

第5章 IPv6 与 IPv4

网络中的通信是建立在协议之上的，如果没有统一的通信协议，则不可能实现各网络的互联及彼此相互通信。而对于连接 Internet 网络的每台计算机，则都必须使用 IP 地址来设置一个相对于所有计算机的唯一标识。例如，为方便远距离人与人之间能够通信，则每个家庭相对全国来说，都有唯一的邮政编码和地址名称。

目前，网络中正使用着 IPv4 协议，但由于该地址协议数量有限及网络用户不断的增加，IPv4 已经无法满足现在的需求。为此第二代版本 IPv6 应运而生，由于它具有许多优点，被誉为下一代 Internet 的标准协议。

本章学习目标：

- 了解 IPv4
- 了解 IP 协议的功能
- 掌握子网掩码的计算
- 了解 IPv6 及基础知识
- 掌握 IPv6 路由协议和过渡技术

5.1 IPv4 简介

现行的 IPv4 自 1981 年 RFC 791 标准发布以来并没有多大的改变。事实证明，IPv4 具有相当强盛的生命力，易于实现且互操作性良好，经受住了从早期小规模互联网络扩展到如今全球范围 Internet 应用的考验。所有这一切都应归功于 IPv4 最初的优良设计。但是，还是有一些发展是设计之初未曾预料到的。

近年来 Internet 呈指数级的飞速发展，导致 IPv4 地址空间几近耗竭。IP 地址变得越来越珍稀，迫使许多企业不得不使用 NAT 将多个内部地址映射成一个公共 IP 地址。地址转换技术虽然在一定程度上缓解了公共 IP 地址匮乏的压力，但它不支持某些网络层安全协议以及难免在地址映射中出现种种错误，这又造成了一些新的问题。而且，靠 NAT 并不可能从根本上解决 IP 地址匮乏问题，随着连网设备的急剧增加，IPv4 公共地址总有一天会完全耗尽。

Internet 主干网路由器维护大型路由表能力的增强。目前的 IPv4 路由基本结构是平面路由机制和层次路由机制的混合，Internet 核心主干网路由器可维护 85000 条以上的路由表项。

地址配置趋向于要求更简单化。目前绝大多数 IPv4 地址配置需要手工操作或使用 DHCP（动态主机配置协议）地址配置协议完成。随着越来越多的计算机和相关设备使用 IP 地址，必然要求提高地址配置的自动化程度，使之更简单化，且其他配置设置能不依赖于 DHCP 协议的管理。

IP 层安全需求的增长。在 Internet 这样的公共媒体上进行专用数据通信一般都要求

加密服务，以此保证数据在传输过程中不会泄露或遭窃取。虽然目前有 IPSec 协议可以提供对 IPv4 数据包的安全保护，但由于该协议只是个可选标准，企业使用各自私有安全解决方案的情况还是相当普遍。

更好的实时 QoS 支持的需求。IPv4 的 QoS 标准，在实时传输支持上依赖于 IPv4 的服务类型字段（TOS）和使用 UDP 或 TCP 端口进行身份认证。但 IPv4 的 TOS 字段功能有限，而同时可能造成实时传输超时的因素又太多。此外，如果 IPv4 数据包加密的话，就无法使用 TCP/UDP 端口进行身份认证。

为了解决上述问题，Internet 工程任务组（IETF）开发了 IPv6。这一新版本，也曾被称为下一代 IP，综合了多个对 IPv4 进行升级的提案。在设计上，IPv6 力图避免增加太多的新特性，从而尽可能地减少对现有的高层和低层协议的冲击。

5.2 IP 协议

IP 协议是 Internet 中的交通规则，接入 Internet 中的每台计算机及处于十字路口的路由器都必须熟知和遵守该交通规则。IP 数据包则是按该交通规则在 Internet 中行驶的车辆，发送数据的主机需要按 IP 协议装载数据，路由器需要按 IP 协议指挥交通，接收数据的主机需要按 IP 协议拆卸数据。IP 数据包携带着地址、满载着数据从发送数据的端用户计算机出发，在沿途各个路由器的指挥下，顺利到达目的端用户的计算机。

5.2.1 IP 协议的特点

IP 协议主机负责为计算机之间传输数据报寻址，并管理这些数据报的分片过程。它对投递的数据报格式有规范、精确的定义。与此同时，IP 协议还负责数据报的路由，决定数据报的发送地址，以及在路由出现问题时更换路由。总的来说，IP 协议具有以下几个特点：

- ❑ 不可靠的数据投递服务。IP 协议本身没有能力证实发送的数据报是否能被正确接收。数据报可能在遇到延迟、路由错误、数据报分片和重组过程中受到损坏，但 IP 不检测这些错误。在发生错误时，也没有机制保证一定可以通知发送方和接收方。
- ❑ 面向无连接的传输服务。IP 协议不管数据报沿途经过哪些结点，甚至也不管数据报起始于哪台计算机，终止于哪台计算机。数据报从源始结点到目的结点可能经过不同的传输路径，而且这些数据报在传输的过程中有可能丢失。有可能正确到达。
- ❑ 尽最大努力投递服务。IP 协议并不随意的丢失数据报，只有当系统的资源用尽、接收数据错误或网络出现故障等状态下，才不得不丢弃报文。

5.2.2 IPv4 地址

IP 地址的长度为 32 位（4B）的无符号的二进制数，它通常采用点分十进制数表示

方法,即每个地址被表示为四个以小数点隔开的十进制整数,每个整数对应一个字节,如 165.112.68.110 就是一个合法的 IP 地址。32 位的 IP 地址由网络号和主机号两部分构成。其中,网络号就是网络地址,用于标识某个网络。主机号用于标识在该网络上的一台特定的主机。位于相同物理网络上的所有主机具有相同的网络号,如图 5-1 所示。

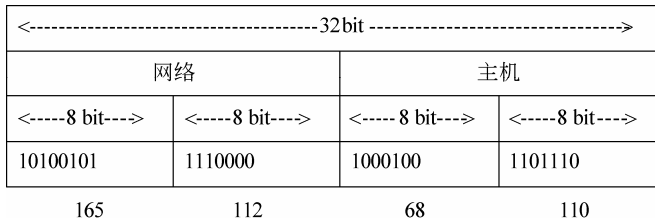


图 5-1 IP 地址的表示

为了适应于不同的规模的物理网络,IP 地址分为 A、B、C、D、E 五类,但在 Internet 上可分配使用的 IP 地址只有 A、B、C 三类。这三类地址统称为单目传送(Unicast)地址,因为这些地址通常只能分配给唯一的一台主机。D 类地址被称为多播(Multicast)地址,组播地址可用于视频广播或视频点播系统,而 E 类地址尚未使用,保留给将来的特殊用途。

不同类别的 IP 地址的网络号和主机号的长度划分不同,它们所能识别的物理网络数不同,每个物理网络所能容纳的主机个数也不同,如图 5-2 所示。

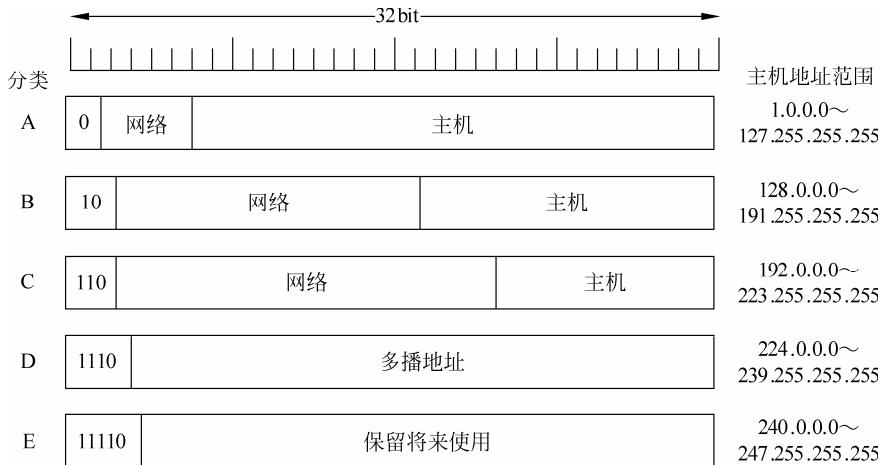


图 5-2 IP 地址的格式与分类

A 类地址用 7 位表示 IP 地址的网络部分,而用 24 位表示 IP 地址的主机部分。因此,它可以用于大型网络。B 类地址用 14 位表示 IP 地址的网络部分,而用 16 位表示 IP 地址的主要部分。它可以用于中型规模的网络。而 C 类地址用 21 位表示 IP 地址的网络部分,而用 8 位表示 IP 地址的主机部分,在一个网络中最多只能连接 256 台设备,因此,只适用于较小规模的网络。D 类地址为多播的功能保留。E 类地址为将来使用而保留。

根据 A、B、C、D、E 的高位数值,可以总结也它们的第一个字节的取值范围,如

A 类地址的第一字节的数值为 1~126。表 5-1 给出了每种地址类别第一字节的取值范围及其规模。

表 5-1 各类地址的取值范围

高位	第一个字节的十进制数	地址类别
0	1~126	A
10	128~191	B
110	192~223	C
1110	224~239	D
11110	240~254	E

5.2.3 IP 数据报的格式

要进行传输的数据在 IP 层首先需要加上 IP 头信息，封装成 IP 数据报。IP 数据报包括一个报文头以及与更高层协议相关的数据。图 5-3 所示为 IP 数据报的具体格式。

0	4	8	16	24	31 bit
版本	报头长度	服务类型	总长度		
标识符			标志	片偏移	
生存周期		协议	头部校验和		
源IP地址					
目的IP地址					
IP选项				填充	
数据…					

图 5-3 IP 数据报格式

IP 数据报的格式可以分为报头区和数据区两大部分，其中，数据区包括高层需要传输的数据，报头区是为了正确传输高层数据而增加的控制信息，这些控制信息包括：

- ❑ **版本** 长度为 4bit，表示与数据报对应的 IP 协议版本号。不同的 IP 协议版本，其数据报格式有所不同。当前的 IP 协议版本号为“4”。所有 IP 软件在处理数据报之前都必须检查版本号，以确保版本正确。IP 软件将拒绝处理版本不同的数据报，以避免错误解释其中内容。
- ❑ **报头长度** 长度为 4bit，指出以 32bit 长计算的报头长度，IP 数据报头中除 IP 选项域外，其他各域均为定长域，各定长域长度为 20B，这样一个不含选项域的普通 IP 数据报其头标长度域值为“5”。总的来说，头标长度应为 32bit 的整数位，假如不是，在头标尾部添“0”凑齐。
- ❑ **服务类型** 服务类型（Service Type）规定对本数据报的处理方式。该域长度为 1B，被分为五个子域。
其中，3bit 的“优先权”（PRECEDENCE）子域指示本数据报的优先权，表示本

数据报的重要程度。优先权取值从 0 到 7, “0” 表示一般优先权, “7” 表示网络控制优先权, 优先权值是由用户指定的。大多数网络软件对此不予理睬, 然而它毕竟提供了一种手段, 允许控制信息享受比一般数据较高的优先权。DTR 这 3 位数据表示本数据报所要的传输类型。其中, D 代表低延迟——Delay; T 代表高吞吐率——Throughput; R 代表高可靠性——Reliability。上述 3 位只是用户的请求, 不具有强制性, Internet 在寻找路径时可能以它们为参考。

