

高等学校计算机网络与通信技术教材

发展中的 通信网络新技术

鲁士文 编著

清华大学出版社 ● 北京交通大学出版社



高等学校计算机网络与通信技术教材

发展中的通信网络新技术

鲁士文 编著

清华大学出版社

北京交通大学出版社

·北京·

内 容 简 介

本书针对高等院校和研究机构的研究生学习通信网络课程和从事相关研究工作的实际需求,深入介绍和讨论了近年来在通信网络领域发展起来的多种革新理论和关键技术,主要内容包括 IPv6、实时流传输、服务质量保证、网络安全、组播和任播、多协议标记交换、多媒体通信、会话启动协议、移动 IP、拥塞控制、社区接入网络、无线网络、自组织移动网络及覆盖网络。

本书融原理、技术和发展为一体,注重介绍新技术和系统设计方法,以提高读者从事研究工作和解决实际问题的能力为主要目标,可供高等学校和科研单位信息技术相关专业的研究生用作通信网络课程的参考教材或开展相关研究课题的参考资料。本书也可供从事通信网络研究和应用开发的人员作为了解通信网络新发展或知识更新的一个媒介。

版权所有,翻印必究。举报电话:010-62782989 13501256678 13801310933

本书封面贴有清华大学出版社防伪标签,无标签者不得销售。

本书防伪标签采用特殊防伪技术,用户可通过在图案表面涂抹清水,图案消失,水干后图案复现;或将表面膜揭下,放在白纸上用彩笔涂抹,图案在白纸上再现的方法识别真伪。

图书在版编目(CIP)数据

发展中的通信网络新技术/鲁士文编著. —北京:清华大学出版社;北京交通大学出版社, 2006.3

(高等学校计算机网络与通信技术教材)

ISBN 7-81082-649-2

I. 发… II. 鲁… III. 通信网-新技术-高等学校-教材 IV. TN915-39

中国版本图书馆 CIP 数据核字 (2005) 第 130323 号

责任编辑:谭文芳 特邀编辑:肖 融

出版者:清华大学出版社 邮编:100084 电话:010-62776969 <http://www.tup.com.cn>

北京交通大学出版社 邮编:100044 电话:010-51686414 <http://press.bjtu.edu.cn>

印刷者:北京东光印刷厂

发行者:新华书店总店北京发行所

开 本:185×260 印张:17.25 字数:435 千字

版 次:2006 年 3 月第 1 版 2006 年 3 月第 1 次印刷

书 号:7-81082-649-2/TN·44

印 数:1~4 000 册 定价:26.00 元

本书如有质量问题,请向北京交通大学出版社质监组反映。对您的意见和批评,我们表示欢迎和感谢。

投诉电话:010-51686043, 51686008; 传真:010-62225406; E-mail: press@center.bjtu.edu.cn。

前 言

当今世界正经历着一场信息革命。从信息的产生、存储、传输到处理等各个方面都在发生着巨大的变化。要想充分利用这一场革命所带来的科技成果并发展我国具有自主知识产权的新技术和新产品，就需要我们的研究生和广大青年科研工作者不断地学习当代通信网络系统中有关基础设施方面的新知识，了解当前通信网络发展趋势，深入理解其体系结构和运行机制，掌握相关理论和算法，熟悉组网技术和网络系统软、硬件的设计方法，培养使用、管理网络 and 开发网络通信系统与网络应用程序的能力，提高解决实际问题能力和创新能力。

通信网络是一个领域广泛、发展迅速的主题。它涉及电话系统、广播电视系统和数据网络及它们在技术上互相融合的方式。正是在这样一个与社会组织和民众生活都息息相关的领域，新的思想和概念以惊人的速度不断涌现，令人耳目一新的产品更是层出不穷。本书针对高等院校和研究机构的研究生学习通信网络课程和从事相关研究工作的实际需求，深入介绍和讨论近年来在通信网络领域发展起来的多种革新理论和关键技术，主要包括 IPv6、实时流传输、服务质量保证、网络安全、组播和任播、多协议标记交换、多媒体通信、会话启动协议、移动 IP、拥塞控制、社区接入网络、无线网络、自组织移动网络及覆盖网络。它是作者在多年从事研究生网络课程教学和开展与数据通信相关的科研项目的基础上编写的，融原理、技术和发展为一体，注重介绍新技术和系统设计方法，以提高读者从事研究工作和解决实际问题的能力为主要目标。

本书编写的目的主要是供高等学校和科研单位信息技术相关专业的研究生用作通信网络课程的教材，也可用作开展相关研究课题的参考资料。由于本书涉及面较广，且内容有较大的深度，读者在使用的过程中可根据具体的教学背景和专业方向进行选用。另外本书也可供从事通信网络研究和应用开发的人员作为了解通信网络的新发展或知识更新的一个媒介。限于时间与水平，不当之处欢迎批评指正。

