

Large-Scale IP Network Solutions

网络核心技术内幕



综合 IP 网络设计解决方案



本书配套光盘内容包括：
与本书配套的电子书

[美] Khalid Raga 著
Mark Turner
希望图书创作室 译



北京希望电子出版社

Beijing Hope Electronic Press
www.bhp.com.cn

TP393
2122-6

(附光盘一张)

Large-Scale IP Network Solutions

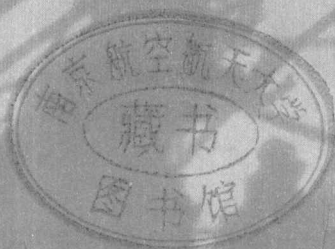
网络核心技术内幕

综合 IP 网络设计解决方案



本书配套光盘内容包括：
与本书配套的电子书

[美] Khalid Raga 著
Mark Turner
希望图书创作室 译



200145337

200145337



北京希望电子出版社

Beijing Hope Electronic Press

www.bhp.com.cn

内 容 简 介

本书是 21 世纪网络工程师设计宝典系列丛书之一。本书的两个作者，一个是 Cisco 公司从事大型网络设计的资深专家，另一个是从事大规模交换和路由产品的开发部经理。作者在本书中吸收了在网络环境中获得的经验，突出了可伸缩性特点，并强调实际网络配置问题。全书分为两个部分，共 16 章。第一部分首先回顾 Internet 的发展历程，概述了 IP 的基础原理和网络交换技术，接着详细探讨现代网络拓扑结构和设计中的约束，以及用于生成可伸缩网络体系结构的工具和技术，最后介绍路由器的发展和技术；第二部分着重介绍了路由选择信息协议版本 1 和版本 2、增强内部网关选择协议、开放最短路径优先、中间系统到中间系统、边界网关协议等路由协议的技术细节和配置，探讨了在路由协议间的迁移的原因和方法，介绍了协议无关多播及其扩展特性和实施问题，介绍了保证 QoS 的机制和可伸缩实施办法，对网络操作和管理的范畴和应用进行了讨论，最后三个不同类型的大型网络设计的实例。

本书内容全面，结构清晰，实用性极强，是从事 IP 网络的设计、配置和维护的网络工程师、管理员必备的指导书，也是高等院校计算机网络专业、信息工程专业人员的专业参考资料，还可以作为世界著名的网络工程师认证考试—CCDA 的自学和教学参考书。

本书配套光盘内容包括与本书配套的电子书。

版 权 声 明

本书英文版名为“Large-Scale IP Network Solutions”，由 Cisco Press 出版，版权归 Cisco Press 所有。本书中文版由 Cisco Press 授权出版。未经出版者书面许可，本书的任何部分不得以任何形式或任何手段复制或传播。

系 列 书：21 世纪网络工程师设计宝典系列（14）

书 名：网络核心技术内幕——综合 IP 网络设计解决方案

文 本 著 作 者：（美）Khalid Raza, Mark Turner 著 希望图书创作室 译

责 任 编 辑：龙启铭

C D 制 作 者：希望多媒体开发中心

C D 测 试 者：希望多媒体测试部

出版、发 行 者：北京希望电子出版社

地 址：北京海淀路 82 号，100080

网址：www.bhp.com.cn

E-mail: lwm@hope.com.cn

电话：010-62562329, 62541992, 62637101, 62637102, 62633308, 62633309

（发行和技术支持）

010-62613322-215（门市） 010-62531267（编辑部）

经 销：各地新华书店、软件连锁店

排 版：希望图书输出中心

CD 生 产 者：文录激光科技有限公司

文 本 印 刷 者：北京双青印刷厂

开本 / 规格：787×1092 16 开本 24.625 印张 564 千字

版次 / 印次：2000 年 5 月第 1 版 2000 年 5 月第 1 次印刷

印 数：0001-5000 册

本 版 号：ISBN7-900044-12-4/TP·12

定 价：58 元（1CD，含配套书）

说明：凡我社光盘配套图书若有缺页、倒页、脱页、自然破损，本社负责调换。

译者的话

世界已经跨入二十一世纪，迎接我们的是势不可挡的知识经济浪潮。在激烈的信息技术竞争中，网络占据了最主要的地位，谁掌握了网络，谁就掌握了主动，永远立于不败之地。

Cisco 系统公司的产品有如网络领域的中枢纽，在相关网络的市场中，占有压倒性的优势。毫不夸张地说，谁掌握了 Cisco 的技术和产品，谁掌握了网络，就能在 IT 领域占有一席之地。

本书的两个作者，一个是 Cisco 公司从事综合网络设计的资深专家，另一个是大规模交换和路由产品的开发部经理。作者在本书中吸收了综合网络环境中获得的经验，突出了可伸缩性特点，并强调实际网络配置问题。

读者会发现本书是设计和维护网络的无价之宝。我们坚信本书的出版必将对 Cisco 用户以及广大的网络技术人员提供用益的帮助。

本书是专为从事 IP 网络的设计、配置和维护的网络工程师、管理员而设计的。本书着重于可伸缩的网络设计，以及使用实际网络配置的具体应用的实现方法。虽然这些配置实例基于 Cisco IOS，但许多思想是适用于任何路由选择平台的。

