DNS (*Domain Name System*) — hierarchiczny, rozproszony system nazw symbolicznych w Internecie
- oparty na oddelegowaniu administracji domenami różnym instytucjom, korzystającym z własnych serwerów DNS, automatycznie wymieniającym między sobą informacje o administrowanych przez siebie domenach
domena — poddrzewo hierarchicznego drzewa nazw
- własności: nadmiarowość, replikacja, buforowanie, duża niezawodność i tolerancja błędów, optymalizacja procesu uzyskiwania odpowiedzi w warunkach rzadkich zmian
- serwer DNS — program, którego zadaniem jest podawanie translacji adresu określonego w zapytaniu klienta, i komunikujący się z innymi serwerami DNS, w celu jej znalezienia
- serwery DNS mogą posiadać redundancję — dla danej domeny można wprowadzić oprócz serwera głównego (**primary**), równoważne serwery dodatkowe (**secondary**)

Serwery systemu DNS

- serwer DNS domyślnie jest **rekurencyjny**; w sytuacji gdy nie zna odpowiedzi na otrzymane zapytanie, sam kontaktuje się z innymi serwerami aby ją uzyskać, i udzielić pytającemu klientowi

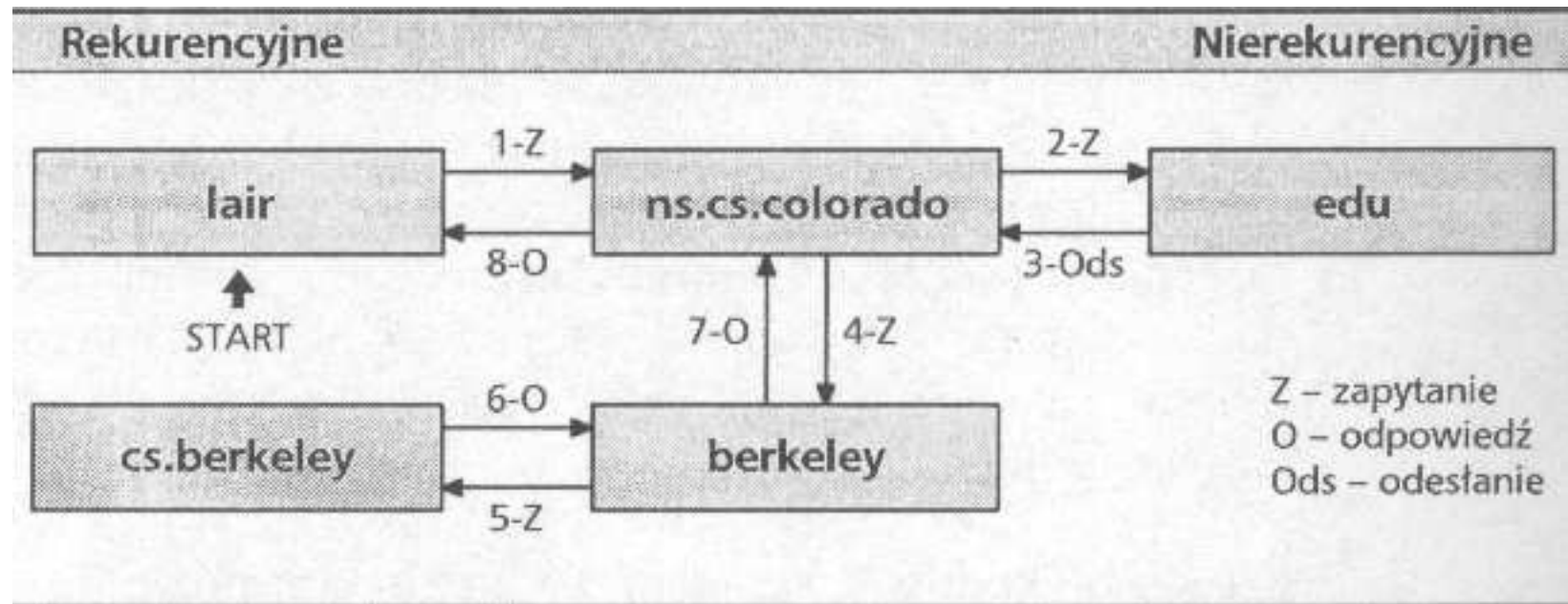
rekurencyjny serwer jest właściwym rozwiązaniem dla sieci lokalnej, ponieważ pozwala klientom zawsze uzyskiwać odpowiedzi na swoje pytania, a ponadto może przechowywać uzyskane odpowiedzi, i udzielać ich potem kolejnym klientom bez ponownego odpytywania rekurencyjnego