作 者

2006 年 1 月于中科院计算所

目 录

第 1 章 下一代因特网协议 IPv6	1
1.1 IPv6 分组格式	2
1.2 扩展头	5
1.2.1 路由选择头	6
1.2.2 逐跳选项头	7
1.2.3 分割头	9
1.2.4 目的地选项头	10
1.2.5 扩展头顺序	10
1.3 IPv6 编址	11
1.3.1 IPv6 地址的表示	11
1.3.2 初始的分配	12
1.3.3 聚合全局单播地址	14
1.3.4 特殊地址格式	15
1.4 自动配置	16
1.4.1 链路本地地址	17
1.4.2 无状态自动配置	18
1.4.3 重复地址检测	19
1.4.4 有状态配置	20
1.4.5 地址的生命期	20
1.4.6 动态主机配置	21
1.5 邻居发现过程	25
1.5.1 基本的算法	25
1.5.2 从路由器得到信息	27
第 2 章 IPv6 对因特网路由选择和高层协议的影响	29
2.1 IPv6 路由选择	29
2.1.1 域间路由选择	30
2.1.2 域内路由选择	36
2.2 ICMP 的演变	41
2.3 传输层分组检验和	44
2.4 在域名服务中的 IPv6	46
2.5 从 IPv4 向 IPv6 过渡的策略	46
第 3 章 IPv6 对实时通信和安全性的支持	54
3.1 IPv6 对实时通信和流传输的支持	54

3.1.1	流标记	54
3.1.2	支持预留	55
3.1.3	等级式编码和优先级	57
3.1.4	实时传输协议	61
3.1.5	资源预留协议	63
3.2	IPv6 对网络安全性的支持	66
3.2.1	加密和身份验证	66
3.2.2	密钥分布	70
第 4 章	组播和任播	72
4.1	组播的概念	72
4.2	链路状态组播	72
4.3	距离向量组播	74
4.3.1	反向通路广播	74
4.3.2	反向通路组播	74
4.4	协议无关的组播	75
4.5	基于 IPv4 实现的组播	77
4.5.1	IP 组播地址	78
4.5.2	扩展 IP 处理组播的功能	79
4.5.3	IGMP 协议	79
4.5.4	组成员状态变迁	80
4.5.5	传播路由信息	81
4.5.6	Mrouted 程序	82
4.5.7	因特网组播主干网 Mbone	83
4.6	采用 IPv6 协议的因特网组播	84
4.6.1	地址结构	84
4.6.2	组标识符的结构	85
4.6.3	组管理	87
4.6.4	组播路由选择	88
4.7	任播	88
第 5 章	多协议标记交换	91
5.1	多协议标记交换的基本思想	91
5.2	数据报和面向连接的网络技术	93
5.3	体系结构	96
5.4	标记分配协议 LDP	99
5.5	MPLS 的应用	100
5.5.1	基于目的地的转发	100
5.5.2	显式路由选择	103
5.5.3	虚拟专用网络和隧道	104
第 6 章	资源预留和实时协议	109

6.1	实时应用分类	109
6.2	资源预留协议	111
6.2.1	服务类别	111
6.2.2	实现机制	111
6.2.3	流描述	112
6.2.4	许可控制	113
6.2.5	预留协议	114
6.2.6	分组分类和调度	121
6.3	RTP 协议	122
6.3.1	基本需求	122
6.3.2	协议机制	123
6.3.3	RTP 头格式	123
6.3.4	信息源的同步	125
6.3.5	RTCP 协议	125
6.4	SIP 协议	129
6.4.1	功能特征	129
6.4.2	体系结构	130
6.4.3	SIP 信令	131
6.4.4	SIP 命令和响应	133
6.4.5	SIP 报文格式	135
第 7 章	移动 IP 和三角路由问题	140
7.1	移动主机的路由选择	140
7.2	移动 IP 的功能设计和工作过程	142
7.2.1	功能特征	143
7.2.2	必要条件和设计目标	143
7.2.3	功能实体及其驻留位置	144
7.2.4	移动 IP 工作过程小结	146
7.3	代理发现	146
7.4	移动检测和移动登记	148
7.4.1	登记功能	149
7.4.2	登记协议	149
7.4.3	登记操作	151
7.5	路由操作和三角路由问题	153
第 8 章	实现 QoS 的技术途径和网络拥塞控制	157
8.1	实现 QoS 的途径	157
8.1.1	过度建设	157
8.1.2	优先级	157
8.1.3	队列	158
8.1.4	拥塞控制与避免	158