本书的主要目的是使读者成为网络专家。要求读者对 TCP/IP、Cisco IOS、IP 路由选择和网络设备如集线器、交换机和路由器有个基本的理解。

全书分为两个部分、十六章。第一部分首先回顾 Internet 的发展历程，概述了 IP 的基础原理和网络交换技术，接着详细探讨现代网络拓扑结构和设计中的约束，以及用于生成可伸缩网络体系结构的工具和技术，最后介绍路由器的发展和技术；第二部分着重介绍了路由选择信息协议版本 1 和版本 2、增强内部网关选择协议、开放最短路径优先、中间系统到中间系统、边界网关协议等路由协议的技术细节和配置，探讨了在路由协议间的迁移的原因和方法，介绍了协议无关多播及其扩展特性和实施问题，介绍了保证 QoS 的机制和可伸缩实施办法，对网络操作和管理的范畴和应用进行了讨论，最后三个不同类型的综合网络设计的实例。

本书是多人努力合作的结果，参加本书翻译、审校、录入的具体人员有：袁勤勇、李学群、李泽、侯志霞、丁岩、于平、李明、王海艳、陆勋等。本书出版过程中还得到北京希望电子出版社的陆卫民、龙启铭等人的大力帮助，在此深表感谢。

在本书的翻译过程中，我们煞费苦心，力求尽善尽美。但是由于时间紧迫，加之译者水平有限，疏漏之处在所难免，恳请广大读者批评指正。

袁勤勇执笔

2000. 3.

作者简介

Khalid Raza 在 5 年多的时间里一直在 Cisco Systems 公司从事综合网络设计有关的工作。他拥有工程管理专业的硕士学位和电子工程专业的学士学位。他在 CCIE 编程方面的贡献包括帮助编写考试的实验部分和 CCIE 再验证的 IP-ISP 新测试。Khalid 的专长在路由选择协议方面，他还在各种各样有关可伸缩 IP 网络设计的会议上发表过论文。

Mark Turner 是 Cisco Systems 公司负责大规模交换和路由的开发部经理。早在二十世纪九十年代，他维护和扩充了澳大利亚教育和研究网络（Australian Academic and Research Network），开始从事综合 IP 网络的设计工作。从那时开始，他已经以工程和网络管理资格为几个大公司设计了企业网（Intranet），最近完成的是 NASA Science Internet 网。这些年来，Mark 对主要项目的贡献有对安全和开放网络的完整设计以及对现有大型网络的物理和逻辑路由结构的逐步升级。

致 谢

我要感谢 Cisco Press 和 Cisco Systems 公司，以及我们的技术编辑们，在他们的允许下，我得以从事这本书的编写。我也要感谢 Atif Khan、Henk Smit 和 Ray Rios 在讨论中对我的提示。同时我也要感谢 Mark Johnson 和 Ray Rios 在我初到 Cisco 时给予我的帮助表示真诚的感谢。最后，我要对 Mike Quinn 和 Joe Pinto 在该项目中表现出的灵活性表示感谢。

—Khalid Raza

我要感谢在 Cisco Systems 公司的很多朋友和同事。他们同样享有概括在书中的一些观点。其中包括 Srihari Ramachandra 和 Ravi Chandra 在边界网关协议 (BGP) 方面的观点；Enke Chen 在可伸缩边界网关协议方面的观点；John Meylor、Dave Meyer 和 John Zwiebel 在多播方面的观点；Dave Rowell 在交换方面的专长；以及与我的合著者有关路由的许多有趣的讨论。

我也要感谢 Jim McCabe 在网络设计中的深刻见解。特别要感谢我的朋友们，他们能够理解我这么长时间以来为什么在深夜和周末忙碌，Charles Goldberg 鼓励我享受生活、Angelo 保证我在写作的同时坚持锻炼、Jude 的战略以及 Em 的其它每件事。

目 录

第一部分 Internet

第1章 数据网络的发展	1
1.1 通信历史概述	1
1.2 Internet 的发展	3
1.3 万维网 (WWW)	10
1.4 今日 Internet	13
1.5 现代 Internet 的体系结构	18
1.6 企业与开放网络的兴衰	20
1.7 Internet 的未来	26
1.8 小结	27
1.9 复习题	27
第2章 IP 基础	29
2.1 IP 基本概念	29
2.2 可变长子网掩码	34
2.3 无类型域间路由	36
2.4 小结	38
2.5 复习题	38
第3章 网络技术	39
3.1 包交换、电路交换和消息交换	39
3.2 局域网技术	42
3.3 广域网技术	45
3.4 城域网 (Metropolitan-Area Network) 及其技术	50
3.5 小结	52
3.6 复习题	52
第4章 网络拓扑结构和设计	53
4.1 需求和约束	53
4.2 工具和技术	54
4.3 层次结构	57
4.4 主干核心网络设计	60
4.5 分布式/区域网络设计	62
4.6 访问设计	63

4.7 小结	64
4.8 复习题	64
第5章 路由器	66
5.1 路由器结构	66
5.2 Cisco 交换算法的演变	68
5.3 路由选择和转发	77
5.4 高速缓存技术案例分析	78
5.5 小结	84
5.6 复习题	85

第二部分 核心和分布式网络

第6章 路由选择信息协议	87
6.1 RIP 概述	87
6.2 RIP 包格式	93
6.3 RIPV1 配置示例	99
6.4 小结	102
6.5 复习题	102
第7章 路由器信息协议版本 2	103
7.1 RIP 操作概述	103
7.2 Cisco 的 RIP 实现	106
7.3 RIP 和缺省路由	114
7.4 小结	114
7.5 复习题	115
第8章 增强型内部网关路由选择协议	116
8.1 EIGRP 的基础和操作	116
8.2 EIGRP 拓扑结构表	122
8.3 EIGRP 配置命令	124
8.4 EIGRP 无类型总结	128
8.5 调整 EIGRP 参数	137
8.6 分割范围和 EIGRP	137
8.7 小结	138
8.8 复习题	139
第9章 开放最短路径优先	140