- ❑ **总长** 总长域为 16bit, 指示整个 IP 数据报的长度, 以字节为单位, 其中包括报头长度及数据区长。因此, IP 数据报总长可达 $2^{16}-1$ (即 65535) B (字节)。
- ❑ **标识** 标识是信源机赋予数据报的标识符, 目标主机利用此域和信源地址判断收到的分组属于哪个数据报, 以使数据报重组。分片时, 该域必须不加修改地复制到新片头中, 数据报标识符的实现原则是对于同一信源机各标识符必须是唯一的。
- ❑ **标志** 标志为 3bit, 用于控制分片和重组。Bit0: 保留, 必须为 “0”。Bit1: 0=可以分片, 1=不分片。Bit2: 0=最后一个分片, 1=还有分片。
- ❑ **片偏移** 13bit, 也是用于分析和重组控制的, 它指出本片数据在初始数据报数据区中的偏移量, 以 8B 为单位。由于各片按独立数据报的方式传输, 其到达信宿机的顺序无法保证, 因此, 重组的片顺序由片偏移域提供。
- ❑ **生存周期** 数据报传输的一大特点是随机寻径, 因此, 从信源机到信宿机的传输延迟也具有随机性。当路由器的路由表出错时, 数据报可能会进入一条循环路径, 无休止地在 Internet 中流动。为避免这个情况, IP 协议对数据报传输延迟要进行控制。为此, 每生成一个数据报, 它都带有一个生存时间, 该时间以秒为单位, 每个处理该数据报的结点必须至少把 TTL 值减 1, 即使处理时间小于 1s。假如数据报在路由器中因等待服务而被延迟, 则从 TTL 中减去等待时间, 一旦 TTL 减至 0, 该数据报将被丢弃。
- ❑ **协议** 1B, 指创建数据报数据区数据的高级协议类型, 如 TCP、OSPF 等。
- ❑ **头部检验和** 2B, 用于保证头部数据的完整性, 其算法很简单, 设 “头部检验和” 初值为 0, 然后对头部数据每 16 位求异或, 结果取反, 便得到检验和。
- ❑ **选项** 主要用于控制和测试两个目的。

5.2.4 IP 数据报的分片与组装

当一个 IP 数据报从一个主机传输到另一个主机时, 它可能通过不同的物理网络。每个物理网络有一个最大的帧大小, 即所谓的最大传输单元 (Maximum Transmission Unit, MTU)。它限制了能够放入一个物理帧中的数据报长度。

IP 用一个进程来对超过 MTU 的数据报进行分片。这个进程建立了一个最大数据量以内的数据报集合。接收主机重新组合原始的数据报。IP 要求每个链路至少支持 68 个 8 位字节 (即 68B) 的 MTU。这是最大的 IP 报文头长度 (60 个 8 位字节, 即 60B) 和非最后分片中可能的最小数据长度 (8 个 8 位字节, 即 8B) 的总和。如果任何一个网络提供了一个比这个还小的值, 则必须在网络接口层实现分片和分片重组。这个过程对于 IP

必须是透明的。IP 实现不必处理大于 576B 的未分片的数据报。

一个未分片的数据报的分片信息字段全为 0，即多个分片标志位为 0，并且片偏移量为 0。分片一个数据报，需执行以下几个步骤：

- ❑ 检查 DF 标志位，查明是否允许分片。如果设置了该位，则数据报将被丢弃，并将一个 ICMP 错误返回给源端。
- ❑ 基于 MTU 值，把数据字段分成两个部分或者多个部分。除了最后的数据部分外，所有新建数据选项的长度必须为 8B 的倍数。
- ❑ 每个数据部分被放入一个 IP 数据报。这些数据报的报文头略微修改了原来的报文头。
- ❑ 除了最后的数据报分片外，所有分片都设置了多个分片标志位。
- ❑ 每个分片中的片偏移量字段设为这个数据部分在原来数据报中所占的位置，这个位置相对于原来未分片数据报中的开头处。
- ❑ 如果在原来的数据报中包括了选项，则选项类型字节的高位字节决定了这个信息是被复制到所有分片数据报，还是只复制到第一个数据报。
- ❑ 设置新数据报的报文头字段及总长度字段。
- ❑ 重新计算报文头部检验和字段。

此时，这些分片数据报中的每个数据报如一个完整 IP 数据报一样被转发。IP 独立地处理每个数据报分片。数据报分片能够通过不同的路由器到达目的。如果它们通过那些规定了更小的 MTU 网络，则还能够进一步对它们进行分片。

在目的主机上，数据被重新组合成原来的数据报。发送主机设置的标识符字段与数据报中的源 IP 地址和目的 IP 地址一起使用。分片过程不改变这个字段。

为了重新组合这些数据报分片，接收主机在第一个分片到达时分配一个存储缓冲区。这个主机还将启动一个计时器。当数据报的后续分片到达时，数据被复制到缓冲区存储器中片偏移量字段指出的位置。当所有分片都到达时，完整的未分片的原始数据包就被恢复了。处理如同未分片数据报一样继续进行。

如果计时器超时并且分片保持尚未认可状态，则数据报被丢弃。这个计时器的初始值称为 IP 数据报的生存期值。它是依赖于实现的。一些实现允许对它进行配置。在某些 IP 主机上可以使用 netstat 命令列出分片的细节。如 TCP/IP for OS/2 中的 netstat-i 命令。

5.2.5 IP 数据报路由选项

IP 数据报选项字段为 IP 数据报源站提供了两种显式提供路由信息的方法。它还为 IP 数据报提供了一种确定传输路由的方法。

- ❑ **不严格的源路由** 不严格的源路由选项也称为不严格的源和记录路由选项，它为 IP 数据报提供了一种显式地提供路由信息的方法。路由器在把数据报转发到目的站时使用该信息。同时还用它来记录路由。
- ❑ **严格的源路由** 严格的源路由选项也称为严格的源和记录路由（Strict Source and Record Route, SSRR）选项，除了中间路由器必须通过一个直接连接的网络把数据报发送到源路由中的下一个 IP 地址外，它使用与不严格的源路由相同的

原则。它不能使用中间路由器。如果不能实现这点,它就发出 ICMP 目的不可达的错误消息。

- ❑ **记录路由** 这个选项提供了一种记录 IP 数据报通过的路径的方法。它的功能类似于源路由选项。但是,这个选项提供了一个空的路由数据字段。这个字段在数据报通过网络时被填入。源主机必须为这个路由信息提供足够的空间。如果数据字段在数据报到达目的主机之前被填充,则在不记录这个路径的情况下继续转发这个数据报。
- ❑ **网际时间戳** 该选项强制目的路由上的一些或所有路由器把一个时间戳放入选项数据中。时间戳按秒度量,并且可以用于调试的目的。但由于大多数 IP 数据在不到 1s 的时间内就被转发及 IP 路由器不需要有同步的时钟,导致时间戳不精确。因此,它不能用于性能度量。

5.3 子网掩码

子网掩码(Subnet Mask)又称网络掩码,是一种用来指明一个 IP 地址的哪些位标识的是主机所在的子网以及哪些位标识的是主机的位掩码。它不能单独存在,必须结合 IP 地址一起使用。

5.3.1 子网掩码概述

互联网是由许多小型网络构成的,每个网络上都有许多主机,这样便构成了一个有层次的结构。IP 地址在设计时就考虑到地址分配的层次特点,将每个 IP 地址都分割成网络号和主机号两部分,以便于 IP 地址的寻址操作。

如果不指定 IP 地址的网络号和主机号各是多少位,就不知道哪些位是网络号、哪些是主机号,这就需要通过子网掩码来实现。子网掩码不能单独存在,必须结合 IP 地址一起使用。子网掩码只有一个作用,就是将某个 IP 地址划分成网络地址和主机地址两部分。子网掩码的设定必须遵循一定的规则。

与 IP 地址类似,子网掩码也是由 32 位的二进制数构成,其左边用若干个连续的二进制数字“1”表示,右边用若干个连续的二进制数字“0”表示,其格式如图 5-4 所示。这样通过左边若干连续个数的“1”及右边若干连续个数的“0”能够区分 IP 地址的网络号和主机号部分。

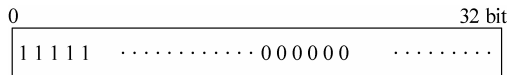


图 5-4 子网掩码格式

子网掩码也通常使用点分十进制数的方法来表示。例如, 255.255.255.0 就表示一个子网掩码,与转后的二进制数关系如图 5-5 所示。并且它是 C 类 IP 地址的默认子网掩码(本书将在后面进行详细讲解)。

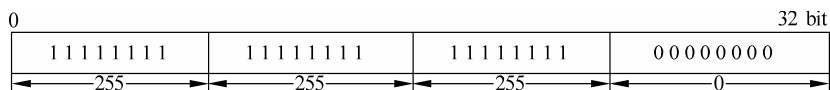


图 5-5 子网掩码十进制与二进制对应关系

常用的子网掩码有数百种，这里只介绍最常用的两种子网掩码，它们分别是“255.255.255.0”和“255.255.0.0”。

□ 子网掩码是“255.255.255.0”的网络

最后面一个数字可以在 0~255 范围内任意变化，因此可以提供 256 个 IP 地址。但是实际可用的 IP 地址数量是 256-2，即 254 个，因为主机号不能全是“0”或全是“1”。

□ 子网掩码是“255.255.0.0”的网络

后面两个数字可以在 0~255 范围内任意变化，可以提供 2552 个 IP 地址。但是实际可用的 IP 地址数量是 2552-2，即 65023 个。

下面介绍子网掩码的相关知识：

1. 子网

对于企业所有主机位于同一网络层次中，不方便管理员对其进行管理。因此，提出了将大网络进一步划分成小网络，而这些小网络就称为“子网”。

IP 地址的子网掩码设置不是任意的。如果将子网掩码设置过大，也就是说子网范围扩大。根据子网寻径规则，很可能发往和本地机不属于同一子网内的计算机，会因为错误的判断而认为目标计算机是在同一个子网内中。那么，数据包将在本子网内循环，直到超时并抛弃，使数据不能正确到达目标的计算机，导致网络传输错误。

如果将子网掩码设置得过小，那么会将本来属于同一子网内的机器之间的通信当作是跨子网传输，数据包都交给默认网关处理，这样势必增加默认网关的负担，造成网络效率下降。因此，子网掩码应该根据网络的规模进行设置。如果一个网络的规模不超过 254 台计算机，采用“255.255.255.0”作为子网掩码就可以了。

2. 掩码

掩码与 IP 地址相对应，具有 32 位地址，当用掩码与 IP 地址进行逐位“逻辑与”(AND)运算后，就能够得知该 IP 地址的网络地址(网络号)。例如，一个 IP 地址为 221.180.60.15，其默认掩码为“255.255.255.0”，与 IP 地址进行 AND 运算后，可得知其网络地址为“201.180.60.0”，如图 5-6 所示。

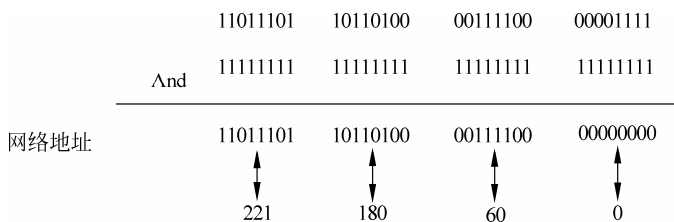


图 5-6 掩码的作用

通常 A 类、B 类和 C 类 IP 地址都有其默认掩码，如图 5-7 所示。

地址类别	默认掩码（十进制）	默认掩码（二进制）
A	255.0.0.0	11111111 .00000000 .00000000 .00000000
B	255.255.0.0	11111111 .11111111 .00000000 .00000000
C	255.255.255.0	11111111 .11111111 .11111111 .00000000

图 5-7 各类地址的默认掩码

5.3.2 子网掩码的计算

在对子网进行划分时，需要使用子网掩码，通过子网掩码，能够表明网络中一台主机所在的子网与其他子网的关系，这就需要计算子网掩码。计算子网掩码的方法有利用划分的子网个数和计算子网中主机的数量两种。