8.1.5	传输整形	158
8.2	QoS 的技术进展	159
8.2.1	MPLS 对 QoS 的支持	159
8.2.2	QoS 路由技术	159
8.2.3	IPv6 对 QoS 的支持	159
8.3	支持 QoS 的现有方法类型	160
8.4	拥塞控制	160
8.4.1	开环控制	161
8.4.2	闭环控制	165
8.5	无线 TCP 及其拥塞问题	168
8.6	用于千兆位网络的传输协议	169
第 9 章	社区接入网络和多媒体应用系统	171
9.1	社区接入网	171
9.1.1	编码/调制技术	171
9.1.2	带宽条件和服务需求	173
9.1.3	非对称数字用户线	174
9.1.4	有线电视网 CATV 和混合光纤同轴系统 HFC	178
9.1.5	光纤到宅边	183
9.1.6	光纤到户	184
9.1.7	无线本地回路	188
9.2	视频点播	190
9.2.1	视频服务器	191
9.2.2	分布式网络	194
9.2.3	机顶盒	196
第 10 章	公用无线网络	198
10.1	无线介质的特征	199
10.2	无线通道的容量限制	201
10.3	数字蜂窝无线电的通道分配	202
10.3.1	全球移动通信系统	202
10.3.2	蜂窝数字分组数据系统	204
10.3.3	码分多址	206
10.4	3G 和 4G	208
第 11 章	无线局域网	211
11.1	无线局域网的组成	211
11.2	无线局域网的协议体系	214
11.3	无线局域网中的扩展频谱技术	219
11.3.1	跳频	220
11.3.2	直接序列	222
11.4	IEEE 802.11 MAC 协议	224

11.4.1	IEEE 802.11 帧结构	226
11.4.2	分布式协调功能 DCF	227
11.4.3	点协调功能 PCF	231
11.5	物理层	232
第 12 章	自组织移动网络	234
12.1	基本概念	234
12.2	自组织移动网络面临的挑战	235
12.3	自组织移动网络的无线介质访问协议	237
12.3.1	带邀请的 MACA 协议	237
12.3.2	感知功率的带信令的多址访问协议	238
12.3.3	双忙音多址访问协议	239
12.3.4	采用简化握手过程的介质访问协议	239
12.4	自组织移动网络的按需距离向量路由选择	241
12.5	基于关联的长活路由选择	242
12.5.1	路由发现阶段	244
12.5.2	路由重构阶段	246
12.5.3	路由删除阶段	248
12.6	ABR 分组头和相关列表	248
第 13 章	覆盖网络	251
13.1	因特网的僵化	252
13.2	路由覆盖	252
13.2.1	IP 试验基地	253
13.2.2	端系统组播	253
13.3	对等网络	256
13.3.1	Gnutella	256
13.3.2	结构性覆盖	257
参考文献		262

第 1 章 下一代因特网协议 IPv6

当前采用的 IP 协议是它的第 4 版 (IPv4), IPv5 的称号被赋给了一个实验的称为流协议的面向连接的因特网协议。现在人们普遍意识到, 或早或晚, IPv4 最终要被一个称为 IPv6 的新协议替代。

对 IP 新版本的需求首先是由在 IPv4 中 32 位地址段的限制引起的。子网划分和无类别域间路由选择有助于控制 Internet 地址空间消耗的速度, 也有助于控制在 Internet 路由器中所需要的路由表信息的增长。然而, 当 Internet 发展到一定程度时, 这些技术就会变得无能为力了。特别地, 人们不可能取得 100% 的地址利用率, 因此, 在远不足 40 亿台主机连到 Internet 之前, 地址空间就要被用尽。即使能够使用所有的 40 亿个地址, 如果要把 IP 地址分配给有线电视的机顶盒, 或者分配给电子仪表, 那么这样多的地址也是不够用的。所有这些可能性都表明, 人们最终肯定需要比 32 位所提供的要大得多的地址空间。

由于 IP 地址运载在每一个 IP 分组的头部, 增加 IP 地址的尺寸势必要改变 IP 分组头。这就意味着要建立一个新的 IP 版本, 因此在 Internet 中的每个主机和路由器都要采用新的软件。这显然不是一件小事, 而是需要非常仔细地考虑的一个主要改变。

定义新版本 IP 的工作产生了滚雪球的效应。网络设计人员总的意见是, 如果要对 IP 做这样大的改变, 也许最好也同时尽可能多地解决 IP 所存在的其他的问题, 例如 IP 协议对多媒体通信和安全性的支持问题。因此, 1990 年, Internet 工程任务组 (Internet Engineering Task Force, IETF) 着手研制一个新的 IP 版本, 其主要目标如下。

