

第3章 局域网技术与综合布线

3.1 局域网基础

3.1.1 局域网参考模型

1980年2月, 电器和电子工程师协会(Institute of Electrical and Electronics Engineers, IEEE)成立了802委员会。当时个人计算机联网刚刚兴起, 该委员会针对这一情况制定了一系列局域网标准, 称为IEEE 802标准。按IEEE 802标准, 局域网体系结构由物理层、媒介访问控制子层(Media Access Control, MAC)和逻辑链路控制子层(Logical Link Control, LLC)组成, 如图3-1所示。

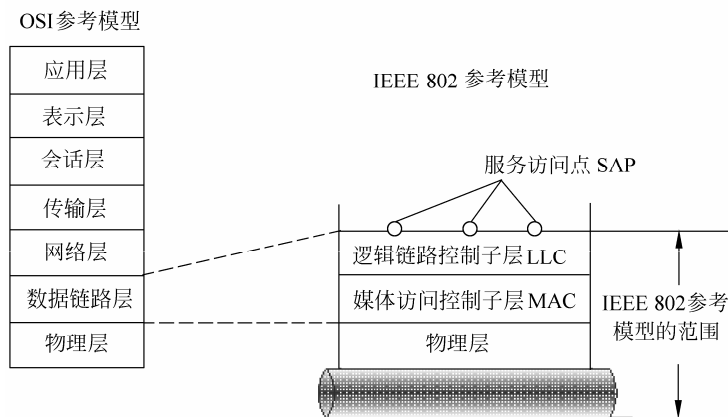
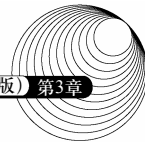


图 3-1 IEEE 802 参考模型

IEEE 802 参考模型的最低层对应于 OSI 模型中的物理层, 包括如下功能。

- (1) 信号的编码/解码。
- (2) 前导码的生成/去除(前导码仅用于接收同步)。
- (3) 比特的发送/接收。

IEEE 802 参考模型的 MAC 和 LLC 合起来对应 OSI 模型中的数据链路层, MAC 子层完成的功能如下。



(1) 在发送时将要发送的数据组装成帧, 帧中包含地址和差错检测等字段。

(2) 在接收时, 将接收到的帧解包, 进行地址识别和差错检测。

(3) 管理和控制对于局域网传输媒介的访问。

LLC 子层完成的功能如下。

(1) 为高层协议提供相应的接口, 即一个或多个服务访问点 (Service Access Point, SAP), 通过 SAP 支持面向连接的服务和复用能力。

(2) 端到端的差错控制和确认, 保证无差错传输。

(3) 端到端的流量控制。

需要指出的是, 局域网中采用了两级寻址, 用 MAC 地址标识局域网中的一个站, LLC 提供了服务访问点 (SAP) 地址, SAP 指定了运行于一台计算机或网络设备上的一个或多个应用进程地址。

目前, 由 IEEE 802 委员会制定的标准已近 20 个, 各标准之间的关系如图 3-2 所示。具体描述如下。

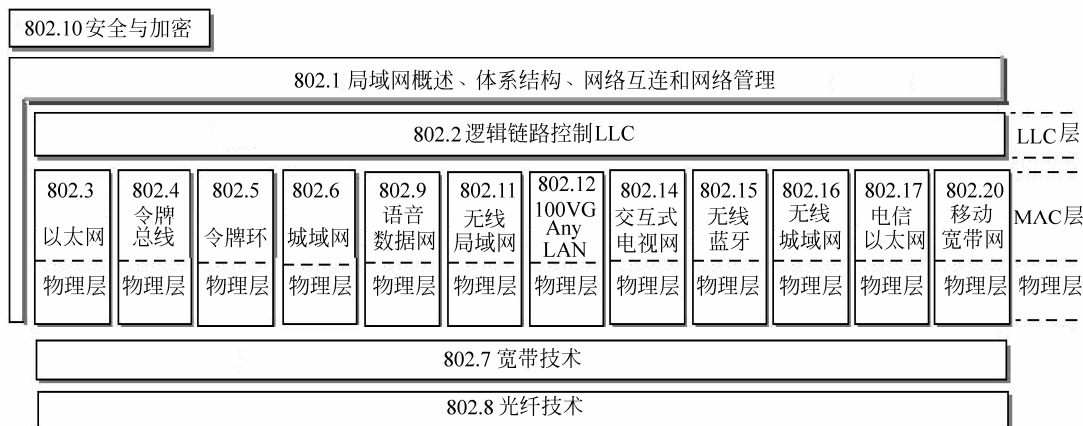
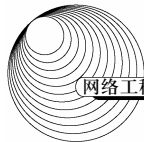


图 3-2 IEEE 802 参考模型各标准之间的关系

- 802.1: 局域网概述、体系结构、网络互连和网络管理。
- 802.2: 逻辑链路控制 (LLC)。
- 802.3: 带碰撞检测的载波侦听多路访问 (CSMA/CD) 方法和物理层规范 (以太网)。
- 802.4: 令牌传递总线访问方法和物理层规范 (Token Bus)。
- 802.5: 令牌环访问方法和物理层规范 (Token Ring)。
- 802.6: 城域网访问方法和物理层规范分布式队列双总线网 (DQDB)。
- 802.7: 宽带技术咨询和物理层课题与建议实施。



- 802.8: 光纤技术咨询和物理层课题。
- 802.9: 综合话音/数据服务的访问方法和物理层规范。
- 802.10: 互操作 LAN 安全标准 (SILS)。
- 802.11: 无线局域网 (wireless LAN) 访问方法和物理层规范。
- 802.12: 100VG ANY LAN 网。
- 802.14: 交互式电视网 (包括 cable modem)。
- 802.15: 简单, 低功耗无线连接的标准 (蓝牙技术)。
- 802.16: 无线城域网 (MAN) 标准。
- 802.17: 基于弹性分组环 (Resilient Packet Ring, RPR) 构建新型宽带电信以太网。
- 802.20: 3.5GHz 频段上的移动宽带无线接入系统。

3.1.2 局域网拓扑结构

拓扑是一种研究与大小、距离无关的几何图形特性的方法。在计算机网络中, 计算机作为节点, 传输媒介作为连线, 可构成相对位置不同的几何图形。网络拓扑结构是指用传输媒介互连各种设备形成的物理布局。参与 LAN 工作的各种设备用媒介互连在一起有多种方法, 不同连接方法的网络性能不同。按照不同的物理布局, 局域网拓扑结构通常分为三种, 分别是总线型拓扑结构、星型拓扑结构和环型拓扑结构。

1. 总线型拓扑结构

总线型拓扑结构是使用同一媒介或电缆连接所有端用户的一种方式, 也就是说, 连接端用户的物理媒介由所有设备共享, 如图 3-3 所示。使用这种结构必须解决的一个问题是确保端用户使用媒介发送数据时不会出现冲突。在点到点链路配置时, 这是相当简单的。如果这条链路是半双工操作, 只需要使用很简单的机制便可保证两个端用户轮流工作。在一点到多点方式中, 对线路的访问依靠控制端的探测来确定。然而, 在 LAN 环境下, 由于所有数据站都是平等的, 不能采取上述机制。为此, 一种在总线共享型网络使用的媒介访问方法, 即带有碰撞检测的载波侦听多路访问 (CSMA/CD) 应运而生。