9.1 OSPF 的基本原理.....	140	12.4 复习题.....	236
9.2 链路状态协议概述.....	141	第 13 章 协议无关组播	237
9.3 LSA 种类.....	146	13.1 组播路由选择协议.....	237
9.4 OSPF 区域的概念.....	155	13.2 操作基础.....	238
9.5 启用和配置 OSPF.....	157	13.3 IGMP 和 PIM 协议描述.....	243
9.6 小结.....	167	13.4 组播伸缩性特征.....	253
9.7 复习题.....	167	13.5 在综合网络中使用组播.....	255
第 10 章 中间系统到中间系统	168	13.6 小结.....	256
10.1 IS-IS 概述.....	168	13.7 复习题.....	257
10.2 IS-IS 的基础和操作.....	168	第 14 章 服务质量特性	258
10.3 用 IS-IS 寻址.....	169	14.1 Qos 概述.....	258
10.4 链路状态概念.....	172	14.2 QoS 策略传播.....	258
10.5 使用 IS-IS 伪节点.....	175	14.3 拥塞管理算法.....	259
10.6 使用 IS-IS 非伪节点.....	176	14.4 拥塞避免算法.....	264
10.7 理解 1 级和 2 级路由.....	176	14.5 在综合网络中配置 QoS.....	275
10.8 IS-IS 包.....	177	14.6 小结.....	275
10.9 IS-IS 扩散.....	179	14.7 案例分析: 在大型网络中 应用不同服务.....	275
10.10 路由总结.....	180	14.8 复习题.....	277
10.11 伸缩 IS-IS.....	180	第 15 章 网络操作和管理	278
10.12 NBMA 网络上的 IS-IS.....	181	15.1 网络管理概述.....	278
10.13 基本 IS-IS 配置.....	182	15.2 网络管理系统.....	279
10.14 小结.....	185	15.3 简单网络管理协议.....	280
10.15 复习题.....	186	15.4 Netflow.....	287
第 11 章 边界网关协议	187	15.5 故障管理.....	289
11.1 BGP 概述.....	187	15.6 配置和安全性管理.....	296
11.2 BGP 操作基础.....	188	15.7 特别的滥用问题.....	304
11.3 BGP4 协议的描述.....	190	15.8 性能和计数管理.....	304
11.4 BGP 的有限状态机.....	203	15.9 小结: 综合网络管理的复核列表.....	307
11.5 路由选择策略和 BGP 决定算法.....	205	15.10 复习题.....	310
11.6 BGP 可升级特征.....	207	第 16 章 设计和配置案例研究	312
11.7 综合网络配置问题.....	211	16.1 案例研究 1: 企业 Alpha.com.....	312
11.8 小结.....	215	16.2 案例研究 2: MKS 零售商店.....	342
11.9 复习题.....	215	16.3 案例研究 3: 一个 ISP 网络.....	365
第 12 章 迁移技术	216	16.4 小结.....	381
12.1 交换协议.....	216		
12.2 路由选择协议的迁移.....	218		
12.3 小结.....	235		

第一部分 Internet

第 1 章 数据网络的发展

本章为你在学习本书中其它更多的技术材料之前提供一个大致的介绍。下面是这一章的有关内容：

通信历史概述 这一节简要地回顾了从早期到现在通信基础设施的发展过程。

Internet 的发展过程 这一节深入研究了 Internet 的发展。以技术细节为主，同时也讨论了政治、社会和经济问题。随着对 Internet 介绍的展开，你将理解在升级综合 IP 网络方面的迫切问题。本节还讨论了诸如网络拓扑结构、路由、网络管理以及分布和多媒体应用的支持等问题。

今天的 Internet 在学习了前面内容以后，你将阅读有关现代 Internet 的体系结构。这一节描述了关键性研究目标、Internet NAP、路由策略和网络地址登记。最后是有关当今 Internet 的拓扑结构的综述。

专用和 OSI 网络的兴衰 这一节展示了迅速萎缩的非 IP 网络协议的世界。对正反两方面进行了说明，并解释了为什么这些不是未来必需的关键技术。

Internet 的未来 在这一节里，将介绍未来 Internet 和综合网络可能发展方向。

本章从历史的角度帮助您理解在数据网络的发展和升级过程中进行了哪些革新及其原因。本章也提出了一个更鲜明的观点：开发过程中的一些观念的使用是周期性的。这种开发的周期性给予了“发明转轮”一个全新的意义。

技术的革新带来了观念的革新，新的解决方案在原先的技术和经济条件下是难以实施的。最极端的例子是光学通信的使用。尽管很多年前人们已成功地使用了火把和闪光，上个世纪光学通信还不被人接受，这是因为使用简单的铜线引导低频的电磁辐射更容易。随着光纤的到来，事情发生了变化。

出于相似的原因，在电子领域也历史性地实现了数据包的交换。然而，最近在分区多路复用技术和光学网络技术的创新可能很快消除对基于高带宽通信的电子交换的需求。

1.1 通信历史概述

现在网络是商务活动和个人生活的核心部分。今天的商务由原来离不开电话服务而演变为离不开数据网络设施。不难理解，为什么公司要花费大量的时间和金钱来“护理”这一关键性的资源。

这种对数据网络的依赖是如何出现的，而且为什么会出现呢？简单地说，这是因为网络强化了所有历史通信机制。近五万年的语音、至少三千五百年的书写通信以及几千年的图像生成都可以通过网络来捕捉和传达到地球的任何地方。

二十世纪八十年代, 光纤光学改善了距离、成本和可靠性问题; CB 无线电教给我们在没有集中通信设施提供者的情况下进行对等通信和自调节通信。军队的需要导致了反击进攻的网络技术的开发。该技术也可以对其它类型的失败进行反击。今天的光交换和多路复用技术再一次说明了它的下一个在性能、规模和可靠性方面的飞跃。商务和个人应用不能承担这种高可靠性技术费用。