- ✧ 具有非常大的地址空间, 即使各个单位和组织对分配的地址利用率不高, 也能支持数十亿以上的主机。
- ✧ 减少路由选择表的尺寸。
- ✧ 简化协议, 允许路由器更快地处理分组。
- ✧ 提供比现在的 IP 更好的安全性 (身份验证和保密)。
- ✧ 更多地关注服务类型, 特别是实时数据。
- ✧ 通过允许指定范围来辅助组播服务。
- ✧ 允许主机移动地理位置 (漫游) 而不用改变其 IP 地址。
- ✧ 允许协议在未来进一步演变。
- ✧ 允许老的和新的协议在若干年内共存。

1992 年 6 月 IETF 公开征求对下一代 IP (IPng) 的建议, 随后收到了若干个提案, 到 1994 年就形成了 IPng 的最后设计。1995 年 1 月 RFC1752 “下一代 IP 建议书” 的发表是一个重要的里程碑。RFC1752 概述了 IPng 的需求, 规定了 PDU 格式, 突出了下一代 IP 在寻址、路由选择和保安等方面采用的方法。这个新一代的 IP 现在已正式地称作 IPv6。有一系列的 Internet 文档描述 IPv6 的细节, 它们包括从总体上描述 IPv6 的 RFC1883, 讨论在 IPv6 头中的流标记的 RFC1809, 以及处理 IPv6 寻址方面的 RFC1884、RFC1886 和 RFC1887。

虽然 IPv6 跟 IPv4 不兼容，但是总的来说，它跟所有其他的 Internet 协议兼容，包括 TCP（Transmission Control Protocol，传输控制协议）、UDP（User Datagram Protocol，用户数据报协议）、ICMP（Internet Control Message Protocol，因特网控制消息协议）、IGMP（Internet Group Management Protocol，因特网组管理协议）、OSPF（Open Shortest Path First，最短通路优先）、BGP（Border Gateway Protocol，边界网关协议）和 DNS（Domain Name System，域名系统），只是在少数地方作了必要的修改（大部分是为了处理长的地址）。IPv6 很好地满足了预定的目标。首要的原因在于，IPv6 有比 IPv4 长得多的地址。IPv6 的地址用 16 个字节表示，地址空间是 IPv4 的 2^{96} 倍，足以给每个人分配 5×10^{28} 个具唯一性的地址。无论未来怎样发展，这么多的地址也是够用的。

IPv6 第二个主要改进的地方是简化了 IP 分组头，它包含 8 个段（IPv4 是 12 个段）。这一改变使得路由器能够更快地处理分组，从而可以改善吞吐率。

第三个主要改进的地方是 IPv6 更好地支持选项。这一改变对新的分组头很重要，因为一些从前是必要的段现在变成可选的了。此外，表示选项的方式也有所不同，使得路由器能够简单地跳过跟它们无关的选项。这一特征加快了分组处理速度。

IPv6 具有重大举措的第四个方面是安全性。身份验证和保安功能是这个新的 IP 的关键特征。

最后一项重要改进的地方是有关资源分配的。取代 IPv4 的服务类型段，IPv6 的流标记段支持对属于一个特别的交通流（对应的发送端可能请求特别的处理）的标记，从而能够支持诸如实时视频这样的特殊交通。

1.1 IPv6 分组格式

IPv6 头由一个必需的基本头和可选的扩展头构成。图 1-1 所示为 IPv6 基本的头格式，分组应该以网络字节顺序传送。

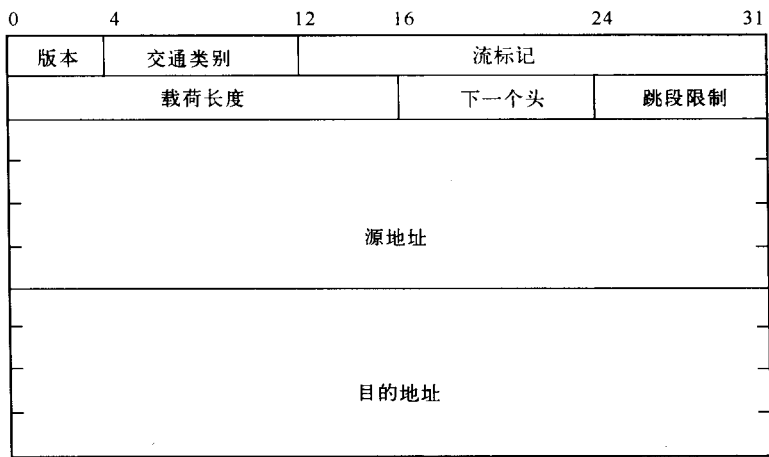


图 1-1 IPv6 基本的头格式

对 IPv6 基本头中的每个段的描述如下。

- (1) 版本。版本段指定协议的版本号，对于 IPv6 应该是 6。版本段的位置和长度都没有

改变，因此协议软件能够很快地识别分组的版本。

(2) 交通类别。交通类别段指定分组的交通类别或优先级。设置交通类别段的目的是支持区分服务。

(3) 流标记。流标记段可以用来标识分组请求的 QoS。在标准中，一个流被定义为“从一个特别的源发送到特别的目的地（单播或组播）并且要求中间的路由器对其做特别处理的一个序列的分组”。一个可能要使用流标记的例子是分组视频系统，该系统需要在一定的时间限制内把其分组投递到目的地，收到这些分组的路由器必须依照它们的请求对它们进行处理，不支持流的主机需要把流标记段置成零。