这种结构具有费用低、数据端用户入网灵活、站点或某个端用户失效不影响其他站点或端用户通信的优点, 缺点是一次仅能有一个端用户发送数据, 其他端用户必须等到获得发送权, 媒介访问获取机制较复杂。尽管有上述一些缺点, 但由于布线要求简单, 扩充容易, 端用户失效和增删不影响全网工作, 所以这种结构是 LAN 技术中使用最普遍的一种。

2. 星型拓扑结构

星型拓扑结构存在中心节点, 每个节点通过点对点的方式与中心节点相连, 任何两个节点

之间的通信都要通过中心节点来接转。图 3-4 所示为目前使用最普遍的以太网星型拓扑结构,处于中心位置的网络设备称为集线器(Hub)。

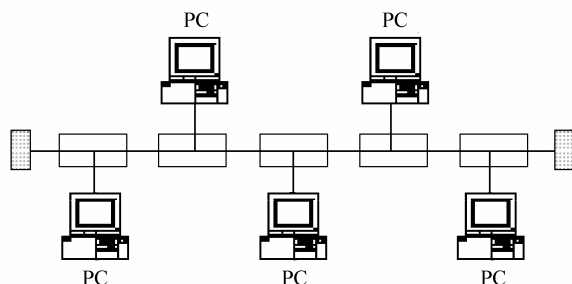


图 3-3 总线型拓扑结构

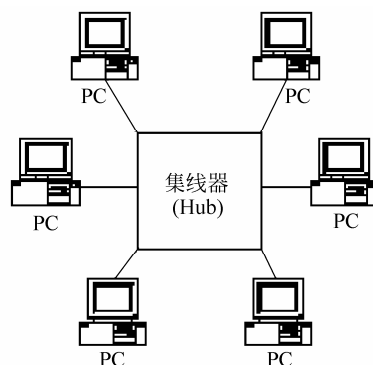


图 3-4 星型拓扑结构

这种结构便于集中控制,因为端用户之间的通信必须经过中心站。这一特点也带来了易于维护和安全等优点,端用户设备因为故障而停机时不会影响其他端用户间的通信。但这种结构非常不利的一点是中心系统必须具有极高的可靠性,因为中心系统一旦损坏,整个系统便会瘫痪。为此,中心系统通常采用双机热备份,以提高系统的可靠性。

3. 环型拓扑结构

环型拓扑结构在 LAN 中使用较多。这种结构中的传输媒介从一个端用户到另一个端用户,直到将所有端用户连成环,如图 3-5 所示。这种结构显然消除了端用户通信时对中心系统的依赖性。

环型拓扑结构的特点是每个端用户都与两个相邻的端用户相连,因而存在着点到点链路,构成闭合的环,但环中的数据总是沿一个方向绕环逐站传递。在环型拓扑中,多个节点共享一条环形通路,为了确定环中的节点在什么时候可以插入传送数据帧,同样要进行介质访问控制。因此,环型拓扑的实现技术中也要解决介质访问控制方法问题。与总线型拓扑一样,环型拓扑一般也采用某种分布式控制方法,环中的每个节点都要执行发送与接收控制逻辑信号。

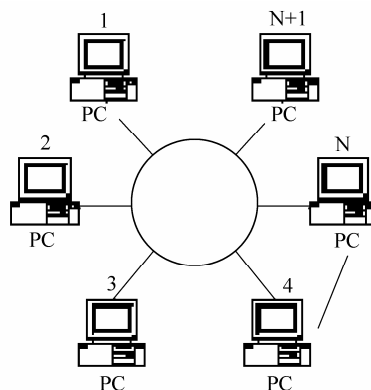
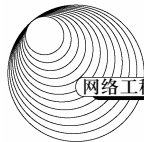


图 3-5 环型拓扑结构



3.1.3 局域网媒介访问控制方法

所有局域网均由共享该网络传输能力的多个设备组成。在网络中,服务器和计算机众多,每台设备随时都有发送数据的需求,这就涉及媒介的争用问题,所以需要有方法来控制设备对传输媒介的访问,以便两个特定的设备在需要时可以交换数据。传输媒介的访问控制方式与局域网的拓扑结构、工作过程有密切关系。目前,计算机局域网常用的访问控制方式有三种,分别是载波侦听多路访问/冲突检测(CSMA/CD)、令牌环访问控制法(Token Ring)和令牌总线访问控制法(Token Bus)。

1. CSMA/CD

CSMA/CD(Carrier Sense Multiple Access With Collision Detection)含有两方面的内容,即载波侦听(CSMA)和冲突检测(CD)。CSMA/CD访问控制方式主要用于总线型拓扑结构,是IEEE 802.3局域网标准的主要内容。CSMA/CD的设计思想如下所述。

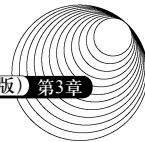
1) 载波侦听多路访问(CSMA)

各个站点都有一个“侦听器”,用来测试总线上有无其他工作站正在发送信息(也称为载波识别),一个站如果要发送数据,首先要侦听(监听)总线,查看信道上是否有信号,如果信道已被占用,则此工作站等待一段时间后再争取发送权;如果侦听总线是空闲的,没有其他工作站发送的信息,就立即抢占总线进行信息发送。查看信号的有无即为载波侦听。CSMA技术中要解决的另一个问题是侦听信道已被占用时如何确定等待多长时间。通常有两种方法,一种是当某工作站检测到信道被占用后,继续侦听下去,一直等到发现信道空闲后,立即发送,这种方法称为持续的载波侦听多点访问;另一种是当某工作站检测到信道被占用后,就延迟一个随机时间后再检测,不断重复这个过程,直到发现信道空闲后,开始发送信息,这称为非持续的载波侦听多点访问。

2) 冲突检测(CD)

当信道处于空闲时,某一个瞬间,如果总线上两个或两个以上的工作站同时想发送数据,那么该瞬间它们都可能检测到信道是空闲的,同时都认为可以发送信息,从而一齐发送,这就产生了冲突(碰撞)。另一种情况是某站点侦听到信道是空闲的,但这种空闲可能是较远站点已经发送了信息包,而由于在传输介质上信号传送的延时,信息包还未传送到此站点的缘故,如果此站点又发送信息,也将产生冲突。因此消除冲突是一个重要问题。

若在帧发送过程中检测到碰撞,则停止发送帧,即会形成不完整的帧(称“碎片”)在媒



介上传输,并随即发送一个 Jam (强化碰撞)信号以保证让网络上的所有站都知道已出现了碰撞。发送 Jam 信号后等待一段随机时间,再重新尝试发送。

在返回去重新发送帧之前,碰撞次数 n 加“1”递增(一开始 $n=0$),判断碰撞次数 n 是否达到 16 (十进制),若 $n=16$,则按“碰撞次数过多”差错处理,若 $n<16$,则计算一个随机量 r , r 的范围为 $0<r<2k$,其中 $k=\min(n, 10)$,即当 $n\geq 10$ 时 $k=10$,当 $n<10$ 时 $k=n$,获得延迟时间 $t=rT$ 。