数据通信由于将各个局域网上的用户孤岛连接到大型机 (IBM 的系统网络结构[SNA]和 Digital 的 DECnet) 上以及它们彼此连接的需要而发展起来的。随着时间的推移, 整个宽广的地理学领域都需要这些服务。然后, 这些服务又扩展到管理控制以及介质和协议转换。在二十世纪八十年代, 路由器成为网络的关键部件, 与此同时, 开发了用于支持多媒体通信的全球网络配置的异步传输模式 (ATM) 单元交换技术。

ATM 的设计者受到了必须支持传统语音网络的需求的限制。这并不奇怪, 因为在当时语音网络的收入超过了其它形式的通信设施。如果开发新的应用, 许多人认为它们可能象电视一样, 可以用于通信或娱乐。

几乎没有人预测到了 Internet 的到来。它毕竟不是以实时声音或视频图像吸引人的, 但家庭个人电脑的普遍存在使得一些应用需求出现了。这些需求包括一对一或一对多的通信需求, 如电子邮件和聊天组, 以及强大的网络浏览器和 Internet 搜索引擎。后二者能把 Internet 转换为人们可以在其中进行旅行、学习、讲授和共享的虚拟世界。用户不需要每秒数兆的速度进入这个世界: 32kbps 就挺高兴, 64kbps 已感到幸福极了, 128kbps 时简直是进入天堂了。

桌面电脑计算能力的提高和对等网络应用的增加, 它们与网络结构之间形成了主要冲突。现代网络体系结构在用户工作站上要求具有智能化和自调节能力, 但他们却只提供了这样一个网络结构: 只具有足够支持包转发的智能性。这种情况与在许多专用解决方案中通过使用复杂的网络设备把简单终端设备与智能化的主机相连接的方法, 如 IBM 的 SNA 形成了鲜明的对照。

注释 一些学校建议网络将变得更加智能化, 而用户工作端则更“傻”。他们相信带宽将比本地内存或 CPU 更便宜, 因此, 计算资源最好保存在中央共享设备上。网络电视和有声服务传送是这种观点的具体体现。

总体来看, 技术开发的效果以及商务、客户和社会的变革需求是清楚的。数据网络以每年百分之二十五的速度增长, 传统的语音服务只增长百分之六, 而 Internet 在短短的几个月内就以成倍的速度增长。使用传统语音这个术语是因为 Internet 现在也传递语音, 而且最近几年一些提供商宣布他们的 Internet 提供 IP 语音服务。

通过传送数据包所获得的收入接近甚至也许超过了传送传统电话语音的收入。语音很快就变成了在数据网络主题中的变种——只是另一个应用程序的包。

根本上来看, 竞争是设计和开发的动力, 不管是电报对电话, 还是传统电话对 Internet 电话。通过采用能够带来更便宜、更好或更灵活服务的新技术, 供应商希望能获得超过他们的竞争对手的商业利润。自然, 设计、规划和实施的好的网络将需要更多的新技术。

1.2 Internet 的发展

从社会、经济、文化和技术的角度看, Internet 已经给大部分人的生活带来了巨大的变化。这种变化将会涉及到更多的人。与电话、电视和汽车一样, Internet 的连通正在快速地成为每个家庭的日用品。然而, 与所发生的变化一样具有戏剧性的是, 至少在开始时, Internet 的发展速度比电话网络慢很多(也许有人会争辩说这只是因为前者依赖于后者), 这一点值得人们记住。

二十世纪六十年代, 有关普遍存在的通信网络的想法已经不新鲜了, 但是不满足于网络使用的角度不仅仅是个人通信, 而应该是计算机程序和其它形式的任何数据的交换, 这在当时还是相当激进的想法。原因是, 当时的通信技术还不够灵活, 不允许这种形式的交换。然后, 出现了 ARPANET。

1.2.1 ARPANET

从 1961 年到 1967 年, 技术人员在 MIT (美国麻省理工学院)、RAND (兰德公司) 和 NPL (英国国家物理实验所) 彼此独立地进行了一系列的后来被称为包网络的试验。1967 年公布的一份有关这些试验的报告是对试验性的广域网交换网络的设计, 称为 ARPANET。

这是技术的转折点。它对于需要成长和发展来满足用户的挑战性需求, 支持对等连网和多媒体通信的网络性能具有深远的影响。在取代传统的局域网环境下分时多路复用技术的同时, ARPANET 通过分层提供了在大范围内利用 TDM (时分复用) 系统的优点的性能, 并实现了两者的无缝结合。

1969 年, 经过了多年在实验室对这一技术的观察, 美国国防部将 ARPANET 投入运行, 该网以每秒 50K 的速度从 SDS、IBM 和 DEC 把四个节点连接起来。网络使用的接口信息处理机 (IMP) 是 Honeywell 的 516 小型机, 该机带有巨大的 24K 的内存, 代码由 BBN 公司提供。随之而来的是 BBN C-30s 和 C-300s, 以及后来在 1984 年被改名的 Packet Switch Nodes (包交换节点)。除了它的历史重要性, ARPANET 还具有两个 Internet 一直沿用至今的鲜明特点: 终端主机来自不同制造商, 和所提议的初始带宽远低于需要的带宽。

其中的一个终端节点是美国史坦福大学研究所 (SRI), 该所提供了第一个网络信息中心和请求评论库。这个开放的设计和讨论的文档库涉及到网络工程, 并提供了一个令人称奇的成功的有关出版和评论的论坛。具有讽刺意味的是, Internet 标准的易用性, 与开放系统互联 (OSI) 论坛所准备的相反, 成为最后 OSI 组消亡的主要原因。