(4) 载荷长度。载荷长度段表明数据的长度（不包括头）。由于该段的长度是 16 位，所以载荷长度被限制为 65 535 字节。但是可以使用在扩展头中的选项发送更大的载荷。

(5) 下一个头。下一个头段用于标识在基本头后面的扩展头的类型。扩展头类似于在 IPv4 中的选项段，但它更灵活，更有效。

(6) 跳段限制。跳段限制段代替在 IPv4 中的 TTL（Time to Live，生存时间）段。其值指定分组在被一个路由器丢弃之前可以经过的跳段数。

(7) 源地址和目的地址。源地址和目的地址段分别标识源主机和目的主机。它们的地址格式将在后面介绍。

IPv6 建立在对 IP 的功能有增加的设计思想上，因此 IPv6 不是 IPv4 的简单演进，而是有实质性的改进。IPv6 头实际上要比 IPv4 头简单，IPv6 头仅有 6 个域和两个地址，而 IPv4 有 10 个固定域、两个地址及一些选项（参见图 1-2）。

0	4	8	16	19	24	31
版本	IP分组头长	服务类型	总长度			
标 识			标志	分割偏移		
生存时间		协议	头检验和			
源IP地址						
目的IP地址						
选项					填充	

图 1-2 IPv4 头格式

IPv6 取消了 IPv4 的头长、服务类型、标识符、标志、分割偏移和头检验和；在 IPv6 中，IPv4 分组格式的 3 个域：长度、协议类型和生存时间被重新命名并稍微改变了定义。IPv6 整个地修改了 IPv4 的选项机制，并增加了两个域：交通类别和流标记。

相对 IPv4 来说，IPv6 中含义和位置都未改变的仅有的域是开头的 4 位，即版本。网络程序可以使用起始的版本域确定对分组的处理方式。如果该域的二进制码是 0100（十进制 4），就当作 IPv4 处理，如果是 0110（十进制 6），就被认为是 IPv6 分组。

IPv4 头的设计是基于 1975 年的技术状态。20 年以后，IPv6 对其作了 3 个方面主要的简化：

- ✎ 对所有的头都分配固定的格式；
- ✎ 去掉头检验；
- ✎ 去掉逐跳分割过程。

IPv6 头不包含任何选项成分，但这并不意味着不可以对特殊分组表示选项。与 IPv4 不同，IPv6 的选项功能不是通过可变量选项取得的，而是把扩展头附加到主头后面。其明显的结果是 IPv6 不再需要一个头长度。

去除头检验的主要优点是减少头处理的代价，因为没有必要在每一中继站都检查和更新检验和的值。其风险是未监测到的差错可能导致对分组作错误的路由选择，但这种风险很小，因为大多数封装过程都包含一个分组检验和。事实上，在 IEEE 802 网络的介质访问控制过程中、在使用 ATM (Asynchronous Transfer Mode, 异步传输模式) 线路的适配层中及在用于串行链路的 PPP 协议 (Point to Point Protocol, 点对点协议) 的成帧过程中都有检验和域。

IPv4 包括一个分割过程，使得发送端可以发送大的分组而不用担心中继的能力。这些大的分组在必要的时候可以被分割成适当大小的片段，接收端等待所有这些片段的到来，并重组分组。但实践表明，这种分割与重组过程会产生一些负面效应。假定在仅能够运载小的分组的中间网络上尝试转发大的分组，由负责中继分组的路由器对 IP 分组进行分割，那么一个分组的成功传输依赖于每个片段的成功传输，只要有一个片段丢失了，整个分组就必须重传，结果将是对网络的低效使用。

IPv6 的规则是，主机通过一个称作最大通路 MTU (Maximum Transfer Unit, 最大传输单元) 发现的过程得知可以被接受的最大片段大小，如果主机发送超过这个大小的分组，这些分组将简单地被拒绝。因此 IPv6 不再像 IPv4 那样设立分割控制域 (包括分组标识符、分割控制标志和片段偏移)。但 IPv6 包括一个端对端的分割规范，而且根据 1996 年的规范，所有的 IPv6 网络都被假定能够运载 536 字节的载荷。在 1997 年的 IPv6 版本中，这个载荷长度又被提升到 1 500 字节，这样，不愿意发现或记住通路 MTU 的主机可以简单地发送小的分组。