1. 利用子网数计算子网掩码

在利用子网数计算子网掩码之前，需要了解具体要划分的子网个数，其具体步骤如图 5-8 所示。

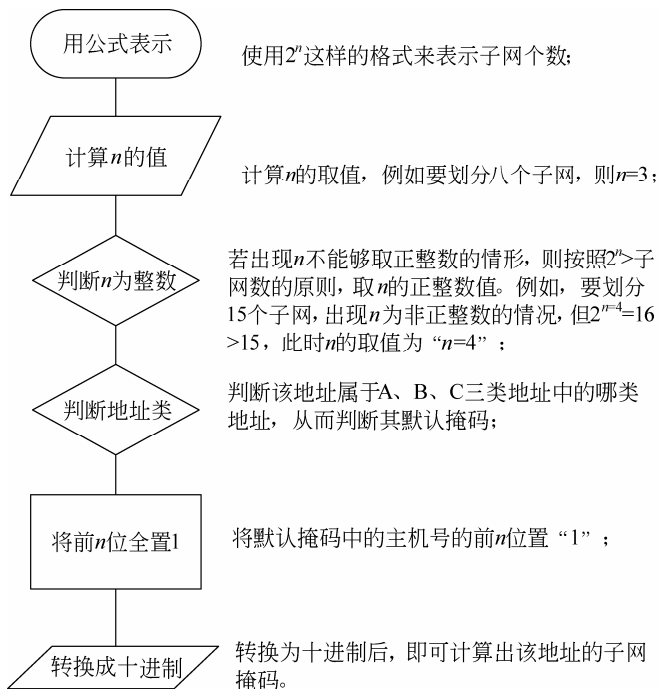


图 5-8 计算子网掩码流程

例如，现将网络地址为“129.65.0.0”的网络划分为 27 个子网，则其子网掩码的计

算方法为:

首先由 $2^n=32>27$ 确定 n 的取值为 5, 然后根据 “129.65.0.0” 的网络地址, 判断其属于 B 类 IP 地址, 其默认掩码为 “255.255.0.0”。最后, 将默认掩码中主机号的前五位置为 1, 即 “11111000”, 转换为十进制为 248, 因此其划分子网的子网掩码为 “255.255.248.0”。

通过前面的介绍, 我们可以得出一个规律, 即从划分子网的个数就能够判断出其子网掩码, 对于 B 类网络来讲, 其子网划分个数与子网掩码即每一个子网的主机数之对应关系如图 5-9 所示。

对于 C 类网络来讲, 其划分子网个数与子网掩码及每个子网所能够容纳的主机数量, 如图 5-10 所示。

子网个数	子网掩码
2	255. 255. 128. 0
4	255. 255. 192. 0
8	255. 255. 224. 0
16	255. 255. 240. 0
32	255. 255. 248. 0
64	255. 255. 252. 0
128	255. 255. 254. 0
256	255. 255. 255. 0

子网个数	子网掩码
2	255. 255. 255. 128
4	255. 255. 255. 192
8	255. 255. 255. 224
16	255. 255. 255. 240
32	255. 255. 255. 248
64	255. 255. 255. 252

图 5-9 B 类网络子网个数与子网掩码对应关系

图 5-10 C 类网络子网个数与子网掩码对应关系

2. 利用主机数计算子网掩码

在利用主机数计算子网掩码时, 必须知道每个子网所需容纳的主机个数, 而不必知道其需要划分的子网个数, 其主要步骤如图 5-11 所示。

例如, 要将网络号为 180.195.0.0 的网络划分成若干子网, 要求其每个子网能够容纳的主机数量为 900 台, 那么其子网掩码计算方法为:

- (1) 由 $2^n=1024>900$, 可以确定 n 的取值为 10;
- (2) 根据 “180.195.0.0” 的网络, 判断属于 B 类 IP 地址, 其默认掩码为 “255.255.0.0”;
- (3) 将默认掩码中主机号的所有位全部转换为 1, 即 “11111111 11111111”, 接着按照由低位到高位顺序将 $n=10$ 位全部转换为 0, 即 “11111100 000000” 转换为十进制为 252。

因此其划分子网的子网掩码为 “255.255.252.0”。

同过前面的计算, 可以得出一个规律, 即按照子网能够容纳的主机数量我们也能够计算出其子网掩码。对于 B 类网络来讲, 其子网能够容纳主机数量与子网掩码的关系,

如图 5-12 所示。

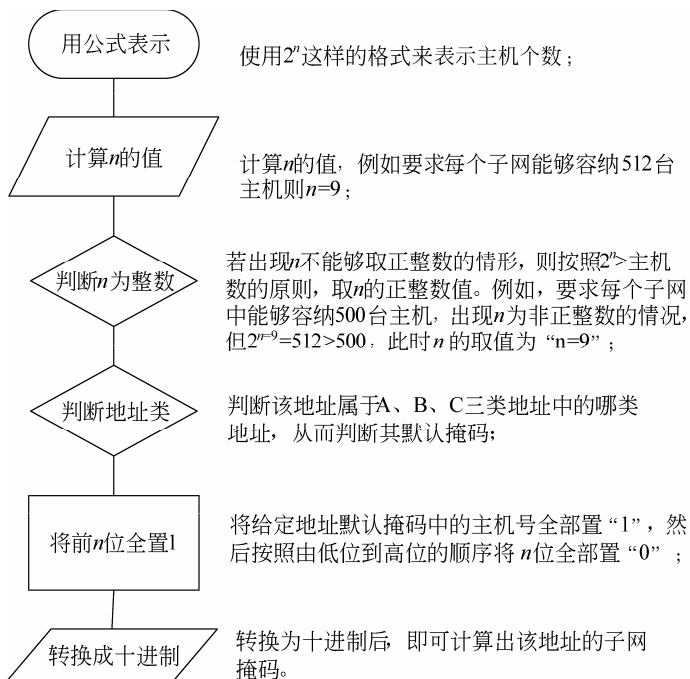


图 5-11 利用主机数计算子网掩码流程

每子网容纳主机数量	子网掩码
32766	255.255.128.0
16382	255.255.192.0
8190	255.255.224.0
4094	255.255.240.0
2046	255.255.248.0
1022	255.255.252.0
510	255.255.254.0
254	255.255.255.0

图 5-12 B 类网络子网掩码每个子网所能够容纳的主机数量的对应关系

对于 C 类网络，其子网掩码每个子网所能够容纳的主机数量的对应关系，如图 5-13 所示。

每子网容纳主机数量	子网掩码
126	255.255.255.128
62	255.255.255.192
30	255.255.255.224
14	255.255.255.240
6	255.255.255.248
2	255.255.255.252

图 5-13 C 类网络子网掩码每个子网所能够容纳的主机数量的对应关系

5.4 IPv6 概述

随着 Internet 的飞速发展,使得网络运营商深切地感受到了,因 IPv4 地址不足而产生的制约性,IP 地址空间不足已经成国际关注的焦点。为解决 IPv4 存在缺陷的问题,新一代 IPv6 协议已经被列入许多国内网络和通信运营商的网络规划中。以 IPv6 协议为核心的大规模下一代网络建设项目已经正式启动。

5.4.1 IPv6 的产生与发展

现有的互联网是在 IPv4 协议的基础上运行的。近年来随着互联网在各个领域内的迅速发展,人们对信息资源的开发和利用进入了一个全新的阶段,这也使得 IPv4 越来越捉襟见肘:IP 地址资源越来越紧张,路由表越来越庞大,路由速度越来越慢等。地址生存期预测工作组预言,现有的 IPv4 地址空间将在 2005~2011 年之间的某个时间耗尽。

虽然各方面都在研究一些补救的方法,如 CIDR 的使用以及 DHCP 使用的增加来缓解地址空间的压力,但这些方法并不能解决 IPv4 的先天不足。除了地址空间消耗问题外,IPv4 其他方面的限制也要求定义一种新的 IP 协议。如即使使用了 CIDR,路由表也会变得太大,以至无法对它们进行管理以及其服务类型没有得到清晰的定义,几乎没有得到使用。鉴于这些问题,1990 年开始着手开发 IP 的新版本 IPv6。

IPv6 报文格式比 IPv4 的报文格式要简单。IPv6 报文头的格式长度从 20B 增加到 40B,它包含两个 16B 的源地址和目的地址,前面是 8B 的控制信息。而 IPv4 报文头有两个 4B 地址,前面是 12B 的控制信息,后面可能还有选项数据。大多数的 IPv6 报文的报文关中减少了控制信息并取消了选项数据,其目的在于优化路由器每个报文的处理时间。报文头中已删除的经常使用的字段,在需要时可以把它们移到可选的扩展报文头中。

IPv6 的编址模型在 RFC2373 中进行了描述。它使用 128 位地址,而不用 IPv4 中的 32 位地址。因此 IPv6 在理论上允许 2^{128} 个地址。即使像现在使用 IPv4 地址空间的效率一样来使用 IPv6 地址空间,也允许地球上每平方米土地上有 50 000 个地址。IPv6 地址

表示为用冒号分开的八个十六进制数的形式。IPv6 定义了单播地址、组播地址和任意播地址三种地址类型。

IPv6 中还引入了流（Flow）的概念，它是从源发送到目的的一系列相关报文，且要求中介路由器的特殊处理。如：实时服务。在源和目的之间可以有多个活动流，以及以任何流关联的流量。每个流分别用 IPv6 报文中的 24 位流标志字段所标记。

总的说来，IPv6 的主要优势体现在以下几方面：

- ☐ 扩大地址空间。
- ☐ 提高网络的整体吞吐量。
- ☐ 改善服务质量（QoS）。
- ☐ 安全性有更好的保证。
- ☐ 支持即插即用和移动性。
- ☐ 更好实现多播功能。

显然，IPv6 的优势能够对上述挑战直接或间接地作出贡献。其中最突出的是 IPv6 大大地扩大了地址空间，恢复了原来因地址受限而失去的端到端连接功能，为互联网的普及与深化发展提供了基本条件。当然，IPv6 并非十全十美、一劳永逸，不可能解决所有问题。IPv6 只能在发展中不断完善，也不可能在一夜之间发生，过渡需要时间和成本，但从长远看，IPv6 有利于互联网的持续和长久发展。

5.4.2 IPv6 的新特性

IPv6 协议的提出，不仅解决了 IPv4 协议存在的如地址空间不足的严重缺陷，而且还针对 Internet 的应用需求，改进了 IPv4 协议报头提供了一些新的机制从而很好的解决了诸如移动性、安全性、多媒体传输等问题。

1. 新的报文结构

IPv6 使用了全新的协议头格式，在 IPv6 报文头中包括基本报头和可扩展报头两部分，其格式如图 5-14 所示。

基本报头	扩展报头1	...	扩展报头n	数据
------	-------	-----	-------	----

图 5-14 IPv6 报头新格式

其中，基本报头与原 IPv4 报头类似，但在 IPv6 报头中添加了一些新的字段，以及改变了某些字段，其格式如图 5-15 所示。

0		8	32bit	
版本	优先级	流标识		
有效负荷长度		下一个报头	跳数限制	
源站IP地址（128 bit）				
目的站IP地址（128 bit）				

图 5-15 IPv6 基本报头格式

扩展报头中提供了许多额外信息，共分为六种扩展报头，且它们是可选的，当有多个扩展报头时，这些扩展报头必须按照一定的排列次序，跟在基本报头之后，图 5-16 所示为其排列次序。

扩展报头类型	描 述
站到站选项	用于路由器的各种信息
源路由报头	指明数据报到达目的站所必须经过的路由器
分片报头	用于数据报的分片管理
身份验证报头	用于接收端对发送端身份的验证
加密报头	用于传输经过加密的报文信息
目的地选项	有关数据报接收端的附加信息

图 5-16 IPv6 扩展报头排列次序

IPv6 采用 128 位 IP 地址，其地址空间为 2^{128} 。此外，一些非根本性的和可选择的字段被移动到了 IPv6 协议头之后的扩展协议中，从而简化了路由器的选择过程，使得网络中的中间路由器在处理 IPv6 协议头时，有更高的效率。