其中 T 为常数,是网络上固有的一个参数,称为“碰撞槽时间”。延迟时间 t 又称“退避时间”,它表示检测到碰撞后要重新发送帧需要一段随机延迟时间,以避免发生碰撞各站的重新发送帧的时间。这种规则又称为“截短二进制指数退避”(truncated binary exponential backoff)规则,即退避时间是碰撞时间的 r 倍。

3) 碰撞槽时间

碰撞槽时间 (Slot time) 即是在帧发送过程中发生碰撞时间的上限,即在这段时间中可能检测到碰撞,而一过这段时间后永远不会发生碰撞,当然也不会检测到碰撞。也就是说,当发送的帧在媒介上传播时,如果超过了 Slot time,就再也不会发生碰撞,直到发送成功,或者说,一过这段时间,发送站就争用媒介成功。

为了帮助读者理解 Slot time,并进一步了解该参数的重要性,这里先分析检测一次碰撞需要多长时间。如图 3-6 所示,假设公共总线媒介长度为 S ,A 与 B 两个站点分别配置在媒介的两个端点上(即 A 与 B 站相距 S),帧在媒介上的传播速度为 $0.7C$ (C 为光速),网络的传输率为 R (bps),帧长为 L (bit)。图 3-6 (a) 表示 A 站正开始发送帧 fA,沿着媒介向 B 站传播;图 3-6 (b) 表示 fA 快到 B 站前一瞬间,B 站发送帧 fB;图 3-6 (c) 表示在 B 站处发生了碰撞,B 站立即检测到碰撞,同时碰撞信号沿媒介向 A 站回传;图 3-6 (d) 表示碰撞信号返回到 A 站,此时 A 站的 fA 尚未发送完毕,因此 A 站能检测到碰撞。从 fA 发送后直到 A 站检测到碰撞为止,这段时间间隔就是 A 站能够检测到碰撞的最长时间,这段时间一过,网络上就不可能发生碰撞,Slot time 的物理意义就是这样描述的,近似可以用以下公式近似表示。

$$\text{Slot time} \approx 2S/0.7C + 2t_{\text{PHY}}$$

其中, C 为光速, $0.7C$ 是信号在媒介上的传输速度, t_{PHY} 为 A 站物理层的延时,因为发送帧和检测碰撞都在 MAC 层进行,因此必须要加上 2 倍的物理层延时时间。

假设 A 站为了在 Slot time 上检测到碰撞至少要发送的帧长为 $L_{\min}$,因为 $L_{\min}/R = \text{Slot time}$,所以 $L_{\min} \approx (2S/0.7C + 2t_{\text{PHY}}) \times R$ 。 $L_{\min}$ 称为最小帧长度,由于碰撞只可能发生在小于或等于 $L_{\min}$ 时,因此 $L_{\min}$ 也可理解为媒介上传播的最大帧碎片长度。

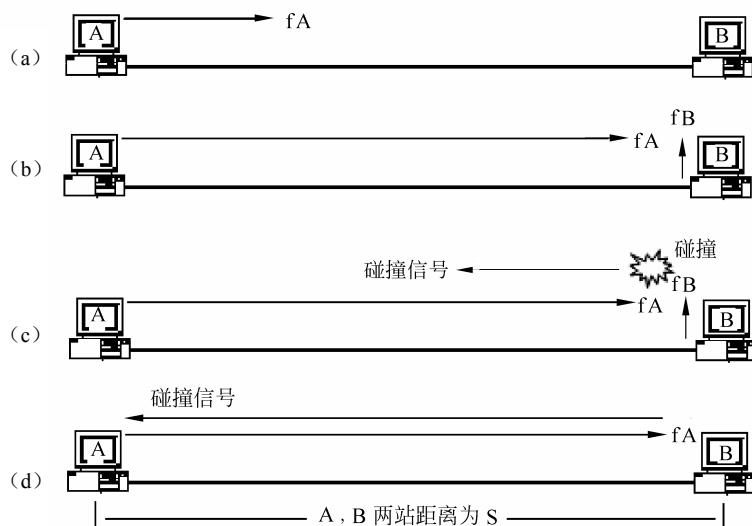
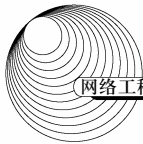


图 3-6 检测碰撞的最长时间

综上所述, Slot time 是 CSMA/CD 机理中一个极为重要的参数,这一参数描述了在发送帧的过程中处理碰撞的如下所述四个方面。

(1) 它是检测一次碰撞所需的最长时间。如果超过了该时间,媒介上的帧将再也不会遭到碰撞而损坏。

(2) 必须要求发送的帧长度有个下限限制,即所谓“最小帧长度”。最小帧长度能保证在网络最大跨距范围内,任何站在发送帧后,若碰撞产生,都能检测到。因为任何站要检测到碰撞必须在帧发送完毕之前,否则碰撞产生后可能漏检,造成传输错误。

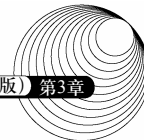
(3) 它是在碰撞产生后,决定了在媒介上出现的最大帧碎片长度。

(4) 作为碰撞后帧要重新发送所需的时间延迟计算的基准。

从公式 $L_{\min} \approx (2S/0.7C + 2t_{\text{PHY}}) \times R$ 可以知道,光速 C 和物理层延时 t_{PHY} 是常数,对于一个具有 CSMA/CD 的公共总线型(或树型)拓扑结构的局域网来说,公式中的其他 3 个参数 $L_{\min}$ 、 S 及 R 作为变量互为正、反比关系。例如,当传输率 R 固定时,最小帧长度与网络跨距具有正变的关系,即跨距越大, $L_{\min}$ 越长;当 $L_{\min}$ 不变时,传输率越高,跨距 S 越小。这些分析对以太网的性能和发展以及高速以太网的特点均有指导性意义。

4) 接收规则

在以太网结构中,发送节点需要通过竞争获得总线的使用权,而其他节点都应处于接收状态。当一个节点完成一组数据接收后,首先要判断接收帧的长度。因为 802.3 协议对帧的最小



长度做了规定,若接收帧长度小于规定帧的最小长度,则必然是冲突后的废弃帧。因此,如果帧太短,则表明冲突发生,接收节点丢弃已接收数据,并重新进入等待接收状态。如果没有发生冲突,接收节点会检查帧目的地址,如果目的地址为单一节点的物理地址,并且是本节点地址,则接收该帧;如果目的地址是组地址,而接收节点属于该组,则接收该帧;如果目的地址是广播地址,也应接收该帧;否则丢弃该接收帧。