IPv6 的最后一项简化是去掉了服务类型 (Type of Service, TOS) 域。在 IPv4 中，主机可以设置 TOS 的值，表示对最短的、最宽的、最可靠的或最安全的通路的期望，但应用程序并没有普遍地使用这个域。在 IPv6 中则提供了处理这些期望的机制。

与 IPv4 类似，IPv6 头包括分组长度、生存时间和协议类型，但这些域的定义都作了稍微修改。

IPv6 的载荷长度代替了 IPv4 的总长度。这里有细微的差别，因为按照定义，载荷长度是在头后面运载的数据的长度。例如，假定载荷是一个 TCP 分组，包括 20 字节的 TCP 头和 400 字节的应用数据，在 IPv4 中，要在这个 TCP 分组的前面加上 1 个 20 字节的 IPv4 头，总长度将是 440 字节；按照 IPv4 的总长度概念，在 IPv6 中将加上一个 40 字节的 IPv6 头。但实际上，在 IPv6 中，载荷长度被设置成 420，包括 TCP 报文段，也包括可能有的全部 IPv6 扩展头。在 IPv6 中，长度域也像 IPv4 那样在 16 位上编码，这就把分组尺寸限制到 64 K 字节。不过 IPv6 使用巨大数据报选项 (属于逐跳选项头) 提供对比较大的分组的传送服务。

在 IPv4 中，生存时间域用于说明分组在被丢弃以前允许在网络中存在的时间，以秒为单位。生存时间的概念是基于对 TCP 协议的理论分析。如果允许分组在网络中无限期地存在着，那么原有的拷贝可能在不可预期的时间退出，从而引起协议错误。IPv4 规范强制每个路由器把生存时间域减少 1 秒，如果在路由器中排队等待的时间较长，则再减去这个等待时间 (大于 1 秒)。但是，要精确地估计一个特定分组的等待时间是很困难的。由于这个时间通常是以毫秒计，而不是以秒计，大多数路由器只是简单地在每一中继处把 TTL 值减 1。这种方式在 IPv6 中已经变成正规的做法，所以相应的域名也改成跳段限制，它以跳段数目

计算，而不以秒的数目计算。

IPv6 把协议（类型）域重新命名为下一个头类型来反映新的 IP 分组结构。在 IPv4 中，IP 分组头后面紧接着就是传输协议数据，如一个 UDP 或 TCP 分组。在 IPv6 的情况下，如果 IP 分组封装 TCP 或 UDP 协议数据单元，那么头末尾的下一个头类型将被设置成协议类型 TCP（6）或 UDP（17）。

IPv6 头中有两个在 IPv4 中不存在的域：流标记和交通类别。这两个域主要是为了便于对实时交通的处理而设计的。交通类别域有 8 位，表示 256 种可能的级别，从最高优先级 0 到最低优先级 255。流标记用于表示需要同样处理的那些分组，它们由一个特定的源发送给一个特定的目的地，并具有指定的一组选择。

IPv4 头允许有选项，可以对某些分组作特别的处理。早先的规范包括对安全性选择的编码，源路由选择、记录路由（用于路由跟踪）及时间印迹，但这些选项并未被普遍采用。人们有足够的理由对某些分组作特别的处理，例如通过源路由选择请求一条特别的路由，或者指定接收端对一个分组作特别的处理。IPv6 规范则说明了如何通过扩展头来实现这类特别的处理。

在 IPv4 中，作为载荷的 TCP 分组紧接在 IP 头的后面。在 IPv6 中，Internet 头和载荷之间可能插入任意数目的扩展头。每个头用一个头类型表示，如图 1-3 所示，在每一个头中都包含紧跟在它后面的头类型，在最后一个扩展层头中则包含载荷的头类型（例如 TCP）。

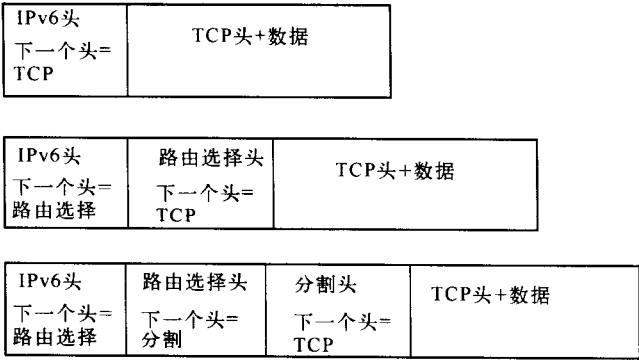


图 1-3 IPv6 头的菊花链示例

1.2 扩展头

当前的 IPv6 规范定义了 6 个扩展头：逐跳选项头、路由选择头、分割头、身份验证头、封装安全载荷头和目的地选择头。

表 1-1 列出了这些扩展头的类型及其编码。这些扩展头在分组中出现的次序应该跟该表中从顶至底的次序相同。

每个扩展头都用一个头类型标识。IPv6 的下一个头域可以包含一个扩展头的类型，也可以包含载荷的协议类型，例如 TCP 或 UDP。因此，头类型一定不能跟协议类型冲突，它们从同样的一组 256 个数字中分配。协议类型域基本上跟 IPv4 相同（虽然有些协议类型略有不同），例如 TCP 是 6，UDP 是 17，OSPF 是 89，ICMP（IPv4）是 1，ICMP（IPv6）是 2；而