2. 新的地址配置方式

随着网络技术的发展，在 Internet 上的结点不再单单是计算机了，它将发展成为包括个人数字助理（PDA）、移动电话（Mobile Phones）、甚至包括冰箱、电视等家用电器，这就要求 IPv6 主机地址配置更加简化。

因此，为了达到简化地址配置的目的，在 IPv6 中除了支持手动地址配置和有状态自动地址配置（使用专用的地址分配服务器动态分配地址）外，还支持一种无状态的地址配置技术。在该技术中，网络上的主机能够给自己自动配置 IPv6 地址，在同一链路上，所有主机不需人工干涉就可以进行通信。

3. QoS（服务质量）保证

在 IPv6 的基本报头中新定义了一种称为流标识的全新字段。它使得网络中的路由器能够对属于一个流标识的数据包进行识别并提供特殊处理。IPv6 的流标识字段，使得路由器可以在不打开传输的内层数据包的情况下就可以识别流，这就是说即使数据包有效负荷已经进行了加密，但仍然可以实现对 QoS 的支持。

4. 支持实时音频和视频传输

在 IPv6 的报头结构中取消了服务类型字段，增加了流标识字段，该字段除了能提供 QoS 外，还使得源主机可以请求对数据做出特殊的处理，如能够支持实时音频和视频的传输。

5. 移动性

在 IPv6 的可扩展报头中，采用了 Routing Header（路由报头）和 Destination Option

Header（目的地选项报头）报头类型，使得 IPv6 对移动性提供了内在的支持。在 IPv6 中能够给移动结点分配一个本地地址，通过此地址总可以访问到它。在移动结点位于本地时，它连接到本地链路并使用其本地地址。在移动结点远离本地时，本地代理（通常是路由器）在该移动结点和正与其进行通信的结点之间传递消息。这样就达到了设备能够在 Internet 上随意改变位置但仍能够维持现有连接的目的。

5.5 IPv6 基础知识

在 IPv4 中地址长度为 32 位，能够容纳最大地址数量为 2^{32} ；而在 IPv6 中地址长度扩展到 128 位，能够容纳的最大地址空间数量为 2^{128} 。显然，IPv6 地址的数量要远远大于 IPv4 地址的数量。

5.5.1 IPv6 编址

原来用于表示 IPv4 地址的点分十进制法显然不在适用于新的 IPv6 地址。人们研究出了更适合于 IPv6 地址的表示方法，即冒号十六进制表示法，它又包括不同的方法。

1. 首选格式

在 IPv6 地址的三种表示方法中，其首选格式为：“x: x: x: x: x: x: x: x”，其中 x 的值是十六进制数值，分别对应于 128 位地址中的八个 16 位地址段。例如：

```
2001: FECD: BA23: CD1F: DCB1: 1010: 9235: 4088;  
FEDC: BA56: DCB2: 3512: BA56: FEDC: 1010: 3220;
```

另外，个别字段前面的“0”可以省略不写，但是每段必须至少有 1 位数字，例如：

```
1080: 0000: 0000: 0000: 0008: 0800: 200C: 417A  
可以简写为：1080: 1080: 0: 0: 0: 8: 800: 200C: 417A。
```

2. 压缩格式

在分配某些形式的 IPv6 地址时，该地址可能是一个包含一长串“0”位的地址，为了简化包含一长串“0”位地址的书写，方便文本形式的描述，指定了一种特殊的语法格式来压缩 0。

在这种特殊的格式中，使用“::”符号表示有多组 16 位“0”，但是“::”符号只能在一个地址中出现一次，可用于压缩一个地址中的前导、末尾或相邻的 16 位“0”。例如：1080: 1080: 0: 0: 0: 8: 800: 200C: 417A 使用“::”符号压缩后表示为“1080: 1080:: 8: 800: 200C: 417A”；FEC0: 1: 0: 0: 0: 0: 0: 1234 使用“::”符号压缩后表示为“fec0: 1:: 1234”。

3. 混合格式

当处理拥有 IPv4 和 IPv6 结点的混合环境时，可以使用该形式。即“x: x: x: x: x:

x: d.d.d.d”，其中，“x”是 IPv6 地址的 96 位高位顺序字节（六个高阶 16 位段）的十六进制值，“d”是 32 位低位顺序字节（四个低阶 8 位段）的十进制值。通常，对于“映射 IPv4 的 IPv6 地址”以及“兼容 IPv4 的 IPv6 地址”来讲，可以采用这种表示法来表示。

例如：

0: 0: 0: 0: 0: 0: 12.1.68.123 用此方法可以表示为 “: : 12.1.68.123”。

4. IPv6 地址前缀表示

IPv6 地址前缀与 IPv4 中的 CIDR（无类域间路由）相似，并写入 CIDR 表示法中。IPv6 地址前缀的表示方法为：“IPv6 地址/前缀长度”，其中 IPv6 地址是指前面所讲格式中的任意一种地址格式，而前缀长度的值为十进制数值，表示前缀由多少个最左侧相邻位构成。例如：

fec0:0:0:1::1234/64，其中地址的前 64 位“FEC0: 0: 0: 1”构成了地址的前缀。在 IPv6 地址中，地址前缀用于表示 IPv6 地址中有多少位表示子网。

5.5.2 IPv6 的地址分类

所有类型的 IPv6 地址都被分配到接口，而不是结点。IPv6 地址是单个或一组接口的 128 位标识符，有三种类型，详细介绍如下：

1. 单播（Unicast）地址

单播地址是指具有单一接口的地址，其中一个单接口有一个标识符。发送给一个单播地址的数据包被送到由该地址标识的接口。对于有多个接口的结点，它的任何一个单播地址都可以用作该结点的标识符。

通常，单播地址在逻辑上划分为子网前缀和接口 ID 两部分，其格式如图 5-17 所示。其中接口 ID 用于标识链路接口，在该链路中其值必须是唯一的，一个接口标识符应与该接口的链路层地址相同，该链路通常由子网前缀来标识。

n bit	$(128-n)$ bit
子网前缀	接口 ID

图 5-17 单播地址逻辑结构

IPv6 单播地址是用连续的位掩码聚集的地址，类似于 CIDR 的 IPv4 地址。IPv6 中的单播地址分配有多种形式，包括全部可聚集全球单播地址、NSAP 地址、IPX 分级地址、站点本地地址、链路本地地址以及运行 IPv4 的主机地址。单播地址中有下列两种特殊地址：

1) 不确定地址

单播地址 0:0:0:0:0:0:0:0 称为不确定地址。它不能分配给任何结点。它的一个应用示例是初始化主机时，在主机未取得自己的地址以前，可在它发送的任何 IPv6 包的源地址字段放上不确定地址。不确定地址不能在 IPv6 包中用作目的地址，也不能用在 IPv6 路

由头中。

2) 回环地址

单播地址 0:0:0:0:0:0:1 称为回环地址。结点用它来向自身发送 IPv6 包。它不能分配给任何物理接口。

2. 任播 (AnyCast) 地址

任播 IPv6 地址也称为任意点播 IPv6 地址, 它是指一组接口的地址 (一般属于不同结点) 具有一个标识符, 发送到任播地址的数据包被送到由该地址标识的、根据路由选择距离度量最近的一个接口上。

任播地址是从单播的地址空间分配而来, 可用任何一种规定的单播地址格式。因此, 在语法上是无法区别单播地址和任播地址的。当一个单播地址分配给多个接口时, 如果把它转为任播地址, 那么被分配该地址的结点, 必须显示地配置, 以便知道这是一个任播地址。其中, 预定的子网路由器任播地址格式如图 5-18 所示。

n bit	$(128-n)$ bit
子网前缀	00000000000000

图 5-18 任播地址预定格式

其中, 子网前缀用来标识一条特定链路; 接口标识为“0”的链路上的一个接口, 其任播地址与单播地址在语法上是相同的。IPv6 任意播地址存在下列限制:

- ☐ 任意播地址不能用作源地址, 而只能作为目的地址;
- ☐ 任意播地址不能指定给 IPv6 主机, 只能指定给 IPv6 路由器;
- ☐ IPv6 任意播地址。

3. 多播 (MultiCast) 地址

多播 IPv6 地址是指一组接口的地址 (通常分属不同结点), 其中这一组接口具有一个标识符。发送到一个多播地址的数据包被送到由该地址标识的每个接口上。简单的说, 多播地址就是一组结点的标识符。其格式如图 5-19 所示。

8 bit	4 bit	4 bit	112 bit
11111111	Flag	Extent	Group ID

图 5-19 多播地址格式

对图中各字段的说明如下:

- ☐ 第一字段是标识字段, 共占 8 位, 所有位全部为“1”, 用于标识该地址是一个多播地址。
- ☐ Flag (标志字段), 共占 4 位, 其格式为“000T”。其中前 3 位为高位是保留位, 其初始值为 0; T 的取值包括 0 和 1。若 T=0, 则表示一个永久分配的多播地址, 由全球 Internet 编号机构进行分配; 若 T=1, 则表示一个非永久的多播地址。
- ☐ Extent (范围字段), 共占 4 位, 主要用于限制多播组的范围。该字段的可能取值及各值所代表意义如图 5-20 所示。

值	含 义
0	保留
1	本地多播组
2	链路本地范围
3、4	未分配
5	站点本地范围
6、7	未分配
8	组织本地范围

图 5-20 Extent 字段取值

- Group ID (组标识), 共占 112 位, 主要用于标识多播组, 可以是永久的也可以是临时的。如常用的多播组有所有结点地址=FF02:0:0:0:0:0:1 (链路本地); 所有路由器地址=FF02:0:0:0:0:0:2 (链路本地); 所有路由器地址=FF05:0:0:0:0:0:2 (站点本地)。

5.5.3 IPv6 数据报

IPv6 包由 IPv6 包头 (40B 固定长度)、扩展包头和上层协议数据单元三部分组成。IPv6 包扩展包头中的分段包头中指明了 IPv6 包的分段情况。其中, 不可分段部分包括 IPv6 包头、Hop-by-Hop 选项包头、目的地选项包头 (适用于中转路由器) 和路由包头; 可分段部分包括认证包头、ESP 协议包头、目的地选项包头 (适用于最终目的地) 和上层协议数据单元。但是需要注意的是, 在 IPv6 中, 只有源结点才能对负载进行分段, 并且 IPv6 超大包不能使用该项服务。

1. IPv6 数据包：包头

IPv6 包头长度固定为 40B, 去掉了 IPv4 中一切可选项, 只包括八个必要的字段, 因此尽管 IPv6 地址长度为 IPv4 的四倍, IPv6 包头长度仅为 IPv4 包头长度的两倍。其中的各个字段格式如下:

Version (版本号): 4 位, IP 协议版本号, 值。

IPv6 数据包中的包头可以分为下面七个组成部分, 详细介绍如下:

- TrafficClass (通信类别) 8 位, 指示 IPv6 数据流通信类别或优先级。功能类似于 IPv4 的服务类型 (TOS) 字段。
- FlowLabel (流标记) 20 位, IPv6 新增字段, 标记需要 IPv6 路由器特殊处理的数据流。该字段用于某些对连接的服务质量有特殊要求的通信, 诸如音频或视频等实时数据传输。在 IPv6 中, 同一信源和信宿之间可以有多种不同的数据流,