如果接收节点进行地址匹配后确认应接收该帧,则下一步进行 CRC 校验。如果 CRC 校验正确,应进一步检查 LLC 数据长度是否正确。如果 LLC 数据长度正确,则 MAC 子层将帧中的 LLC 数据送往 LLC 子层,进入“成功接收”的结束状态。如果 LLC 数据长度不对,则进入“帧长度错”的结束状态。如果帧校验中发现错误,首先应判断接收帧的长度是不是 8 位的整数倍。如果帧长度是 8 位的整数倍,表示传输过程中没有发生比特丢失或对位错,此时应进入“帧校验错”结束状态;如果帧长度不是 8 位的整数倍,则进入“帧比特错”结束状态。

从以上讲解中可以看出,任何一个节点发送数据都要通过 CSMA/CD 方法去争取总线使用权,从它准备发送到成功发送的发送等待延时时间是不确定的。因此人们将以太网所使用的 CSMA/CD 方法定义为一种随机争用型介质访问控制方法。

CSMA/CD 方式的主要特点是原理比较简单、技术上较易实现、网络中各工作站处于同等地位、不需要集中控制,但这种方式不能提供优先级控制,各节点争用总线,不能满足远程控制所需要的确定延时和绝对可靠性的要求。另外,此方式效率高,但当负载增大时,发送信息的等待时间较长。

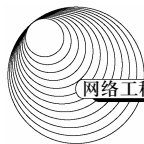
2. Token Ring

Token Ring 是令牌通行环(Token Passing Ring)的简写。其主要技术指标是网络拓扑为环形布局、基带网、数据传送速率 4Mbps、采用单个令牌(或双令牌)的令牌传递方法。环型拓扑结构网络的主要特点是只有一条环路、信息单向沿环流动、无路径选择问题。

令牌环技术的基础是使用了一个称为令牌的特定比特串,当环上的所有站都处于空闲时,令牌沿环传递,当某一站想发送帧时必须等待,直至检测到经过该站的令牌为止。这时该站改变令牌中的一个比特,从而抓住令牌,然后将令牌加在发送数据帧的帧首,变成发送数据帧。此时在环上不再有令牌,因此其他想发送帧的站必须等待。这个发送数据帧将在环上环行一周,然后由发送站将其清除。在下列两个条件都符合时,发送站将在环上释放一个新的令牌。

- (1) 本站已完成帧的发送。
- (2) 本站所发送的帧的前沿已回到本站(帧已绕环运行一整圈)。

其他站在环上监听,并不断地转发经过的帧。每个站都能对经过的帧进行检测,如果检测到的目的地址是本身的地址,并且本站有足够的缓存,就复制该帧,从而达到接收的目的,同时继续转发该信息包。环上的帧信息绕网一周后,由源发送点予以收回。按这种方式工作,发



送权会一直在源站点控制之下, 只有发送信息包的源站点放弃发送权, 把 Token 置“空”后, 其他站点才有机会得到令牌发送自己的信息。

令牌方式在轻负载时, 由于发送信息之前必须等待令牌, 加上规定由源站收回信息, 大约有 50% 的环路在传送无用信息, 所以效率较低。然而在重负载环路中, 令牌以“循环”方式工作, 故效率较高, 各站机会均等。令牌环的主要优点在于它提供的访问方式的可调整性, 它可提供优先权服务, 具有很强的实时性。其主要缺点是有令牌维护要求, 为避免令牌丢失或令牌重复, 使得这种方式的控制电路较复杂。

3. Token Bus

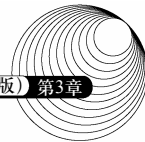
Token Bus 是 Token Passing Bus (令牌通行总线) 的简写。这种方式主要用于总线型或树型拓扑结构网络中。1976 年, 美国 Data Point 公司研制成功的 ARCnet (Attached Resource Computer) 网络综合了令牌传递方式和总线网络的优点, 在物理总线结构中实现了令牌传递控制方法, 从而构成一个逻辑环路。此方式也是目前微机局域中的主流介质访问控制方式。

ARCnet 网络就是使总线型或树型传输介质上的各工作站形成一个逻辑上的环, 即将各工作站置于一个顺序的序列内 (例如可按照接口地址的大小排列)。方法可以是在每个站点中设一个网络节点标识寄存器 NID, 初始地址为本站点地址, 网络工作前要对系统初始化, 以形成逻辑环路, 其过程主要是: 网中最大站号 n 开始向其后继站发送“令牌”信息包, 目的站号为 $n+1$, 若在规定时间内收到肯定的信号 ACK, 则 $n+1$ 站连入环路; 否则再加 1, 继续向下询问 (该网中最大站号为 $n=255$, $n+1$ 后变为 0, 然后继续 1、2、3、……, 凡是给予肯定回答的站都可连入环路, 并将给予肯定回答的后继站号放入本站的 NID, 从而形成一个封闭逻辑环路, 经过一遍轮询过程后, 网络各站标识寄存器 NID 中存放的都是其相邻的下游站地址。

逻辑环形成后, 令牌的控制方法类似于 Token Ring。在 Token Bus 中, 信息是按双向传送的, 每个站点都可以“听到”其他站点发出的信息, 所以令牌传递时都要加上目的地址, 明确指出下一个控制站点。这种方式与 CSMA/CD 方式的不同在于, 除了当时得到令牌的工作站之外, 其余所有工作站只收不发, 只有当收到令牌后才能开始发送。所以, 虽然拓扑结构是总线型结构, 但可以避免冲突。

Token Bus 方式的最大优点是具有极好的吞吐能力, 且吞吐量随数据传输速率的增高而增加, 并随介质的饱和而稳定下来, 但并不下降; 各工作站不需要检测冲突, 故信号电压允许较大的动态范围, 连网距离较远; 有一定实时性, 在工业控制中得到了广泛应用, 如 MAP 网就是采用宽带令牌总线。其主要缺点在于其复杂性和时间开销较大, 工作站可能必须等待多次无效的令牌传送后才能获得令牌。

应该指出, ARCnet 网实际上采用称为集中器的 HUB 硬件联网, 物理拓扑有星型和总线型两种连接方式。



3.1.4 无线局域网简介

21 世纪迎来了信息时代,网络已经渗透到了个人、企业以及政府。现在的网络建设已经发展到无所不在,不论用户在任何时间还是任何地点,都可以轻松上网。网络无所不在其实并不简单,光靠光纤、铜缆是不够的,毕竟在许多场合不允许铺设线缆。因此,需要推广一种新的解决方案,使网络的无所不在能够得以实现,这种解决方案就是无线数据网络。

1. 无线数据网络种类

无线数据网络解决方案包括无线个人网、无线局域网、无线城域网和无线广域网。

1) 无线个人网 (Wireless Personal Area Network, WPAN)

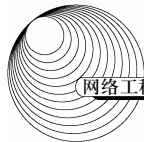
WPAN 主要用于个人用户工作空间,典型距离覆盖几米,可以与计算机同步传输文件,可以访问本地外围设备,如打印机等。WPAN 通常形象描述为“最后 10 米”的通信需求,目前主要技术为蓝牙 (Bluetooth)。