HBH（逐跳选项，IPv6）是 0，RH（路由选择头，IPv6）是 43，FH（分割头，IPv6）是 44。

表 1-1 扩展头类型及其编码

头 编 码	头 类 型
0	逐跳选项头
43	路由选择头
44	分割头
51	身份验证头
52	封装安全载荷头
60	目的地选项头

1.2.1 路由选择头

在 IPv6 中对选项处理的最典型的例子是路由选择头，它的作用与 IPv4 中的源路由选项相同。这个头主要包含分组将被中继经过的中间节点的地址列表，源路由选择可以是严格的，也可以是松散的。

如图 1-4 所示，路由选择头由一组参数和随后的一个地址列表组成。前 32 位包含 4 个 8 位整数。

- ☞ 下一个头：标识在头的菊花链中紧跟在路由选择头后面的头的类型。
- ☞ 头扩展长度：用 64 位字的数目表示的头扩展长度，不包括开头 64 位（4 个 8 位整数和 32 个保留位）。
- ☞ 路由选择类型：设置为 0。
- ☞ 剩余段：表示在列表中剩余段的数目，该值的范围是从 0~23。

长度 8	8	8	8	位
下一个头	头扩展长度	路由选择类型=0	剩余段	
保留				
地址[1]				
地址[2]				
.....				
地址[n]				

图 1-4 类型 0 路由选择头

紧接着的 32 位是保留域，应该设置为 0。

路由选择头的剩余部分是一组 128 位地址的列表，编号从 1~N。在 IPv4 中，源路由编码在可选头域中，即使它们不被包括在源路由内明确说明的中继站列表中，所有的路由器都需要对其进行检查，因此对源路由分组的处理是非常缓慢的，该选项在实践中用得不多。在 IPv6 中，路由器仅查看路由选择头，确定它们是否能识别出它们在主头的目的地域中的地址。没有被明确地列在源路由列表中的中间路由器将转发分组而不作任何附加的处理，这应该能够产生比较好的性能。

在分组的目的地域中识别出自己的地址的站将检查路由选择头，检查在列表中是否至少还有 1 个域。如果结果是否定的，那么分组就已经到达源路由的终点，该站将跳过路由选择头去处理下一个头，其类型在下一个头参数中表明，如果检查的结果是肯定的，那么在列表

中至少还剩下 1 个城，该站将处理源路由选择。

在源路由中的下一地址的位置可从头扩展长度 ($H=2N$) 和剩余段数 (L) 参数推导出来。每个地址是 128 位长，头扩展长度 (路由地址段空间大小) 是指定地址段 (segment) 的 64 位字的数目。因此在列表中地址的数目 N 等于该长度 (H) 的一半，要处理的下一地址在列表中的位置号是 $N-L$ 。

所有 IPv6 规范的实现必须都能够处理类型 0 路由选择头，而且这些实现也必须准备好碰到并处理其他的类型，然后采取某种默认操作。

事实上，类型 0 路由选择头是一般路由选择头的子类型。如图 1-5 所示，一般路由选择头的组成是：32 位头后随类型特有的数据。32 位头就是在前面介绍过的在类型 0 头中的 4 个 8 位参数 (下一个头，头扩展长度，路由选择类型和剩余段)。“路由选择类型”说明使用的版本。“类型特有的数据”的格式和处理源路由的规则在每个路由选择类型的规范中解释。

下一个头	扩展头长	路由选择类型	剩余段
类型特有的数据			

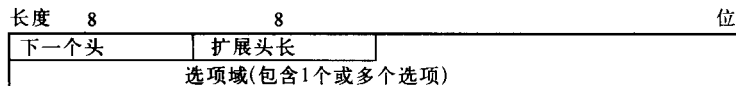
图 1-5 一般路由选择头

如果一个 IPv6 系统必须处理一个路由选择头，它将首先检查路由选择头类型和剩余段的数目。如果类型未知，分组应该被拒绝，并给源发送端返回一个 ICMP 错误报文 (参数问题)，其 ICMP 编码值为 0，使参数 (说明是什么样的错误) 指向路由选择类型段。但剩余段值 0 表明该分组已经到达最后的目的地，即使系统不懂得指定的路由选择类型，它也应该接受这样的分组。