1970 年, 网络控制协议 (NCP) 的细节得到了丰富, 而且, ARPANET 开始使用该协议进行通信。从这点上说, 应用设计人员可以开始工作了。1972 年电子邮件引入网络, 并占据了主动地位, 根据通信量的大小, 占据第二的是 FTP, 这种状态一直持续到 1995 年三月万维网超过 FTP 为止。

网络控制协议 (NCP) 的局限性显现了出来: 可升级性和地址能力的限制以及网络的可靠性数据传输的可靠性是主要的问题。

TCP/IP 引入到终端系统, 并且它解决了今日 Internet 操作的基础问题。该设计允许自

主拥有和管理网络，唯一需要 Internet 网络配合的重要问题是网络地址的分配。需要进一步考虑的问题（和 TCP/IP 的属性）如下：

- 假设网络/尽力传递 (best-effort delivery)
- 维护在包交换中的非单位流量流状态 (no per-traffic flow state)
- 提供监测和丢弃循环包的途径。
- 提供多个传送中的包和路由。
- 包含在不影响其它性能的情况下的高效的实现方法。
- 支持包的拆分和重组。
- 提供用于监测包复制并纠正错误的方法。
- 包括操作系统无关性。
- 提供流和拥塞的控制。
- 提供操作系统无关性。
- 支持可扩展设计。

错误纠正决定（如一些针对实时数据的错误纠正）最好交给应用程序来处理，这种要求迫切时，用户数据报协议 (UDP) 应运而生了。例如，在实时的音频应用程序中，如果被破坏的声音段的播放时间已经过去的话，就不值得去纠正该破坏了的声音段。

独立的协议实现最初是由 DARPA 委托开发的，但是随着时间的推移，商业投资商开始为许多操作系统和平台实现该协议，其中包括 IBM 的 PC 机。

当时的一个不幸的选择是使用了 32 位的地址空间，其中，前 8 位指明网络 ID，其余的 24 位指明主机 ID。初始的协议规范在引入 LAN 之前就发布了，是提出以太网思想后一年，至少比 PC 成为桌面商品早了 15 年。

在二十世纪七十年代后期，通过对 LAN 技术的引入，人们很快发现 8/24 位网络/主机地址分配模式将不能运行了。许多要进入实验的小网络操作员要求比 24 位小很多的主机地址空间。这时引入了分为 A、B 和 C 三级的可分类地址分配模式。这种模式的出笼是另一个不幸的决策，因为人们很快认识到 A 类地址空间是一个太大的分配块。二十世纪九十年代引入了无类型域间路由 (CIDR) 和无类型路由，路由器的可分类性能导致了在子路由中的混乱和冲突路由状态——或者更确切地说，是可分类性能的不足。

注释 无类型路由允许自由地给 IP 地址空间的网络 and 主机部分分配地址，这一点优于与预定义的类型 A、B 和 C 界限一致的地址分配。

1983 年 1 月是从 NCI 到 TCP/IP 转变的标志性的一个月。当时仅仅几百台主机，这种转变是可行的，而且运行得惊人的顺利。在 Internet 的今天，无论如何——即使在普通的集体网络中——平稳的迁移计划是任何新技术更新程序的关键部分。

也是在 1983 年，在往 TCP/IP 转变的推动下，ARPANET 发生了巨大的分离。一半仍然叫做 ARPANET，用于研究；另一半 MILNET 用于不分类的军事活动。MILNET 后来被集成到国防数据网。

到了 1984 年，主机数量超过了 1000 台。仅这一点，将机器地址与功能或位置关联就变得几乎不可能了，因此发明了域名系统 (DNS)。DNS 支持网络节点名的层次分配和分解（将网络节点名映射为 IP 地址）。然而，随着时间的推移，DNS 演变为不仅仅来实现

从网络节点名到 IP 地址的映射：它的支持多种资源记录的能力结合其可升级性和分布广泛的部署，意味着它可以用于所有的资源定位功能。

在二十世纪八十年代初期，一些商业的 ARPANET 的代替品（许多将电子邮件作为主应用程序）开始出现。随着时间的推移，它们都通过消息网关与 ARPANET 或 Internet 建立了连接。但是，这些网关与下面的 Internet 协议套件的主要设计能力不一致：提供自主操作网络的无缝连接的特性。最终，无缝的 TCP/IP 占了上风，其它的网络有的消失，有的发展为学院或商业的 Internet 服务提供商（ISP）。

另一个有趣的发展趋势在二十世纪八十年代初期开始出现：LAN 变得普遍起来。各个 ARPANET 站点要求不只跟网络的一台计算机连接，而是与多台计算机连接。Internet 开始看起来象是 LAN 的集合。而这些 LAN 通过路由器和广域网与 ARPANET 核心连接起来。（如图 1-1）。1982 年，在 RFC827 中描述了外部网关协议（EGP），并第一次出现了现代大型 IP 网络层次结构的迹象。