蓝牙技术源于 1994 年 Ericsson 提出的无线连线与个人接入的想法。1997 年 Ericsson、IBM、INTEL、NOKIA 和 TOSHIBA 商议建立一种全球化的无线通信个人接入与无线连线新手段,定名为“蓝牙”(Bluetooth)。Bluetooth 是一位在 10 世纪统一了丹麦和挪威的丹麦国王的名字,发明者无疑希望蓝牙技术也能够像这位国王一样,把移动电话、笔记本电脑和手持设备紧密地结合在一起。1998 年 5 月“蓝牙特别兴趣组织”BSIG (Bluetooth Special Interest Group) 正式发起成立,简称蓝牙 SIG。同期,1998 年 3 月在 IEEE 802.11 项目组中,对 WPAN 感兴趣的人士成立了研究小组,命名为 IEEE 802.15 工作组,主要工作是在 WPAN 内对无线媒介接入控制 (MAC) 和物理层 (PHY) 进行规范。为了保持两个标准的互操作性,蓝牙 SIG 采纳了 WPAN 的标准,即 IEEE 802.15 标准。这样蓝牙 1.0 版本可以达到与 802.15 之间 100% 的互操作性。

1999 年 11 月, Motorola、Lucent、Microsoft 及 3Com 加盟 BSIG,成为 BSIG 的发起成员,使蓝牙技术的发展获得了更强有力的支持,并显示出更明朗的前景。现今,BSIG 的参加成员已达 2500 多个,其发展势头令人瞩目。目前,蓝牙信道带宽为 1MHz,异步非对称连接最高数据速率达 723.2Kbps,连接距离多半为 10m 左右。为了适应未来宽带多媒体业务的需求,蓝牙速率亦拟进一步增强,新的蓝牙标准 2.0 版拟支持高达 10Mbps 以上的速率 (4 Mbps、8 Mbps、12Mbps、20Mbps)。

2) 无线局域网 (Wireless LAN, WLAN)

WLAN,顾名思义是一种借助无线技术取代以往的有线布线方式构成局域网的新手段。WLAN 可提供传统有线局域网的所有功能,是计算机网络与无线通信技术相结合的产物。WLAN 利用射频无线电或红外线,借助直接序列扩频或跳频扩频、GMSK、OFDM 等技术,甚



至将来的超宽带传输技术 UWBT, 实现固定、半移动及移动的网络终端对互联网进行较远距离的高速连接访问, 支持的传输速率为 2~54Mbps。WLAN 通常形象描述为“最后 100 米”的通信需求, 如企业网和驻地网等。

1997 年 6 月, IEEE 推出了 802.11 标准, 开创了 WLAN 先河。目前, WLAN 领域主要是 IEEE 802.11x 系列。IEEE 802.11 是 1997 年 IEEE 最初制定的一个 WLAN 标准, 主要用于解决办公室无线局域网和校园网中用户终端的无线接入, 其业务范畴主要限于数据存取, 速率最高只能达 2Mbps。由于它在速率、传输距离、安全性、电磁兼容能力及服务质量方面均不尽如人意, 从而产生了其系列标准, IEEE 802.11x 系列标准中应用最广泛的是 802.11b。802.11b 将速率扩充至 11Mbps, 并可在 5.5Mbps、2Mbps 及 1Mbps 之间进行自动速率调整, 亦提供了 MAC 层的访问控制和加密机制, 从而达到了与有线网络相同级别的安全保护, 还提供了可供选择的 40 位及 128 位的共享密钥算法, 从而成为目前 IEEE 802.11 系列的主流产品。而 802.11b+还可将速率增强至 22Mbps。802.11a 工作在 5GHz 频带, 数据传输速率将提升到 54Mbps。

目前, IEEE 802.11 系列得到了许多半导体器件制造商的支持, 这些制造商成立了一个无线保真联盟 Wi-Fi (Wireless Fidelity)。Wi-Fi 实质上是一种商业认证, 表明具有 Wi-Fi 认证的产品要符合 IEEE 802.11 无线网络规范。无疑, Wi-Fi 为 802.11 标准的推广起到了积极的促进作用。

3) 无线城域网 (Wireless MAN, WMAN)

WMAN 是一种有效作用距离比 WLAN 更远的宽带无线接入网络, 通常用于城市范围内的业务点和信息汇聚点之间的信息交流和网际接入, 有效覆盖区域为 2~10km, 最大可达 30km, 数据传输速率最快可高达 70Mbps。目前主要技术为 IEEE 802.16 系列。

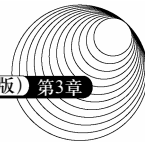
IEEE 802.16 标准于 2001 年 12 月获得批准, 其主题为“Air Interface For Fixed Broadband Wireless Access System”, 即“宽带固定无线接入系统的空中接口”。IEEE 802.16 标准对无线接入设备的媒介接入控制层和物理层制定了技术规范, 可支持 1~2GHz、10GHz 以及 12~66GHz 等多个无线频段。

借鉴于 Wi-Fi 模式, 一个同样由多个顶级制造商组成的全球微波接入互操作联盟 WiMax (Wireless Interoperability Microwave Access) 宣告成立。WiMax 的目标是帮助推动和认证采用 IEEE 802.16 标准的器件和设备具有兼容性和互操作性, 促进这些设备的市场推广。

4) 无线广域网 (Wireless WAN, WMAN)

无线广域网主要解决超出一个城市范围的信息交流无线接入需求。IEEE 802.20 和 3G 蜂窝移动通信系统构成了 WMAN 的标准。

2002 年 11 月, IEEE 802 标准委员会成立了 IEEE 802.20 工作组, 即移动宽带无线接入 (Mobil Broadband Wireless Access, MBWA) 工作组, 其主要任务是制定适用于各种工作在 3.5GHz 频段上的移动宽带无线接入系统公共空中接口的物理层和媒介访问控制层的标准协议。



这个标准初步规划是为以 250km/h 速度前进的移动用户提供高达 1Mbps 的高带宽数据传输,这将为高速移动用户创造使用视频会议等对带宽和时间敏感的应用的条件。拟议中的 802.20 标准的覆盖范围同现在的移动电话系统一样,都是全球范围的,而传输速度却达到了 Wi-Fi 水平,与现在的移动通信网络相比具有明显的优势。

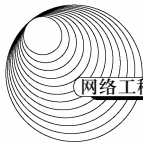
ITU 早在 1985 年就提出工作在 2GHz 频段的移动商用系统为第三代移动通信系统,国际上统称为 IMT-2000 系统(International Mobile Telecommunications-2000),简称 3G(3rd Generation)。ITU 所设定的 3G 标准的主要特征包括国际统一频段、统一标准;实现全球的无缝漫游;提供更高的频谱效率、更大的系统容量,是目前 2G 技术的 2~5 倍;提供移动多媒体业务。其设计目标为高速移动环境支持 144Kbps,步行慢速移动环境支持 384Kbps,室内环境下支持 2Mbps 的数据传输,从而为用户提供包括话音、数据及多媒体等在内的多种业务。3G 的三大主流无线接口标准分别是 W-CDMA、CDMA2000 和 TD-SCDMA,其中,W-CDMA 标准主要起源于欧洲和日本;CDMA2000 系统主要是由美国高通北美公司为主导提出的;时分同步码分多址接入标准 TD-SCDMA 由中国提出,并在此无线传输技术(RTT)的基础上与国际合作,完成了 TD-SCDMA 标准,成为 CDMA TDD 标准的一员,这是中国移动通信界的一次创举,也是中国对第三代移动通信发展的贡献。在与欧洲、美国各自提出的 3G 标准的竞争中,中国提出的 TD-SCDMA 已正式成为全球 3G 标准之一,这标志着中国在移动通信领域已经进入世界领先之列。