彼此之间以非“0”流标记区分。如果不要路由做特殊处理,则该字段值置为“0”。

- ❑ **PayloadLength** (负载长度) 16 位负载长度。负载长度包括扩展头和上层 PDU, 16 位最多可表示 65535 字节负载长度。超过这一字节数的负载, 该字段值置为“0”, 使用扩展头逐个跳段 (Hop-by-Hop) 选项中的巨量负载 (JumboPayload) 选项。
- ❑ **NextHeader** (下一包头) 8 位, 识别紧跟 IPv6 头后的包头类型, 如扩展头 (有的话) 或某个传输层协议头 (诸如 TCP, UDP 或者 ICMPIPv6)。
- ❑ **HopLimit** (跳段数限制) 8 位, 类似于 IPv4 的 TTL (生命期) 字段。与 IPv4 用时间来限定包的生命期不同, IPv6 用包在路由器之间的转发次数来限定包的生命期。包每经过一次转发, 该字段减 1, 减到 0 时就把这个包丢弃。
- ❑ **SourceAddress** (源地址) 128 位, 发送方主机地址。
- ❑ **DestinationAddress** (目的地址) 128 位, 在大多数情况下, 目的地址即信宿地址。但如果存在路由扩展头的话, 目的地址可能是发送方路由表中下一个路由器接口。

2. IPv6 数据包: 扩展包头

IPv6 包头设计中对原 IPv4 包头所做的一项重要改进就是将所有可选字段移出 IPv6 包头, 置于扩展头中。由于除 Hop-by-Hop 选项扩展头外, 其他扩展头不受中转路由器检查或处理, 这样就能提高路由器处理包含选项的 IPv6 分组性能。

通常, 一个典型的 IPv6 包, 没有扩展头。仅当需要路由器或目的结点做某些特殊处理时, 才由发送方添加一个或多个扩展头。与 IPv4 不同, IPv6 扩展头长度任意, 不受 40B 限制, 以便于日后扩充新增选项, 这一特征加上选项的处理方式使得 IPv6 选项能以真正的利用。但是为了提高处理选项头和传输层协议的性能, 扩展头总是 8B 长度的整数倍。目前, RFC2460 中定义了以下六个 IPv6 扩展包头, 详细介绍如下:

(1) Hop-by-Hop (逐个跳段) 选项包头: Hop-by-Hop 选项包头包含分组传送过程中, 每个路由器都必须检查和处理的特殊参数选项。其中的选项描述一个分组的某些特性或用于提供填充。这些选项包括 Pad1 选项 (选项类型为 0), 填充单字节。PadN 选项 (选项类型为 1), 填充两个以上字节。JumboPayload 选项 (选项类型为 194), 用于传送超大分组。使用 JumboPayload 选项, 分组有效载荷长度最大可达 4 294 967 295 字节。负载长度超过 65 535 字节的 IPv6 包称为“超大包”。路由器警告选项 (选项类型为五), 提醒路由器分组内容需要做特殊处理。路由器警告选项用于组播收听者发现和 RSVP (资源预定) 协议。

(2) 目的地选项包头: 目的地选项包头指名需要被中间目的地或最终目的地检查的信息。它有两种用法, 一种是如果存在路由扩展头, 则每一个中转路由器都要处理这些选项; 另一种是, 如果没有路由扩展头, 则只有最终目的结点需要处理这些选项。

(3) 路由包头: 类似于 IPv4 的松散源路由。IPv6 的源结点可以利用路由扩展包头指定一个松散源路由, 即分组从信源到信宿需要经过的中转路由器列表。

(4) 分段包头: 提供分段和重装服务。当分组大于链路最大传输单元 (MTU) 时,

源结点负责对分组进行分段，并在分段扩展包头中提供重装信息。

(5) 认证包头：提供数据源认证、数据完整性检查和反重播保护。认证包头不提供数据加密服务，需要加密服务的数据包，可以结合使用 ESP 协议。

(6) ESP 协议包头：提供加密服务。

3. IPv6 数据包：上层协议数据单元

上层数据单元即 PDU，全称为 ProtocolDataUnit。PDU 由传输头及其负载（如 ICMP IPv6 消息或 UDP 消息等）组成。而 IPv6 包有效负载则包括 IPv6 扩展头和 PDU，通常所能允许的最大字节数为 65535B，大于该字节数的负载可通过使用扩展头中的 Jumbo Payload 选项进行发送。

5.5.4 ICMPv6

ICMPv6 是 Internet Control Message Protocol Version 6 的简称，译为第六版互联网控制信息协议。

Internet 控制信息协议（ICMP）是 IP 协议的一个重要组成部分。通过 IP 包传送的 ICMP 信息主要用于涉及网络操作或错误操作的不可达信息。ICMP 包发送是不可靠的，所以主机不能依靠接收 ICMP 包解决任何网络问题。

ICMP 在 IPv6 定义中重新修订。此外，IPv4 组成员协议（IGMP）的多点传送控制功能也嵌入到 ICMPv6 中。ICMPv6 功能主要包括以下几个方面：

- ❑ 通告网络错误 比如，某台主机或整个网络由于某些故障不可达。如果有指向某个端口号的 TCP 或 UDP 包没有指明接受端，这也由 ICMP 报告。
- ❑ 通告网络拥塞 当路由器缓存太多包，由于传输速度无法达到它们的接收速度，将会生成“ICMP 源结束”信息。对于发送者，这些信息将会导致传输速度降低。当然，更多的 ICMP 源结束信息的生成也将引起更多的网络拥塞，所以使用起来较为保守。
- ❑ 协助解决故障 ICMP 支持 Echo 功能，即在两个主机间一个往返路径上发送一个包。Ping 是一种基于这种特性的通用网络管理工具，它将传输一系列的包，测量平均往返次数并计算丢失百分比。
- ❑ 通告超时 如果一个 IP 包的 TTL 降低到零，路由器就会丢弃此包，这时会生成一个 ICMP 包通告这一事实。TraceRoute 是一个工具，它通过发送小 TTL 值的包及监视 ICMP 超时通告可以显示网络路由。

与 IPv4 一样 IPv6 的报头和扩展报头并没有提供报错功能。在 RFC2463 中定义了 ICMPv6，并且在实现 IPv6 时必须实现 ICMPv6。

在邻结点发现和多播侦听发现的不同情况下，ICMPv6 协议提供了一个数据包结构的框架。ICMPv6 报文主要分为两类，详细介绍如下：

- ❑ 差错报文 由目标结点或者中间路由器发送，用于报告在转发和传送 IPv6 数据包过程中出现的错误。在所有的 ICMPv6 差错报文中，8 位类型字段中的最高位都为 0。因此，对于 ICMPv6 差错报文的类型字段，其有效值范围就是 0~127。

ICMPv6 的差错报文包括以下几种类型：目标不可达（Destination unreachable）、数据包过长（Packet Too Big）、超时（Time Exceeded）和参数问题（Parameter Problem）。

- ❑ **信息报文** 提供诊断功能和附加的主机功能，比如多播侦听发现（MLD）和邻结点发现（ND）。在所有的 ICMPv6 报文中，8 位类型字段中的最高位都为 1。因此其有效值的范围就是 128~255。根据 RFC2463，ICMPv6 的信息报文包括会送请求报文（Echo request）和会送应答报文（Echo Reply）。

ICMPv6 的报头由前一个报头中的下一个报头字段的值 58 来标识。ICMPv6 报头中的字段包括：

- ❑ **类型** 表示 ICMPv6 报文的类型，此字段长度为 8 位。在 ICMPv6 差错报文中，此字段的最高位为 0，在 ICMPv6 信息报文中，此字段的最高位为 1。
- ❑ **代码** 区分某一给定类型报文中的多个不同报文，此字段的长度为 8 对某一类型的第一个（或者只有一个）报文，代码字段的值为 0。
- ❑ **检验和** 存放 ICMPv6 报文的检验和。此字段的长度为 16 位。当计算检验和时，IPv6 的伪报头被加到 ICMPv6 报文之前。
- ❑ **报文主体** 包括 ICMPv6 报文专有的（message-specific）数据。

5.6 邻居发现（ND）协议

IPv6 邻居发现（ND）是一组确定邻居结点之间关系的消息和过程。ND 代替了在 IPv4 中使用的“地址解析协议（ARP）”、“Internet 控制消息协议（ICMP）”、路由器发现和 ICMP 重定向，并提供了其他功能。ND 在 RFC 2461“用于 IP 版本 6（IPv6）的邻居发现”（Neighbor Discovery for IP Version 6 （IPv6））中进行了描述。

邻居发现（ND）协议的使用主要可分为三个方面，包括 ND 由主机使用、ND 由路由器使用和 ND 由结点使用。其中，在 ND 由主机使用中，主要用于探索邻居路由器、探索地址、地址前缀和其他配置参数；在 ND 由路由器使用中，主要用于公告它们的存在、主机配置参数以及处于链路的前缀，通知主机更好的下一个跃点地址，以便转发用于特定目标的数据包；在 ND 由结点使用中，主要用于解析 IPv6 数据包所转发到的邻居结点的链路层地址，确定邻居结点的链路层地址何时发生变化，确定 IPv6 数据包是否可以发送到邻居和能否收到来自邻居的数据包。邻居发现（ND）协议的描述过程如表 5-2 所示。

表 5-2 ND 描述过程

过程	说明
路由器探索	主机探索附加链路上的本地路由器（等价于 ICMPv4 路由器发现）并自动配置默认路由器（等价于 IPv4 中的默认网关）的过程
前缀探索	主机探索用于本地目标的网络前缀的过程
参数探索	主机探索其他操作参数的过程，包括链路最大传输单元（MTU）和用于出站数据包的默认 hop 限制
地址自动配置	为全状态地址配置服务器（例如，动态主机配置协议版本 6（DHCPv6））上的接口配置 IP 地址的过程

续表

过程	说明
地址解析	结点将邻居结点的 IPv6 地址解析为它的链路层地址（等价于 IPv4 中的 ARP）的过程。所解析的链路层地址将成为结点的邻居高速缓存（等价于 IPv4 中的 ARP 高速缓存）中的项。在运行 Windows XP 的计算机上可以通过使用 <code>ipv6 nc</code> 命令查看邻居高速缓存的内容。邻居高速缓存将显示用于邻居高速缓存项的接口标识、邻居结点的 IPv6 地址、相应的链路层地址，以及邻居高速缓存项的状态
下一个跃点确定	结点确定根据目标地址而将数据包转发到的邻居的 IPv6 地址的过程。转发或下一个跃点的地址不是数据包被发送到的目标地址，就是邻居路由器的地址。用于目标的已解析的下一个跃点地址将成为结点的目标高速缓存中的项（也称为路由高速缓存）。在运行 Windows XP 的计算机上可以使用 <code>ipv6 rc</code> 命令查看路由器高速缓存的内容。路由高速缓存将显示目标地址、接口标识和下一个跃点地址、接口标识和在发送到目标时用作源地址的地址，以及用于目标的路径 MTU
邻居无法建立连接检测	结点确定 IPv6 数据包不能被发送的邻居结点和从邻居结点接收到的过程。在邻居的链路层地址被确定后，将跟踪邻居高速缓存中项的状态。如果邻居不再接收和发送回数据包，则邻居高速缓存项最终将被删除。路径无法建立连接检测为 IPv6 提供了一种机制，用来确定邻居主机或路由器在本地网段上不再可用
重复地址检测	结点确定被认为在使用的地址已经不再由邻居结点使用的过程（等价于 IPv4 中无偿 ARP 帧的使用）
重定向功能	路由器通知主机更好地到达目标的第一个跃点 IPv6 地址的过程（等价于 IPv4 ICMP 重定向消息的功能）

5.7 DHCPv6 协议

DHCPv6（Dynamic Host Configuration Protocol for IPv6，支持 IPv6 的动态主机配置协议）是针对 IPv6 编址方案设计的，为主机分配 IPv6 前缀、IPv6 地址和其他网络配置参数的协议。与其他 IPv6 地址分配方式（手工配置、通过路由器公告消息中的网络前缀无状态自动配置等）相比，DHCPv6 具有以下优点：