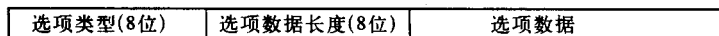
1.2.2 逐跳选项头

逐跳选项头用于沿着通路的所有路由器都必须查看的信息。某些管理或诊断功能需要给所有的路由器传递附加的信息，这正是逐跳选项头的目的，它用头类型 0 标识。在 IPv6 中，下一个头值 “null” 意味着存在逐跳选项头，这时即使目的地址不是本地节点地址也要对它进行处理。

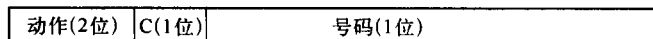
如图 1-6 (a) 所示，逐跳选项头的组成包括下一个头类型 (8 位)、一个扩展头长域 (8 位) 及一个或多个选项。扩展头长域表示在该选项头中 64 位字的数目，其中不包括开头的 64 位 (8 个字节)。例如，如果选项头仅由 8 个字节组成，那么长度域的值将是 0；如果选项头有 32 个字节组成，那么长度域的值将是 3。



(a) 按跳段逐级处理的选项头



(b) 每个选项的组成



(c) 选项类型标识符的结构

图 1-6 逐跳选项头的组成

选项域包含一个选项列表。如图 1-6 (b) 所示，每个选项的编码是可变数目的字节。选项类型段是表示选项类型的 8 位标识符，在它的后面是 8 位整数，表示在选项数据段中的字节的数目。选项类型标识符的结构如图 1-6 (c) 所示，如果处理节点不识别该选项时，两个高序位编码表示必须采取的动作。第 3 位表示该选项在路途中是否可以改变。最后 5 位表示选项号码本身。路途中改变位 (C) 表示该选项可以被途中的中继站修改，类似于在路由选择头中的剩余段域。这类选项不被端对端的检验和计算。

某些选项仅仅提供关于分组上下文的附加信息，或者表示倾向性。如果它们不被识别，则可以被安全地忽略。处理节点可以只是跳过该选项数据域，其长度由选项数据长度字节标出，然后再继续处理在头中的剩余选项。与此相反，有一些选项是关键性的，必须把分组丢弃。然而，当一个站丢弃一个分组时，一般的规则是往回发送一个 ICMP 报文，这不一定是发送端所期待的事情。动作位用于指定所请求的动作 (参见表 1-2)，当发送一个 ICMP 报文时，应该把其编码设置成 2 (表示参数问题)，并将报文的参数指向未被识别的选项类型。

表 1-2 对未能识别的选项的处理

位	动 作
00	跳过这一选项
01	丢弃分组，不发送 ICMP 报文
10	丢弃分组，发送 ICMP 报文，即使目的地址是组播也如此
11	丢弃分组，在目的地址非组播的情况下发送 ICMP 报文

在当前的规范中已经定义了一个巨大载荷选项 (参见图 1-7)，其选项类型是 194。此外，一个路由器警报选项正在讨论之中。

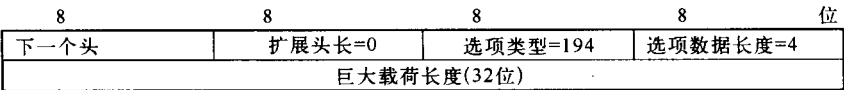


图 1-7 巨大载荷选项格式

巨大载荷选项用于发送非常大的分组，其长度不能仅用 16 位来编码。当使用这个选项时，把 IPv6 长度域置成 0，处理节点解码，求得实际的分组长度，并编码为 32 位整数。为方便这个长度域的处理，选项类型域 (194) 的对准需求被确定为 $4n+2$ (2 表示下一个头和扩展头长引起的 2 个字节偏移)，因此，表示分组实际长度的长度域本身起始于 32 位边界。

如果长度小于 65 535 字节，不应使用巨大载荷选项。如果分组运载一个分割头，也不应使用巨大载荷选项。

在一些情况下，发送给目的地的信息会影响所有中途的路由器。例如，有的多播路由算法使用管理分组标记数据将要遵循的分布树，而资源预留协议 (Resource Reservation Protocol, RSVP) 使用报文标记随后要在其上执行预留的通路。通过使用逐跳选项，一个源可以把一些分组标记成包含在前往目的地的通路上所有路由器都应该查看的信息。当然，要达到这个目标，可以在每定义一个新算法时，定义一个新的逐跳选项，但这样做可能还是不够的。例如，可能需要对管理报文进行身份验证，因此仅在剥离逐跳选项头和身份验证头之后才可以对这些报文进行处理。其解决方案是定义一个类属“路由器警报”选项，目的是警示