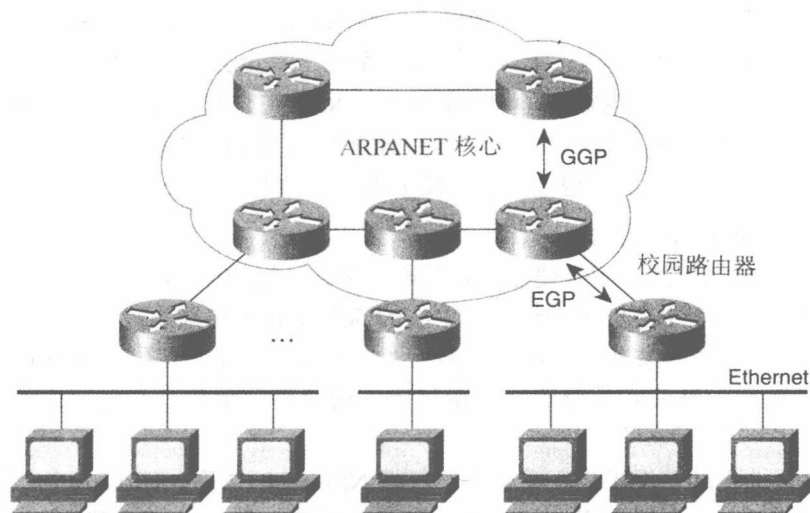


图 1-1 ARPANET 网络层次结构：现代 Internet 结构的前奏

ARPANET 的主干由一些少量的核心路由器组成，并由一个单一的管理组织操作（Internet 网络操作中心）。大量的非核心路由器将 ARPANET 的客户与其主干连接起来，并由客户自己操作。这些非核心路由器通常指向一个核心路由器中的缺省路由，而它们本身并不是缺省的。换句话说，核心路由器包含 Internet 中每个网络的路由入口。

注释 非缺省（或自由缺省（default-free））路由器必须包含它需要访问的每个网络的显式路由。

在该核心内使用的路由协议是距离向量网关对网关协议（GGP: Gateway-to-Gateway Protocol）。GGP 是一种远距离路由协议，能够保证相邻节点之间路由选择更新的可靠性传输。GGP 遭受了早期远距离向量协议的不良收敛特性，但是到 1988 年时，ARPANET 核心路由已经升级到 BBN Butterfly，它运行称为 SPREAD 的早期的链路状态路由协议。

SPREAD 特别令人感兴趣的是它的作为路由计量的网络延时方面，它通过检查在链路接口中的传输队列来确定网络延时。这提供了对拥塞链路选择路由的性能——这些拥塞链路有很长的延时。当对拥挤链路动态选择路由时，必须采取步骤来避免路由摆动。因为网络变得越来越复杂，这个问题的解决并不容易，当今使用的两个主要的链路状态 IP 路由协议 IS-IS 和 OSPF 没有一个能够动态地时拥挤进行路由选择，这是非常重要的。

在某种程度上，SPREAD 路由协议的脆弱性破坏了链路状态的广告（SPREAD 使用链路状态广告 LSA 的循环序号空间）。通过限制 LSA 的生成率和它们的寿命可以避免在序列号回绕时出现的“哪个是最新的？”问题。不幸的是，设计者没有解决由序列号破坏而引起错误的硬件问题。这导致了由于 LSA 风暴的网络瘫痪，它将关闭整个 ARAP 网，并增强了链路状态路由协议的实现、配置和算法的关键性能。

ARPANET 在 1990 年遭淘汰了，与此同时，NSF 和其它联邦网络形成了 Internet 核心。

1.2.2 NSFNET（国家科学基金会网络）

1984 年，设计人员宣布了 JANET（联合科研网）的研制计划，JANET 是英国用于连接学术和研究团体的网络。1985 年，美国的国家科学基金会宣布了相似的计划，在接下来的几年里，其它一些国家也宣布了相似的计划。NSFNET 的设计分三个等级，形成单一自由缺省（或非缺省）核心的国家主干、区域分布的一组中间级网络和大量的校园访问网络组成了 NSFNET。在 1995 年退出之前，NSFNET 一直是 Internet 核心的主要成分。

56K bps 和 Fuzzball 路由器

NSFNET 的早期目的之一是使研究人员和科学家可以访问 NSF 大型计算机。其初始核心由连接到美国六个主要的大型计算机中心的 56K bps 的链路组成。每个大型计算机中心配备有 LSI-11 微型计算机，它运行昵称为“fuzzball”的路由代码和 HELLO 路由协议。

和传统的使用中继段数作为路由度量相比，远距离向量协议在包延迟的使用上也值得称道。

UNIX 路由选择软件用于相互重分布，包括 NSFNET 主干的 HELLO 协议和每个区域网的 IGP（通常为 RIP）之间的度量转换。通常，按照在发生错误配置时维护网络稳定性的预先安排的策略来过滤路由重分布。

Carnegie Mellon 大学同时维护了 NSFNET 的 fuzzball 和 ARPANET 的 PSN。事实上，CMU 成了第一个对等网络的 Internet 的交换和连接点。

图 1-2 所示的网络层次结构具有现代 Internet 层次结构的所有元素——核心网络、分布网络和 ISP 的访问网络——同时，每个 ISP（ARPANET 和 NSFNET）与早期的 Internet NAP（CMU, Pittsburgh）等效。

T1 和 NSS 路由器

在典型的 Internet 方式中，它自己的成功带来了坏处，很快地原有的网络主干达到了饱和。有关第二个 NSFNET 主干的建议的请求很快发布，而且被接受的建议发展成为 Merit 公司（一个在 Michigan 大学运作的区域网络）、IBM 和 MCI 的合作伙伴。