- ❑ 不仅可以分配 IPv6 地址，还可以分配 IPv6 前缀，便于全网络的自动配置和管理。
- ❑ 更好地控制地址的分配。通过 DHCPv6 不仅可以记录为主机分配的地址/前缀，还可以为特定主机分配特定的地址/前缀，以便于网络管理。
- ❑ 除了 IPv6 前缀、IPv6 地址外，还可以为主机分配 DNS 服务器、域名等网络配置参数。

1. DHCPv6 网络构成

在 DHCPv6 典型组网一般由三种角色组成，包括 DHCPv6 客户端、DHCPv6 服务器和 DHCPv6 中继，详细介绍如下：

- ❑ **DHCPv6 客户端** 动态获取 IPv6 地址、IPv6 前缀或其他网络配置参数的设备。
- ❑ **DHCPv6 服务器** 负责为 DHCPv6 客户端分配 IPv6 地址、IPv6 前缀和其他网络配置参数的设备。DHCPv6 服务器不仅可以为 DHCPv6 客户端分配 IPv6 地址，还可以为其分配 IPv6 前缀。DHCPv6 服务器为 DHCPv6 客户端分配 IPv6 前缀后，

DHCPv6 客户端向所在网络发送包含该前缀信息的 RA 消息,以便网络内的主机根据该前缀自动配置 IPv6 地址。

- ❑ **DHCPv6 中继** DHCPv6 客户端通过本地链路范围的组播地址与 DHCPv6 服务器通信,以获取 IPv6 地址和其他网络配置参数。如果服务器和客户端不在同一个链路范围内,则需要通过 DHCPv6 中继来转发报文,这样可以避免在每个链路范围内都部署 DHCPv6 服务器,既节省了成本,又便于进行集中管理。

2. DHCPv6 中继工作过程

通过 DHCPv6 中继动态获取 IPv6 地址/前缀和其他网络配置参数的过程中,DHCPv6 客户端与 DHCPv6 服务器的处理方式与不通过 DHCPv6 中继时的处理方式基本相同。DHCPv6 中继的转发过程如下:

- ❑ DHCPv6 客户端向所有 DHCPv6 服务器和中继的组播地址 FF02::1:2 发送请求;
- ❑ DHCPv6 中继接收到请求后,将其封装在 Relay-forward 报文的中继消息选项(Relay Message Option)中,并将 Relay-forward 报文发送给 DHCPv6 服务器;
- ❑ DHCPv6 服务器从 Relay-forward 报文中解析出客户端的请求,为客户端选取 IPv6 地址和其他参数,构造应答消息,将应答消息封装在 Relay-reply 报文的中继消息选项中,并将 Relay-reply 报文发送给 DHCPv6 中继;
- ❑ DHCPv6 中继从 Relay-reply 报文中解析出服务器的应答,转发给 DHCPv6 客户端;
- ❑ DHCPv6 客户端根据 DHCPv6 服务器分配的 IPv6 地址/前缀和其他参数进行网络配置。

5.8 IPv6 中的 DNS 协议

互联网上的应用很多,但大都离不开域名系统(DNS)的支持,域名系统的主要作用是用来进行域名与 IP 地址的转换,即域名解析,比如浏览网站、Email、FTP 等都需要先进行域名解析。IPv6 网络中的 DNS 非常重要,一些 IPv6 的新特性和 DNS 的支持密不可分。下面从 IPv6 DNS 的体系结构、IPv6 的地址解析、IPv6 地址自动配置和即插即用、IPv4 到 IPv6 的过渡等几方面对 IPv6 DNS 进行介绍。

1. IPv6 域名系统的体系结构

IPv6 网络中的 DNS 与 IPv4 的 DNS 在体系结构上是一致的,都采用树型结构的域名空间。IPv4 协议与 IPv6 协议的不同并不意味着需要单独两套 IPv4 DNS 体系和 IPv6 DNS 体系,相反的是,DNS 的体系和域名空间必须是一致的,即 IPv4 和 IPv6 共同拥有统一的域名空间。在 IPv4 到 IPv6 的过渡阶段,域名可以同时对应于多个 IPv4 和 IPv6 的地址。以后随着 IPv6 网络的普及,IPv6 地址将逐渐取代 IPv4 地址。

2. DNS 对 IPv6 地址层次性的支持

IPv6 可聚合全局单播地址是在全局范围内使用的地址,必须进行层次划分及地址聚

合。IPv6 全局单播地址的分配方式如下：顶级地址聚合机构 TLA（即大的 ISP 或地址管理机构）获得大块地址，负责给次级地址聚合机构 NLA（中小规模 ISP）分配地址，NLA 给站点级地址聚合机构 SLA（子网）和网络用户分配地址。IPv6 地址的层次性在 DNS 中通过地址链技术可以得到很好的支持。下面从 DNS 正向地址解析和反向地址解析两方面进行分析。

1) 正向解析

IPv4 的地址正向解析的资源记录是“A”记录。IPv6 地址的正向解析目前有两种资源记录，即，“AAAA”和“A6”记录。其中，“AAAA”较早提出 4，它是对“A”记录的简单扩展，由于 IP 地址由 32 位扩展到 128 位，扩大了四倍，所以资源记录由“A”扩大成四个“A”。“AAAA”用来表示域名和 IPv6 地址的对应关系，并不支持地址的层次性。

“A6”在 RFC2874<5>中提出，它是把一个 IPv6 地址与多个“A6”记录建立联系，每个“A6”记录都只包含了 IPv6 地址的一部分，结合后拼装成一个完整的 IPv6 地址。“A6”记录支持一些“AAAA”所不具备的新特性，如地址聚合，地址更改（Renumber）等。

首先，“A6”记录方式根据 TLA、NLA 和 SLA 的分配层次把 128 位的 IPv6 的地址分解成为若干级的地址前缀和地址后缀，构成了一个地址链。每个地址前缀和地址后缀都是地址链上的一环，一个完整的地址链就组成一个 IPv6 地址。这种思想符合 IPv6 地址的层次结构，从而支持地址聚合。

其次，用户在改变 ISP 时，要随 ISP 改变而改变其拥有的 IPv6 地址。如果手工修改用户子网中所有在 DNS 中注册的地址，是一件非常烦琐的事情。而在用“A6”记录表示的地址链中，只要改变地址前缀对应的 ISP 名字即可，可以大大减少 DNS 中资源记录的修改。并且在地址分配层次中越靠近底层，所需要改动的越少。

2) 反向解析

IPv6 反向解析的记录和 IPv4 一样，是“PTR”，但地址表示形式有两种。一种是用“.”分隔的半字节 16 进制数字格式（Nibble Format），低位地址在前，高位地址在后，域后缀是“IP6.INT.”。另一种是二进制串（Bit-string）格式，以“\<”开头，16 进制地址（无分隔符，高位在前，低位在后）居中，地址后加“>”，域后缀是“IP6.ARPA.”。半字节 16 进制数字格式与“AAAA”对应，是对 IPv4 的简单扩展。二进制串格式与“A6”记录对应，地址也像“A6”一样，可以分成多级地址链表示，每一级的授权用“DNAME”记录。和“A6”一样，二进制串格式也支持地址层次特性。

总之，以地址链形式表示的 IPv6 地址体现了地址的层次性，支持地址聚合和地址更改。但是，由于一次完整的地址解析分成多个步骤进行，需要按照地址的分配层次关系到不同的 DNS 服务器进行查询。所有的查询都成功才能得到完整的解析结果。这势必会延长解析时间，出错的机会也增加。因此，需要进一步改进 DNS 地址链功能，提高域名解析的速度才能为用户提供理想的服务。

3. IPv6 中的即插即用与 DNS

IPv6 协议支持地址自动配置，这是一种即插即用的机制，在没有任何人工干预的情

况下, IPv6 网络接口可以获得链路局部地址、站点局部地址和全局地址等, 并且可以防止地址重复。IPv6 支持无状态地址自动配置和有状态地址自动配置两种方式。

IPv6 结点通过地址自动配置得到 IPv6 地址和网关地址。但是, 地址自动配置中不包括 DNS 服务器的自动配置。如何自动发现提供解析服务的 DNS 服务器也是一个需要解决的问题。正在研究的 DNS 服务器的自动发现的解决方法可以分为无状态和有状态两类。

在无状态的方式下, 需要为子网内部的 DNS 服务器配置站点范围内的任播地址。要进行自动配置的结点以该任播地址为目的地址发送服务器发现请求, 询问 DNS 服务器地址、域名和搜索路径等 DNS 信息。这个请求到达距离最近的 DNS 服务器, 服务器根据请求, 回答 DNS 服务器单播地址、域名和搜索路径等 DNS 信息。结点根据服务器的应答配置本机 DNS 信息, 以后的 DNS 请求就直接用单播地址发送给 DNS 服务器。

另外, 也可以不用站点范围内的任播地址, 而采用站点范围内的多播地址或链路多播地址等。还可以一直用站点范围内的任播地址作为 DNS 服务器的地址, 所有的 DNS 解析请求都发送给这个任播地址。距离最近的 DNS 服务器负责解析这个请求, 得到解析结果后把结果返回请求结点, 而不像上述做法是把 DNS 服务器单播地址、域名和搜索路径等 DNS 信息告诉结点。从网络扩展性, 安全性, 实用性等多方面综合考虑, 第一种采用站点范围内的任播地址作为 DNS 服务器地址的方式相对较好。

在有状态的 DNS 服务器发现方式下, 是通过类似 DHCP 这样的服务器把 DNS 服务器地址、域名和搜索路径等 DNS 信息告诉结点。当然, 这样做需要额外的服务器。

5.9 IPv6 路由协议及安全

IPv6 的概念现在已并不陌生, 下面仅对 IPv6 路由协议在安全问题上, 从以下三个方面做一个深入的研究。

1. 协议安全

在协议安全层面上, IPv6 路由协议全面支持认证头 (AH) 认证和封装安全有效负荷 (ESP) 信息安全封装扩展头。AH 认证支持 hmac_md5_96、hmac_sha_1_96 认证加密算法, ESP 封装支持 DES_CBC、3DES_CBC 以及 Null 等三种算法。

2. 网络安全

IPv6 路由协议的网络安全包括以下四个方面, 详细介绍如下:

- ❑ 端到端的安全保证。在两端主机上对报文进行 IPSec 封装, 中间路由器实现对有 IPSec 扩展头的 IPV6 报文进行透传, 从而实现端到端的安全。
- ❑ 对内部网络的保密。当内部主机与因特网上其他主机进行通信时, 为了保证内部网络的安全, 可以通过配置的 IPSec 网关实现。因为 IPSec 作为 IPV6 路由协议的扩展报头不能被中间路由器而只能被目的结点解析处理, 因此 IPSec 网关可以通过 IPSec 隧道的方式实现, 也可以通过 IPV6 路由协议扩展头中提供的路由头和逐跳选项头结合应用层网关技术来实现。后者的实现方式更加灵活, 有利

于提供完善的内部网络安全，但是比较复杂。

- ❑ 通过安全隧道构建安全的 VPN。此处的 VPN 是通过 IPV6 路由协议的 IPSec 隧道实现的。在路由器之间建立 IPSec 的安全隧道，构成安全的 VPN 是最常用的安全网络组建方式。IPSec 网关的路由器实际上就是 IPSec 隧道的终点和起点，为了满足转发性能的要求，该路由器需要专用的加密板卡。
- ❑ 通过隧道嵌套实现网络安全。通过隧道嵌套的方式可以获得多重的安全保护。当配置了 IPSec 的主机通过安全隧道接入到配置了 IPSec 网关的路由器，并且该路由器作为外部隧道的终结点将外部隧道封装剥除时，嵌套的内部安全隧道就构成了对内部网络的安全隔离。

3. 其他安全保障

IPv6 路由协议的 IPSec 为网络数据和信息内容的有效性、一致性以及完整性提供了保证，但是数据网络的安全威胁是多层面的，它们分布在物理层、数据链路层、网络层、传输层和应用层等各个部分。