- serwer DNS może być również **nierekurencyjny**; w przypadku nieznajomości odpowiedzi serwer taki nie pyta się innych serwerów, tylko odpowiada tzw. odsyłaczem (ang. *referral*), podającym adres innego, bardziej właściwego dla danej domeny serwera DNS

serwery DNS wyższego poziomu w hierarchii Internetu (np. serwery główne takich domen jak .com albo .pl) są z zasady nierekurencyjne, więc tym bardziej nie przechowują informacji, które ich nie dotyczą

Dla domen pośrednich pomiędzy siecią lokalną a domeną główną Internetu musimy wybrać pomiędzy pracą rekurencyjną a nierekurencyjną serwera DNS. Jednak nierekurencyjny serwer nie może obsługiwać normalnych klientów, nieprzygotowanych na otrzymanie na swoje zapytanie odpowiedzi w postaci odsyłacza.

Przykład sekwencji odwołań do serwerów DNS dla zapytania o nazwę `mammoth.cs.berkeley.edu` wykonanego na komputerze `lair.cs.colorado.edu`:



Serwery DNS

primary — jest tylko jeden taki serwer dla strefy (ang. *zone*); strefa jest częścią domeny administrowaną przez serwer

secondary — takich może być dla danej strefy wiele, automatycznie aktualizują one swoje dane i ich odpowiedź jest równoważna odpowiedzi serwera *primary*

caching-only — nie jest właściwym źródłem informacji o żadnej strefie, nie posiada własnych informacji tylko realizuje funkcję rekurencyjnego odpytywania innych serwerów i przechowuje informacje przez dozwolony okres; można go uważać za rodzaj aktywnego klienta; jeśli nie chcemy zakładać w danym systemie serwera DNS, ale chcemy zaoszczędzić na ruchu sieciowym do zewnętrznych serwerów DNS, to możemy założyć właśnie serwer *caching-only*

Jeden serwer (uruchomiona instancja programu) może być serwerem primary dla jednej strefy (lub kilku), i serwerem secondary dla grupy innych stref, albo może być czystym serwerem caching-only.

System DNS zaprojektowany w połowie lat 1980-tych jest przykładem rozproszonego systemu o dużej niezawodności, poprawnie zabezpieczającego serwis translacji adresów symbolicznych w warunkach rozległego, niedeterministycznego Internetu. System poprawnie przetrwał rewolucyjny rozwój Internetu od początku lat 1990-tych, kiedy całkowicie zmieniły się technologie, prędkość, niezawodność, i wymagania stawiane Internetowi.

Jednej cechy, której nie przewidziano w tym systemie jest bezpieczeństwo. System powstał w czasie, gdy nie istniały ani obecne zagrożenia ani wymagania dotyczące bezpieczeństwa. Rozproszoność i redundancja tego systemu powoduje łatwość przeprowadzania ataków takich jak podszywanie się.

Literatura

Computer Networking: Principles, Protocols and Practice

<http://cnp3book.info.ucl.ac.be/1st/html/>

Overview of DNS Protocol

<https://www.plixer.com/blog/overview-of-dns-protocol-part-1-of-3/>