2. 无线局域网扩频技术

无线局域网采用电磁波作为载体传送数据信息。对电磁波的使用有两种常见模式,即窄带和扩频。窄带微波(Narrowband Microwave)技术适用于长距离点到点的应用,可以达到 40km,最大带宽可达 10Mbps;但受环境干扰较大,不适合用来进行局域网数据传输。所以目前无线局域网的数据传输通常采用无线扩频技术(Spread Spectrum, SST)。

常见的扩频技术包括跳频扩频(Frequency-Hopping Spread Spectrum, FHSS)和直接序列扩频(Direct Sequence Spread Spectrum, DSSS)两种,它们工作在 2.4~2.4835GHz。这个频段称为 ISM 频段(Industrial Scientific Medical Band),主要开放给工业、科学、医学三方面使用。该频段是依据美国联邦通信委员会(FCC)定义出来的,在美国属于免执照(Free License),并没有使用授权的限制。

跳频技术将 83.5MHz 的频带划分成 79 个子频道,每个频道带宽为 1MHz。信号传输时在 79 个子频道间跳变,因此传输方与接收方必须同步,获得相同的跳变格式,否则接收方无法恢复正确的信息。跳频过程中如果遇到某个频道存在干扰,将绕过该频道。由于受跳变的时间间隔和重传数据包的影响,跳频技术的典型带宽限制为 2~3Mbps。无线个人网采用的蓝牙技术就是跳频技术,该技术提供非对称数据传输,一个方向速率为 720Kbps,另一个方向速率仅为



57Kbps。蓝牙技术也可以传送 3 路双向 64Kbps 的语音。

直接序列扩频技术是无线局域网 802.11b 采用的技术,将 83.5MHz 的频带划分成 14 个子频道,每个频道带宽为 22MHz。直接序列扩频技术用一个冗余的位格式来表示一个数据位,这个冗余的位格式称为 chip,因此它可以抗拒窄带和宽带噪音的干扰,提供更高的传输速率。直接序列扩频技术(DSSS)提供的最高带宽为 11Mbps,并且可以根据环境因素的限制自动降速至 5.5Mbps、2Mbps、1Mbps。

3. 无线局域网拓扑结构

无线局域网组网分两种拓扑结构,即对等网络和结构化网络。

对等网络(Peer to Peer)用于一台计算机(无线工作站)和另一台或多台计算机(其他无线工作站)的直接通信,该网络无法接入有线网络,只能独立使用。对等网络中的一个节点必须能“看”到网络中的其他节点,否则就认为网络中断,因此对等网络只能用于少数用户的组网环境,比如 4~8 个用户,并且他们离得足够近。

结构化网络(Infrastructure)由无线访问点 AP(Access Point)、无线工作站 STA(Station)以及分布式系统(DSS)构成,覆盖的区域分基本服务区(Basic Service Set, BSS)和扩展服务区(Extended Service Set, ESS)。

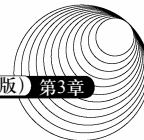
无线访问点也称无线集线器,用于在无线工作站(STA)和有线网络之间接收、缓存和转发数据。无线访问点通常能够覆盖几十至几百用户,覆盖半径达上百米。基本服务区由一个无线访问点以及与其关联的无线工作站构成,在任何时候,任何无线工作站都与该无线访问点关联。一个无线访问点所覆盖的微蜂窝区域就是基本服务区。无线工作站与无线访问点关联采用 AP 的基本服务区标识符(BSSID),在 802.11 中,BSSID 是 AP 的 MAC 地址。扩展服务区是指由多个 AP 以及连接它们的分布式系统组成的结构化网络,所有 AP 必须共享同一个扩展服务区标识符(ESSID),也可以说扩展服务区 ESS 中包含多个 BSS。

无线局域网产品中的楼到楼网桥(Building to Building Bridge)为难以布线的场点提供了可靠、高性能的网络连接。使用无线楼到楼网桥可以得到高速度、长距离的连接,事实上,可以得到超过两路 T1 线路的流量。无线楼到楼网桥可以提供点到点、点到多点的连接方式,用户可以选择最符合需求的天线,如传输近距离的全向性天线或传输远距离的扇形指向性天线。

4. 无线局域网的几个主要工作过程

1) 扫频

STA 在加入服务区之前要查找哪个频道有数据信号,分主动和被动两种方式。主动扫频是指 STA 启动或关联成功后扫描所有频道;一次扫描中,STA 采用一组频道作为扫描范围,如果发现某个频道空闲,就广播带有 ESSID 的探测信号;AP 根据该信号做响应。被动扫频是指



AP 每 100ms 向外传送灯塔信号, 包括用于 STA 同步的时间戳、支持速率以及其他信息, STA 接收到灯塔信号后启动关联过程。

2) 关联 (Associate)

关联过程用于建立无线访问点和无线工作站之间的映射关系, 实际上是把无线网变成有线网的连线。分布式系统将该映射关系分发给扩展服务区中的所有 AP。一个无线工作站同时只能与一个 AP 关联。在关联过程中, 无线工作站与 AP 之间要根据信号的强弱协商速率, 速率变化包括 11Mbps、5.5Mbps、2Mbps 和 1Mbps。

3) 重关联 (Reassociate)

重关联就是当无线工作站从一个扩展服务区中的一个基本服务区移动到另外一个基本服务区时, 与新的 AP 关联的整个过程。重关联总是由移动无线工作站发起。

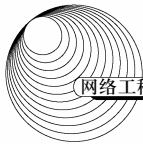
4) 漫游 (Roaming)

漫游指无线工作站在一组无线访问点之间移动, 并提供对于用户透明的无缝连接, 包括基本漫游和扩展漫游。基本漫游是指无线 STA 的移动仅局限在一个扩展服务区内部。扩展漫游是指无线 STA 从一个扩展服务区中的一个 BSS 移动到另一个扩展服务区中的一个 BSS。802.11 并不保证这种漫游的上层连接, 常见做法是采用 Mobile IP 或动态 DHCP。