对于物理层的安全隐患，可以通过配置冗余设备、冗余线路、安全供电、保障电磁兼容环境以及加强安全管理来防护。

对于物理层以上层面的安全隐患，可以采用以下防护手段：通过诸如 AAA、TACACS+、RADIUS 等安全访问控制协议控制用户对网络的访问权限来防止针对应用层的攻击；通过 MAC 地址和 IP 地址绑定、限制每端口的 MAC 地址使用数量、设立每端口广播包流量门限、使用基于端口和 VLAN 的 ACL、建立安全用户隧道等来防范针对二层网络的攻击；通过进行路由过滤、对路由信息的加密和认证、定向组播控制、提高路由收敛速度、减轻路由振荡的影响等措施来加强三层网络的安全性。

路由器和交换机对 IPSec 的完善支持保证了网络数据和信息内容的有效性、一致性以及完整性，并且为网络安全提供了诸多解决办法。

5.10 IPv6 过渡技术

IPv6 是 IPv4 的替代协议，是下一代互联网的核心和基础协议。从 1994 年 IETF 批准 RFC1752 至今，IPv6 已经走过了 13 年的发展历程。今天，IPv6 已经站在了大规模商用的门坎上，业界公认，IPv6 将全面推动以移动通信和数字家电为代表的一系列新业务和新技术的发展，为社会全面信息化提供可靠的基础协议。

IPv6 的实际部署并不能一蹴而就，由于 IPv4 是目前 Internet 上的统治性协议，现阶段 IP 网络上的绝大多数应用都是基于 IPv4 的，基于 IPv6 的应用需要逐步开发，客户群需要逐渐培养，而且基于 IPv4 技术的累积投资巨大，因此决定了 IPv4 向 IPv6 的过渡将是一个漫长的过程。如何在整个过渡期间做到业务平滑迁移，如何保护运营商既有投资和 IPv4 用户既有利益，是电信运营商和设备供应商需要研究的共同课题。

5.10.1 从 IPv4 向 IPv6 过渡的三个阶段

市场需求是网络演进的根本动力，而市场需求则来源于最终用户的业务体验，IPv6

商用也必然遵循这一规律。从业务体验方面来看,首先对 IPv6 有现实需求的是 3G、WiMax、WiFi 等移动宽带业务,其次是 IPTV、视频会议、VoIP 等对资源要求高的业务,待这些业务充分发展、IPv6 多业务承载网逐步成熟后,IPv6 技术将会逐渐进入智能家居、远程办公等人们的日常生活领域,最终过渡到全 IPv6 的下一代互联网。从网络部署方面来看,首先部署 IPv6 的应该是业务量需求大的核心城市,以及大学、科研院所、政府机关、大企业等 VIP 客户,待网络层次成型、产业链逐渐成熟后,再逐步将 IPv6 引入业务量稀疏地区和普通用户。

在 IPv6 向 IPv4 过渡的初期,基于 IPv6 的业务还带有相当程度的试验和探索性质,Internet 上的绝大部分应用仍然是基于 IPv4 的,Internet 的主体也仍然采用 IPv4,IPv6 网络仅在局部构成中小型网络。业界形象地将这个阶段称为“IPv6 孤岛,IPv4 海洋”。这一阶段的问题首先是使各个相对孤立的 IPv6 网络之间能够相互通信;其次,IPv6 孤岛需要访问的资源大部分仍然处在基于 IPv4 的 Internet 主体网络中,需要解决 IPv6 网络中的主机访问 IPv4 网络的问题;而且此时 IPv6 的地址分配还处在 IPv4 地址的阴影当中,大量 IPv6 特定网段的地址或者 IPv4 向 IPv6 的映射网段地址构成 IPv6 路由表的重要组成部分。

当基于 IPv6 的业务普遍实现大规模应用后,从 IPv4 到 IPv6 的过渡就进入了 IPv4 与 IPv6 并存的中期阶段。在这一阶段,IPv6 网络已经形成规模,构成了独立于 IPv4 网络的从接入网、驻地网到核心网的完整网络结构。这一阶段最重要的任务是如何解决 IPv4 和 IPv6 网络中的主机和资源互访过程中出现的种种问题,其次是如何将 IPv4 网络中的结点升级到 IPv6,以及如何将 IPv4 中的网络资源和业务、应用迁移到 IPv6 网络中去。

最后,当 IPv4 网络中的绝大部分业务和应用都迁移到 IPv6 网络以后,IPv4 到 IPv6 的过渡就进入了后期阶段。这一阶段与初期阶段相反,网络的状况是“IPv4 孤岛,IPv6 海洋”。要解决的问题变成了如何连接互不连接的 IPv4 网络,如何让 IPv4 网络中的主机能够访问 IPv6 网络中的资源。

目前,IPv6 的实际部署处于存在着若干 IPv6 孤岛,仅有少量类似中国的 CNGI、日本的 e-Japan 这样的纯 IPv6 试验性骨干网,同时 Internet 的主体仍然使用 IPv4 的初期阶段。在此阶段内,运营商和方案设备提供商们最为关注的是从 IPv4 到 IPv6 的过渡技术,以及现有电信设备对 IPv6 的支持和升级能力。

5.10.2 IPv4 to IPv6 过渡技术

在 IPv4 到 IPv6 过渡的初期阶段,可以看到有三类过渡需求:第一,需要有一些网络结点能够同时支持 IPv4 和 IPv6,特别是连接 IPv4 和 IPv6 网络的网关设备必须具有这种能力;第二,必须使 IPv6 孤岛网络能够穿越通过基于 IPv4 的网络主体实现互连互通;第三,IPv4 和 IPv6 网络之间必须能够相互访问对方网络中的资源。对应于这三类需求,可以分别采用双栈技术、隧道技术和互通技术来应对。

1. 双栈技术

“双栈”是指单个结点同时支持 IPv4 和 IPv6 协议栈,这样的结点既可以基于 IPv4

协议直接与 IPv4 结点通信,也可以基于 IPv6 协议直接与 IPv6 结点通信,因此它可以作为 IPv4 网络和 IPv6 网络之间的衔接点。很明显,无论是隧道技术中隧道的封装和解封装设备,还是互通技术中的 NAT-PT (Network Address Translation-Protocol Translator, NAT 协议转换器) 设备或者 ALG (Application Level Gateway, 应用层网关) 设备,本身都必须是双栈设备,因此双栈技术是各种过渡技术的基础。

由于双栈设备需要同时运行 IPv4 和 IPv6 两个协议栈,因此需要同时保存两套命令集,同时计算、维护与存储两套表项,对网关设备而言,还需要对两个协议栈进行报文转换和重封装,所以运行双栈的设备明显要比只运行一个协议栈的设备负担更重,对设备的性能要求更高,维护和优化的工作也复杂。

双栈技术除了用在 IPv4 和 IPv6 间的网关设备上以外,还可以用来组建小型的 IPv4 和 IPv6 混合型网络。在这种网络中,所有的网络结点都是双栈主机,都可以直接访问 IPv4 或者 IPv6 网络中的资源,这样的双栈网络不存在互通问题,有一定的方便性。但是它需要为网络中的每个 IPv6 结点同时分配一个 IPv4 地址,不但仍然受制于 IPv4 地址资源不足的问题,而且对每个结点的性能要求都比较高,势必会增加用户建网和维护的成本,因而仅适合于 IPv4 to IPv6 过渡的初期或者后期,在 IPv6 或者 IPv4 的小型孤岛上组建这种网络。

2. 隧道技术

隧道技术用来将不直接相连的 IPv6 或者 IPv4 孤岛互相连接起来,这种连接可能有两种情况:一种是隧道的两端是 IPv6 孤岛,需要穿越 IPv4 海洋进行连接,另一种是隧道的两端是 IPv4 孤岛,需要穿越 IPv6 海洋进行连接,无论哪种情况,都需要在隧道的入口对报文进行重新封装,然后把封装过的报文通过中间网络转发到隧道出口,在隧道的出口对报文进行解封装后,再将恢复后的报文转发到目的地。

在实际使用中,常见的隧道技术有 GRE 隧道、IPv6 over IPv4 手工配置隧道、6PE/6VPE 隧道、6to4 隧道、ISATAP 隧道等几种,详细介绍如下:

1) GRE 隧道

通用路由封装协议 GRE (Generic Routing Encapsulation) 是一种常用的隧道封装协议。使用 GRE 封装 IPv6 报文时,整个 IPv6 数据报文都在隧道的入口路由器上作为 GRE 的载荷被封装起来,待传递到隧道的出口路由器上,再解除 GRE 封装,将恢复后的 IPv6 报文在 IPv6 网络中继续转发。在整个转发过程中, GRE 隧道对 IPv6 网络来说相当于一条物理链路,中间转发 GRE 报文的路由器对 IPv6 报文和 IPv6 网络的存在一无所知,因此 GRE 隧道仅要求隧道的入口和出口路由器是双栈路由器、可以较好地支持 GRE 技术即可。

GRE 隧道技术成熟,对除了入口和出口路由器以外的其他设备没有双栈要求;且 GRE 协议本身安全性较好;在未来“IPv4 孤岛, IPv6 海洋”时期, GRE 隧道仍可以用来封装连接分离的 IPv4 网络,技术寿命也很长。但是,由于 GRE 仅能提供点对点连接,因此它的使用范围较为狭窄。

2) IPv6 over IPv4 手工配置隧道

手工配置隧道直接使用 IPv4 封装 IPv6 报文。隧道入口的路由器从 IPv6 侧收到一个

IPv6 报文后, 根据 IPv6 报文的地址查找 IPv6 转发表, 如果该报文下一跳地址为隧道逻辑接口, 则将该报文根据隧道配置的源和目的 IPv4 地址, 将 IPv6 的报文封装到 IPv4 的报文中。封装后的 IPv4 报文的源地址和目的地址分别是隧道入口和出口的 IPv4 地址, 并用 IPv4 报头的“协议”字段标识其负载为 IPv6 报文。报文通过 IPv4 网络转发到隧道的出口路由器, 在此再将 IPv6 分组取出转发给目的 IPv6 结点。

手工隧道技术原理简单, 技术成熟稳定。但是由于是纯手工配置, 大量使用时带来的维护量较大, 可扩展性不好。而且即使与同为手工配置的 GRE 隧道相比, 由于 IPv4 报文本不提供安全认证和报文验证, 安全性也不如 GRE 隧道。

3) 6PE/6VPE 隧道

6PE 是基于 MPLS 的隧道技术, 其核心思想是借助成熟的 BGP MPLS VPN 技术平台实现在启用 MPLS 的 IPv4 骨干网上传输 IPv6 数据报文, 为 IPv6 网络孤岛提供互连能力。6PE 隧道技术的 VPN 路由发布和报文转发原理与常见的 IPv4 骨干网上的 MPLS L3 VPN 类似。

6PE 路由器与同处于 IPv6 网络内的 CE 路由器之间通过 IPv6 IGP 路由协议交换路由信息。6PE 路由器为 IPv6 路由加上私网标签 (由 MP-IBGP 协议随机自动生成, 被传递到对端 6PE 并保留到转发表中), 并将此路由的“Next-hop”属性更改为映射后的自身 Loopback 地址 (为与 CE 的路由保持相同的地址族, 6PE 的 IPv4 Loopback 地址被映射成 IPv6 地址, 地址形式为: “::FFFF:IPv4-Address”), 然后加上 MPLS 外层标签通过 MPLS LSP 隧道发布给对端 6PE 设备, 对端 6PE 接收并保留私网标签, 然后将路由的下一跳属性改变为映射后的自身 Loopback 地址, 再以 IPv6 普通路由的形式发布给自己一侧的 IPv6 CE 设备, 两个 IPv6 网络的路由通过这种方式就完成了交互。