虽然合作伙伴中的所有成员都紧密地合作，但在网络操作、计算机网络设备开发和长距离通信服务方面形成了各种鲜明的观点。

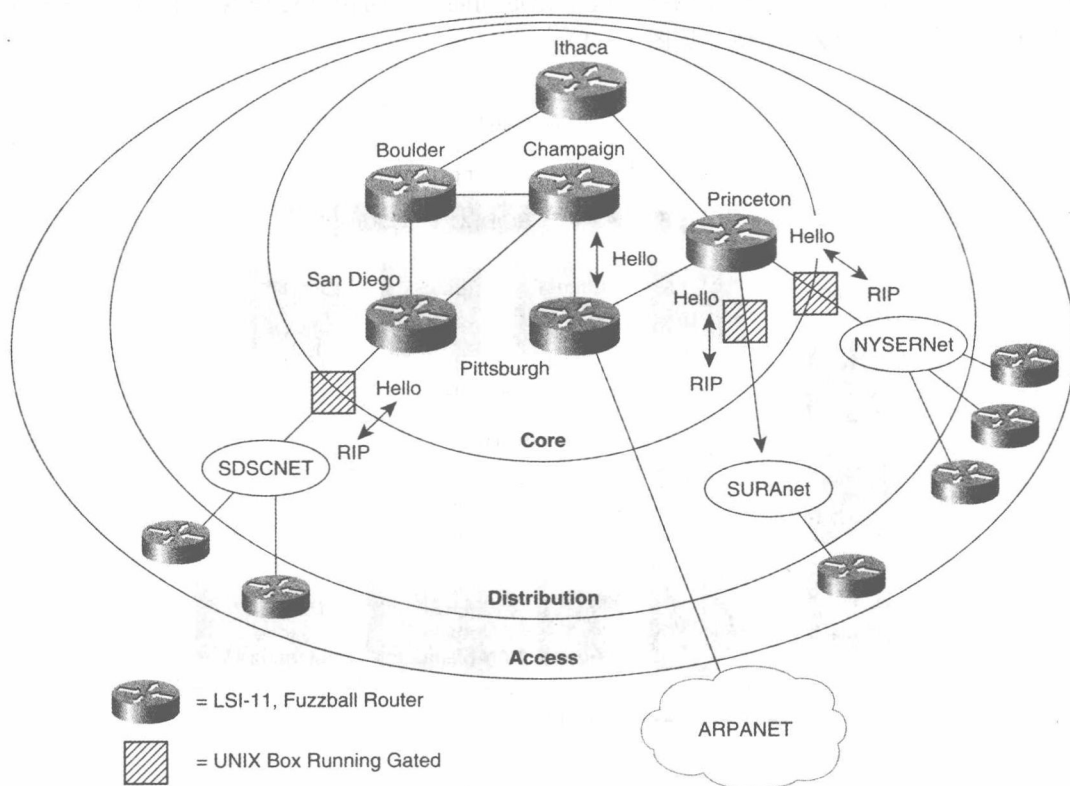


图 1-2 1998 升级前的原始 NSFNET 网（为简化起见，仅显示了七个区域网络中的三个）

第二个 NSFNET 主干更新了带有 13 个节点交换子系统（NSS）的原始主干的 fuzzball 路由器。每个 NSS 包含 14 个 IBM RT，这些 IBM RT 以双令牌环连接作为内部处理器通信的内部总线。（参见图 1-3）：

- 一个 RT 是路由选择和控制处理器。正如它的名字所表示的，这个处理器执行路由选择算法计算，创建 IP 路由表，并负责框的整体控制。
- 五个 RT 是包交换处理器。
- 四个 RT 包含 WAN 连接的线路卡（开始时为 448Kbps，后来达到 T1）
- 一个 RT 为外部 PSP，包含 LAN 连接的以太网卡。PSP 负责线路接口之间的包转发和 PSP 之间的调节平行交换的设计。
- 其余的三个 RT 作为后备系统。

每个 NSS 还包含两台 IBM PS/2 80。一个用于内部令牌环桥接管理器；另一个运行 NetView 和 LAN 管理器，并负责网络监测功能。

为 IP 修改的 OSI IS-IS 链路状态路由协议的 IBM 实现，（参见第 10 章）运行在组成 NSFNET 主干的 NSS 路由器上。使用内部模式的 EGP 实现了核心网和区域网之间的路由交换。EGP 是可接触性协议，而不是路由选择协议：它不对 EGP 更新中的网络的距离度量进行翻译以作出路由选择决定。因此，EGP 把 Internet 的拓扑结构限制为树状，NSFNET

为核心。路由阻尼 —— 一种在网络高度不稳定期间减少路由器处理器使用的方法 —— 在当时处于讨论中，但没有实现（见图 1-4）。

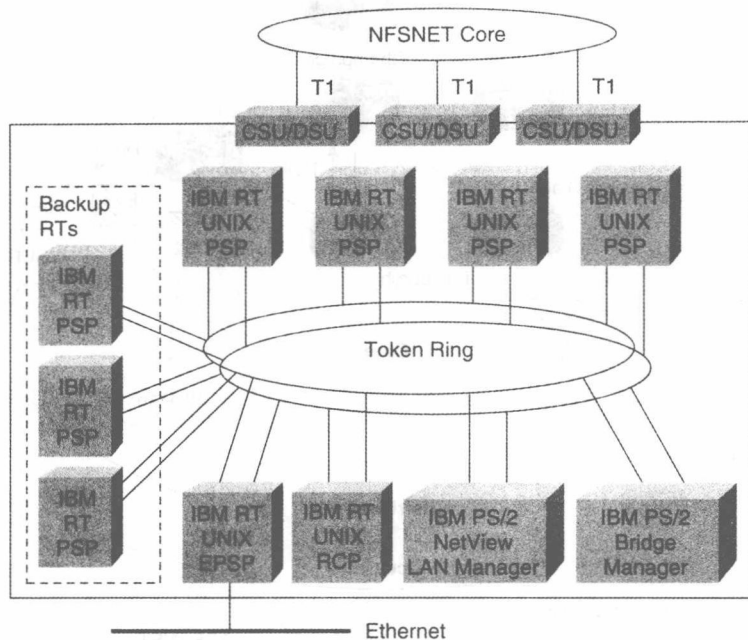


图 1-3 NSS 路由器

支持 NSFNET 的主干路由器是一个进化性的试验，它要求网络操作人员和代码开发者之间的紧密合作。事实证明，来自 Merit、IBM 和 MCI 的团队具有巨大的力量，因为它们提供了前进中的工程来支持每年以 500 个百分点增长的网络。