5. 无线局域网的访问控制方式

802.3 标准的以太网使用 CSMA/CD 访问控制方法。在这种介质访问机制下, 准备传输数据的设备首先检查载波通道, 如果在一定时间内没有侦听到载波, 那么这个设备就可以发送数据。如果两个设备同时发送数据, 冲突就会发生, 并被所有冲突设备所检测到。这种冲突便延缓了这些设备的重传, 使得它们在间隔某一随机时间后才发送数据。而 802.11b 标准的无线局域网使用的是带冲突避免的载波侦听多路访问方法(CSMA/CA), 冲突检测(Collision Detection)变成了冲突避免(Collision Avoidance)。因为在无线传输中侦听载波及冲突检测都是不可靠的, 侦听载波有困难。另外, 通常无线电波经天线送出去时, 自己是无法监视到的, 因此冲突检测实质上也不做到。在 802.11 中侦听载波由两种方式来实现, 一个是实际去听是否有电波在传, 然后加上优先权控制; 另一个是虚拟的侦听载波, 告知大家待会有多久的时间我们要传东西, 以防止冲突。

CSMA/CA 访问控制方式将时间域的划分与帧格式紧密联系起来, 保证某一时刻只有一个站点发送数据, 实现了网络系统的集中控制。因传输媒介不同, CSMA/CD 与 CSMA/CA 的检测方式也不同。CSMA/CD 通过电缆中电压的变化来检测, 当数据发生碰撞时, 电缆中的电压就会随之发生变化; 而 CSMA/CA 采用能量检测(ED)、载波检测(CS)和能量载波混合检测三种方式检测信道是否空闲。



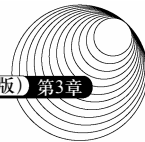
3.2 以太网

3.2.1 以太网简介

以太网(Ethernet)是 Xerox 公司在 1972 年开创的。1972 年秋,一位刚从麻省理工学院毕业的学生 Bob Metcalfe 来到 Xerox palo Alto 研究中心(PARC)计算机实验室工作, Metcalfe 的第一件工作是把 Xerox ALTO 计算机连到 ARPANet 上(ARPANet 是现在的 Internet 的前身)。在访问 ARPANet 的过程中,他偶然发现了 ALOHA 系统(这是一个源于夏威夷大学的地面无线电广播系统,其核心思想是共享数据传输信道)的一篇文章, Metcalfe 认识到,通过优化就可以把 ALOHA 系统的速率提高到 100%。1972 年底, Metcalfe 和 David Boggs 设计了一套网络,把不同的 ALTO 计算机连接起来。Metcalfe 把他的这一研究性工作命名为 ALTO ALOHA。1973 年 5 月 22 日,世界上第一个个人计算机局域网 ALTO ALOHA 投入了运行,这一天, Metcalfe 写了一段备忘录,称他已将该网络改名为以太网(Ethernet),其灵感来自于“电磁辐射是可以通过发光的以太来传播的”这一想法。最初的以太网以 2.94Mbps 的速度运行,运行速度慢,原因是以太网的接口定时采用 ALTO 系统时钟,即每 340ns 才发送一个脉冲。当然,因为以太网的核心思想是使用共享的公共传输信道,在公共传输信道上进行载波监听,这已比初始的 ALOHA 网络有了巨大的改进,经过一段时间的研究与发展,1976 年,以太网已经发展到能够连接 100 个用户节点,并在 1000m 长的粗缆上运行。由于 Xerox 急于将以太网转化为产品,因此将以太网改名为 Xerox Wire。1976 年 6 月, Metcalfe 和 Boggs 发表了题为“以太网:局域网的分布型信息包交换”的著名论文,1977 年底, Metcalfe 和他的三位合作者获得了“具有冲突检测的多点数据通信系统”的专利,多点传输系统被称为 CSMA/CD(载波监听多路访问/冲突检测)。从此,以太网就正式诞生了。

1979 年,在 DEC、Intel 和 Xerox 共同将此网络标准化时,也将 Xerox Wire 网络又恢复成“以太网”这个原来的名字。1980 年 9 月,三方公布了第三稿“以太网:一种局域网的数据链路层和物理层规范 1.0 版”,这就是著名的以太网蓝皮书,也称 DIX(DEC、Intel、Xerox 的第一个字母)版以太网 1.0 规范,一开始规范规定在 20MHz 下运行,经过一段时间后降为 10MHz,并重新定义了 DIX 标准,并以 1982 年公布的以太网 2.0 版规范终结。1983 年,以太网技术(802.3)与令牌总线(802.4)和令牌环(802.5)共同成为局域网领域的三大标准。1995 年,IEEE 正式通过了 802.3u 快速以太网标准,以太网技术实现了第一次飞跃。1998 年,802.3z 千兆以太网标准正式发布,2002 年 7 月 18 日正式通过了万兆以太网标准 802.3ae。

从 20 世纪 80 年代开始,以太网就成为最普遍采用的网络技术,它一直“统治”着世界各



地的局域网和企业骨干网,并且正在向城域网发起攻击。根据 IDC 的统计,以太网的端口数约为所有网络端口数的 85%,而且以太网这种强大的优势仍然有继续保持下去的势头。纵观以太网的强劲发展历程,可以发现以太网主要得益于以下几个特点。

(1) 开放标准,获得众多厂商的支持。目前,几乎所有硬件制造商生产的设备以及几乎所有软件开发商开发的操作系统和应用协议都与以太网兼容。

(2) 易于移植和升级,可最大限度保护用户投资。对于所有以太网技术,其帧的结构几乎是一样的,这就提供了一个非常好的升级途径。快速以太网技术提供了从 10M 向 100M 以太网的平滑升级。千兆和万兆以太网的出现,增加带宽的同时也扩展了可升级性。只要将低速以太网设备用交换机连接到千兆和万兆以太网设备上,就可实现一个物理线速向另一个物理线速的适配。这样的升级方式就使得千兆和万兆能无缝地与现在的以太网集成。

(3) 价格便宜,管理成本低。无论在局域网、接入网还是即将进入的城域网、广域网,以太网技术在价格上与其他技术相比都具有优越性。若全面采用以太网解决方案,价格将更具有吸引力。另外,以太网存在时间长,标准化程度高,一般网络管理人员都比较熟悉,因此它的运行维护管理成本也比较低。

(4) 结构简单,组网方便。以太网技术的实现原理统一采用了 CSMA/CD 介质访问控制方法,不同版本以太网的帧结构和网络拓扑结构也是一致的,对布线系统的要求较低,网络连接设备的配置比较简单。

3.2.2 以太网综述

1. 10M 以太网

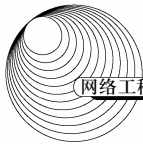
根据传输媒介的不同,10M 以太网大致有 4 个标准,各个标准的 MAC 子层媒介访问控制方法和帧结构以及物理层的编码译码方法(曼彻斯特编码)均是相同的,不同的是传输媒介和物理层的收发器及媒介连接方式。依照技术出现的时间顺序,这 4 个标准依次如下。

1) 10Base 5

1983 年,IEEE 802.3 工作组发布 10Base 5 “粗缆”以太网标准,这是最早的以太网标准。10Base 5 以太网传输媒介采用 $\phi 10$ 、50 Ω 粗同轴电缆,拓扑结构为总线型,电缆段上工作站之间的距离为 2.5m 的整数倍,每个电缆段内最多只能有 100 台终端,但每个电缆段不能超过 500m。网络设计遵循“5-4-3”法则,根据该法则,整个网络的最大跨距为 2500m。