IPv6 报文转发时, CE 设备根据报文的地址发送给 6PE 设备, 6PE 设备在 IPv6 路由表中进行查找, 得到该数据报文对应 IPv6 路由的下一跳地址 (即对端 6PE 的 Loopback 地址) 和私网标签, 在 IPv6 报文外先封装私网标签, 再根据 MPLS LSP 标签转发表中与下一跳对应的标签封装外层标签, 然后将 MPLS 报文通过 LSP 上各个 P 路由器逐跳转发, 倒数第二跳 P 路由器弹出外层标签并继续转发给相应 6PE 路由器, 在 6PE 路由器上根据内层标签将 IPv6 数据包转发至目的 CE 设备。

4) ISATAP 自动隧道

ISATAP (Intra-Site Automatic Tunnel Addressing Protocol, 内部隧道自动地址协议) 自动隧道技术主要用于 IPv4 网络中的双栈主机通过 ISATAP 隧道访问 IPv6 网络的场景, 但它不仅是一种自动隧道技术, 而且同时可以进行地址自动配置。

配置了 ISATAP 隧道以后, IPv6 网络将作为隧道外层封装的 IPv4 网络看作一个 NBMA (Non-Broadcast Multiple Access, 非广播多接入) 链路, 这是 ISATAP 最大的特点。

5) 6to4 自动隧道和 6to4 中继

6to4 隧道是为了解决隧道自动配置问题而设计的, 虽然它也是使用 IPv4 报文来重封装 IPv6 报文, 但是隧道的源和目的 IPv4 地址不需要手工配置, 而是嵌在 6to4 结点的 IPv6 地址内部的。

6to4 结点 (可能是路由设备, 也可能是单个主机) 的 IPv6 地址前缀是统一的 2002::/16

地址空间。

每个 6to4 结点至少使用一个分配给 6to4 设备的全球 IPv4 单播地址，这个地址连接在 2002::/16 前缀后面，成为 2002:IPv4-address::/48 的特定结点前缀。

每个 6to4 结点前缀后面的 16bits 用于它所连接的 IPv6 网络内的子网划分，即可以划分出 $2^{16}=65\,536$ 个 IPv6 子网，每个子网拥有 $2^{64}-2$ 个 IPv6 地址。

标准 (RFC3068) 规定，6to4 结点本身必须是 IPv6 网络的边缘结点 (这个 IPv6 网络里可以只有 6to4 结点自己)。由于 6to4 结点的 IPv6 地址是由 IPv4 地址映射生成的，所以 6to4 结点间在 IPv4 网络内不必以任何形式发布 IPv6 路由信息，仅使用 IPv4 全局路由就可以保证彼此间的可达性。

3. 互通技术

互通技术是为了解决 IPv4 和 IPv6 网络内主机和资源互访的问题而提出的。其实双栈技术本身也是一种最简单的互通技术，但是由于 IPv4 地址资源的限制，根本不可能为每一台使用 IPv6 的主机同时分配一个 IPv4 地址，因此还必须开发其他技术，以便那些只有 IPv6 地址的主机也能访问 IPv4 网络中的资源。当前最常用的互通技术是 NAT-PT 技术，而 NAT-PT 技术又可以分为静态 NAT-PT 和结合 ALG 技术的动态 NAT-PT。

1) 静态 NAT-PT 技术

静态 NAT-PT 技术是在 NAT-PT 网关静态配置 IPv6 和 IPv4 地址的绑定关系。当 IPv4 主机和 IPv6 主机报文互通，NAT-PT 网关根据配置的绑定关系进行转换，且任何一侧主机都可以主动向另一侧发起连接。

静态 NAT-PT 技术原理简单，适合永久在线或需要提供稳定连接的应用场合。但是当有很多主机需要转换时，配置和维护显得复杂，而且消耗较多的 IPv4 地址，因而不适合大规模网络中使用。

2) 动态 NAT-PT 技术

动态 NAT-PT 技术改进了静态 NAT-PT 消耗大量 IPv4 地址地缺点，采用动态地址映射和上层协议映射的方法，使大量的 IPv6 地址可以通过很少的 IPv4 地址进行转换。

NAT-PT 网关向 IPv6 网络通告一个 96 位的地址前缀，当 IPv4 网络中的主机访问 IPv6 网络时，这个地址前缀加上 32 位的 IPv4 地址就成为转换后的 IPv6 地址。

动态 NAT-PT 技术应用中，若从 IPv4 端首先发起连接，IPv4 主机无从知道 IPv6 主机随机映射后的 IPv4 地址或上层协议端口，连接无法进行。ALG (Application Level Gateway) 即应用层网关技术配合动态 NAT-PT 技术可以解决这个问题，它包括 DNS-ALG, FTP-ALG, SIP-ALG 等多种应用

PCB 首先向它的 DNS 服务器请求 www.A.com.cn 这个域名所对应的 IPv4 地址，IPv4 的 DNS 服务器发现没有此记录，于是转发这个 DNS 请求给 IPv6 的 DNS 服务器，请求报文的源和目的地址分别为 1.3.3.5 和 1.1.1.9，报文中包含类型为 A 的查询报文和需要解析的名字 A。请求报文在 NAT-PT 网关把报文头部源和目的地址转换为 10::1.3.3.5 和 2::3。同时，NAT-PT 网关中的 DNS-ALG 将报文的类型由 A 转换成 A6，然后将此报文向 IPv6 DNS 转发。

IPv6 DNS 收到报文后，查询自己的记录表，解析出 A 的 IPv6 地址为 1::3，于是向

IPv4 DNS 做出回应。回应报文的源和目的地址在 NAT-PT 网关被转换为 1.1.1.9 和 1.3.3.5。同时, DNS-ALG 将其中的 DNS 应答类型改为 A, 并从 IPv4 地址池中分配一个地址 1.4.2.6, 替换应答中的 IPv6 地址 1::3, 并记录二者之间的映射信息。PCB 收到 IPv4 DNS 应答消息, 从而知道 PCA 的 IPv4 地址为 1.4.2.6, PCB 就可以主动发起到 PCA 的连接了。

动态 NAT-PT 技术仅使用很少的 IPv4 地址, 在不修改 IPv4 网络的情况下, 就可实现纯 IPv4 网络与纯 IPv6 网络的相互访问, 是一个很优秀的 IPv4 与 IPv6 网络互通的过渡技术。但是动态 NAT-PT 技术比较复杂, 对 NAT-PT 设备的操作系统设计水平、硬件处理能力及系统稳定性提出了很高的要求。

5.11 扩展练习

1. 启用 IPv6 协议

IPv6 是 Internet Protocol Version 6 的缩写, 是一种互联网协议, 它可以提高网速。但是, 由于一些优化工具把 IP Helper 服务禁用, 所以要使用该协议提高网速, 首先应启用 IPv6 的服务。本例介绍启用 IPv6 协议, 步骤如下:

1 单击【开始】按钮, 执行【控制面板】命令, 在弹出的窗口中, 双击【管理工具】图标, 如图 5-21 所示。



图 5-21 双击【管理工具】图标

2 在【管理工具】窗口的右侧窗格中, 双击【服务】图标, 如图 5-22 所示。



图 5-22 双击【服务】图标

3 在【服务】窗口中, 双击 IPv6 Helper Service 选项, 如图 5-23 所示。

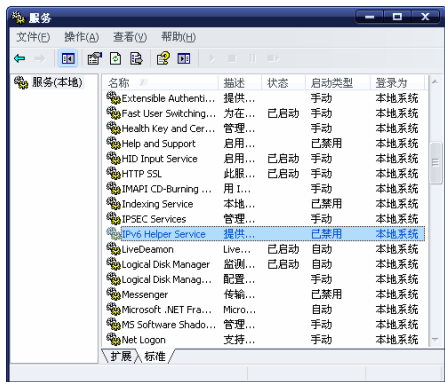


图 5-23 双击 IPv6 选项

4 在【IPv6 Helper Service 的属性(本地计算机)】窗口中, 设置【启动类型】为【自动】, 如图 5-24 所示。



图 5-24 设置 IPv6 启动类型

5 单击对话框中的【应用】按钮，然后，单击【服务状态】栏下面的【启动】按钮，如图 5-25 所示。



图 5-25 单击【启动】按钮

6 单击【启动】按钮后，将出现【服务控制】对话框，如图 5-26 所示，即可启用 IPv6 协议。



图 5-26 启用 IPv6 协议



单击【开始】按钮，执行【运行】命令，在弹出的对话框中，输入 cmd 命令。然后，在弹出的窗口中，输入 IPv6 install 命令，可安装 IPv6 协议。如果输入 IPv6 uninstall 命令，即可卸载该协议。

2. 使用 Ping 命令

Ping 命令是 Windows 系统自带的一个可执行命令，它是一个通信协议，是 TCP/IP 协议的一部分。用它可以检查网络是否是连通的，也可以用它帮助用户分析判断网络故障。本例来介绍 Ping 命令的用法，步骤如下：

1 单击【开始】按钮，执行【运行】命令，在弹出的对话框中，输入 cmd 命令，如图 5-27 所示。

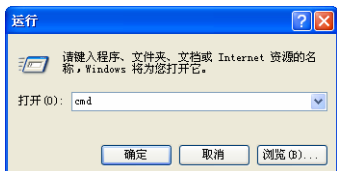


图 5-27 输入 cmd 命令

2 在弹出的 MS-DOS 窗口中，输入 ipconfig 命令，即可查出本机的 IP 地址及网关地址，如图 5-28 所示。

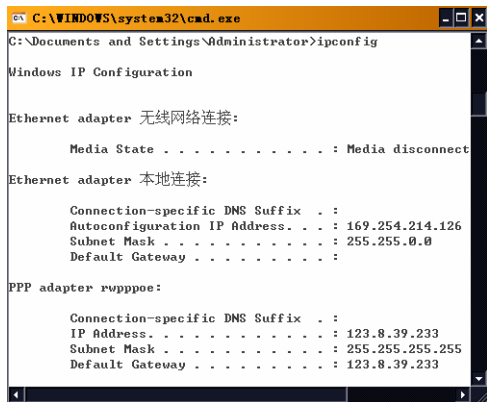


图 5-28 查询 IP 地址

3 在 MS-DOS 窗口中，输入 ping 169.254.214.126 命令，如果网卡配置没有问题，则应显示如图 5-29 所示的内容。

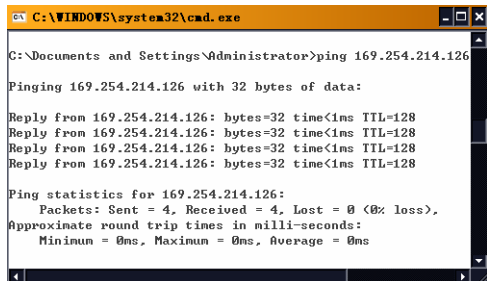


图 5-29 检测网卡配置

4 在 MS-DOS 窗口中，输入 ping 123.8.39.233 命令，如果显示如图 5-30 所示的内容，则表明局域网中的网关路由器正在正常运行。

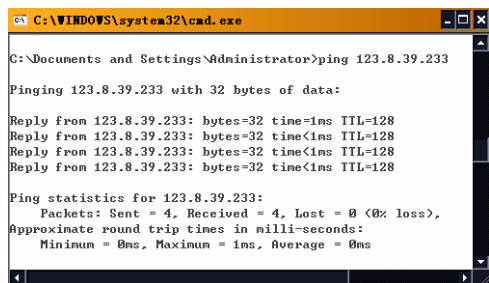


图 5-30 检测网关路由器



若输入 ping 命令后，在 MS-DOS 窗口中显示内容为：Request timed out，则表明网卡安装或配置有问题。将网线断开再次执行此命令，如果显示正常，则说明本机使用的 IP 地址可能与另一台正在使用的机器 IP 地址重复了。如果仍然不正常，则表明本机网卡安装或配置有问题，需继续检查相关网络配置。