对 NSFNET 的快速增长，以及不依赖于 NSFNET 主干而直接相连的区域网络需求，促使 NSFNET 网操作人员提出了基本的策略路由选择系统。

其中，策略路由选择系统包括同时基于网络前缀（因为网络路由仍然分类，所以实际上是网络编号）和自治系统编号的对从区域网络到 NSFNET 作广告的所有网络的过滤。策略路由选择系统也设置所有接受的路由度量。这些功能是与分布策略路由选择数据库（PRDB）配合完成的。

注释 NSFNET 和 ARPA 与 MILNET（统称为防卫数据网）之间的路由规定了相似的策略。

这期间人们很好地认识到了 EGP 所存在的限制，为了解决这些限制，开始了替换域间路由协议的边界网关协议的设计工作。BGP 配置在 NSFNET 上，BGP 的版本 4 仍然是 Internet 的核心。

1990 年，在 NSFNET 核心中加入了一个新的 NSS 节点 Atlanta，使路由器总数达到 14 个。另外，Merit、IBM 和 MCI 成立了一个新的非赢利公司：Adanvanced Network and Service(高级网络和服务公司)。这个公司继续把 NSFNET 作为独立的实体进行运作。1991

年，ANS CO+RE 系统作为一个赢利企业脱颖而出。

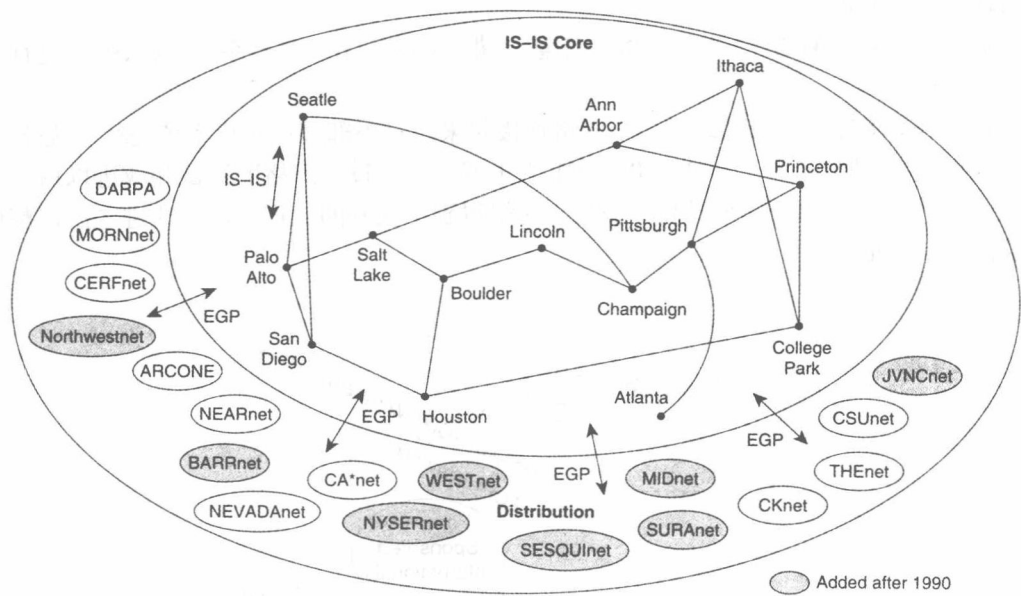


图 1—4 NSFNET 网的 T1 主干（灰色部分是 1990 年后加入的区域）

Internet 的管理和互操作性问题

1989 年底，Internet 拥有 300,000 多个主机和 2000 多个网络。当然地，用户开始关注网络管理技术；并开发了简单的网络管理协议（SNMP），提供了一种可扩展的查询和控制路由器以及与网络相连的实体方法。一些人认为 SNMP 是通往 OSI 通用管理信息协议的里程碑。SNMP 出现后得到了广泛的使用，而 CMIP 经历了与打字机相同的命运。

互操作性问题也成了人们关注的焦点——1988 年 8 月的 Interop 技术贸易展览上，供应商展示了他们的网络产品的性能和互操作性。这个展览会同时也成了人们交换观点的论坛，部分原因是因为供应商之间的竞争要稍微强于 IETF 标准会议。现在，这个展览会每年在美国举行一次，而其它的同类展览会合并到了欧洲和澳大利亚。

SNMP 在 1990 年的展览会上的唯一平台上第一次亮相：该平台为 Internet toster。

1. 2. 3 NSFNET—T3 和 ENSS 路由器

1993 年底，在 NSFNET 网络主干上的传输量增加到要求 45MB（T3）的连接。虽然那时 ARPANET 已经退役三年了，但是由于在区域网络方面的增长以及由“四大联邦网络”中的剩余网络（能源部、国防部和 NASA）操作的对等网络，事实上增加的 NSFNET 连接足于抵消 ARPANET 退役带来的损失。NSFNET 也鼓励学院和研究机构使用它的“连接”程序（参见图 1-5）。

NSFNET NSS 路由器也进行了重新设计以满足在包交换方面的大幅度增长的需求。新的路由器由一个配有 T3 接线卡的 IBM RS/6000 工作站组成。每个接线卡运行简化的