- “5”即是网络中任意两个端到端的节点之间最多只能有 5 个电缆段。
- “4”即是网络中任意两个端到端的节点之间最多只能有 4 个中继器。
- “3”即是网络中任意两个端到端的节点之间最多只能有 3 个共享网段。

10Base 5 代表的具体意思是:工作速率为 10Mbps,采用基带信号,每一个网段最长为 500m。



2) 10Base 2

1986年, IEEE 802.3工作组发布 10Base 2“细缆”以太网标准。10Base 2以太网传输媒介采用 $\phi 5$ 、50 Ω 粗同轴电缆, 拓扑结构为总线型, 电缆段上工作站之间的距离为0.5m的整数倍, 每个电缆段内最多只能有30台终端, 但每个电缆段不能超过185m。10Base2以太网设计遵循“5-4-3”法则, 整个网络的最大跨距为925m。

10Base 2代表的具体意思是: 工作速率为10Mbps, 采用基带信号, 每一个网段最长约为200m。

3) 10Base T

1991年, IEEE 802.3工作组发布 10Base T“非屏蔽双绞线”以太网标准。10Base T以太网传输媒介采用100 Ω UTP双绞线, 拓扑结构为星型, 所有站点均连接到一个中心集线器(Hub)上, 但每个电缆段不能超过100m。10Base T以太网设计遵循“5-4-3”法则, 整个网络的最大跨距为500m。

10Base T代表的具体意思是: 工作速率为10Mbps, 采用基带信号, T表示传输媒介为双绞线(Twisted pair)。

4) 10Base F

1993年, IEEE 802.3工作组发布 10Base F“光纤”以太网标准。10Base F以太网传输媒介采用多模光纤, 拓扑结构为星型, 所有站点均连接到一个支持光纤接口的中心集线器上, 每个电缆段不能超过2000m。10Base F以太网设计也遵循“5-4-3”法则, 但由于受CSMA/CD碰撞域的影响, 整个网络的最大跨距为4000m。

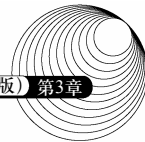
10Base F代表的具体意思是: 工作速率为10Mbps, 采用基带信号, F表示传输媒介光纤(Fiber)。

2. 100M 以太网

1995年, IEEE通过了802.3u标准, 将以太网的带宽扩大为100Mbps。从技术角度上讲, 802.3u并不是一种新的标准, 只是对现存802.3标准的升级, 习惯上称为快速以太网。其基本思想很简单, 即保留所有旧的分组格式、接口以及程序规则, 只是将位时从100ns减少到10ns, 并且所有快速以太网系统均使用集线器。快速以太网除了继续支持在共享媒介上的半双工通信外, 1997年, IEEE通过了802.3x标准后, 还支持在两个通道上进行双工通信。双工通信进一步改善了以太网的传输性能。另外, 100M以太网的网络设备的价格并不比10M以太网的设备贵多少。100Base-T以太网在近几年的应用得到了非常快速的发展。

1) 100Base-T4

100Base-T4传输载体使用3类UTP, 它采用的信号速度为25MHz, 需要4对双绞线, 不使用曼彻斯特编码, 而是三元信号, 每个周期发送4bit, 这样就获得了所要求的100Mbps, 还



有一个 33.3Mbps 的保留信道。该方案即所谓的 8B6T (8bit 被映射为 6 个三进制位)。

2) 100Base-TX

100Base-TX 传输载体使用 5 类 100 Ω UTP, 其设计比较简单, 因为它可以处理速率高达 125MHz 以上的时钟信号, 每个站点只需使用两对双绞线, 一对连向集线器, 另一对从集线器引出。它采用了一种运行在 125MHz 下的称为 4B/5B 的编码方案, 该编码方案将每 4bit 的数据编成 5bit, 挑选时每组数据中不允许出现多于 3 个“0”, 然后再将 4B/5B 码进一步编成 NRZI 码进行传输。这样要获得 100Mbps 的数据传输速率, 只需要 125M 的信号速率。

3) 100Base-FX

100Base-FX 既可以选用多模光纤, 也可以选用单模光纤。在全双工情况下, 多模光纤传输距离可达 2km, 单模光纤传输距离可达 40km。

3. 千兆以太网

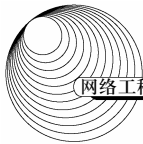
工作站之间用 100Mbps 以太网连接后, 对于主干网络的传输速度就会提出更高的要求, 1996 年 7 月, IEEE 802.3 工作组成立了 802.3z 千兆以太网任务组, 研究和制定了千兆以太网的标准, 这个标准满足以下要求: 允许在 1000Mbps 速度下进行全双工和半双工通信; 使用 802.3 以太网的帧格式; 使用 CSMA/CD 访问控制方法来处理冲突问题; 编址方式和 10Base-T、100Base-T 兼容。这些要求表明千兆以太网和以前的以太网完全兼容。1997 年 3 月, 又成立了另一个工作组 802.3ab 来集中解决用 5 类线构造千兆以太网的标准问题, 而 802.3z 任务组则集中制定使用光纤和对称屏蔽铜缆的千兆以太网标准。802.3z 标准于 1998 年 6 月由 IEEE 标准化委员会批准, 802.3ab 标准计划也于 1999 年通过批准。

1) 1000Base LX

1000Base LX 是一种使用长波激光作为信号源的网络介质技术, 在收发器上配置波长为 1270~1355nm (一般为 1300nm) 的激光传输器, 既可以驱动多模光纤, 也可以驱动单模光纤。1000Base LX 所使用的光纤规格包括 62.5 μ m 多模光纤、50 μ m 多模光纤、9 μ m 单模光纤。其中, 使用多模光纤时, 在全双工模式下, 最长传输距离可以达到 550m; 使用单模光纤时, 全双工模式下的最长有效距离为 5km。系统采用 8B/10B 编码方案, 连接光纤所使用的 SC 型光纤连接器与快速以太网 100Base FX 所使用的连接器的型号相同。

2) 1000Base SX

1000Base SX 是一种使用短波激光作为信号源的网络介质技术, 收发器上所配置的波长为 770~860nm (一般为 800nm) 的激光传输器不支持单模光纤, 只能驱动多模光纤, 具体包括 62.5 μ m 多模光纤、50 μ m 多模光纤。使用 62.5 μ m 多模光纤, 全双工模式下的最长传输距离为 275m; 使用 50 μ m 多模光纤, 全双工模式下最长有效距离为 550m。系统采用 8B/10B 编码方案, 1000Base SX 所使用的光纤连接器与 1000Base LX 一样, 也是 SC 型连接器。



3) 1000Base CX

1000Base CX 是使用铜缆作为网络介质的两种千兆以太网技术之一,另外一种就是将在后面介绍的 1000Base T。1000Base CX 使用的一种特殊规格的高质量平衡双绞线对的屏蔽铜缆,最长有效距离为 25m,使用 9 芯 D 型连接器连接电缆,系统采用 8B/10B 编码方案。1000Base CX 适用于交换机之间的短距离连接,尤其适合于千兆主干交换机和主服务器之间的短距离连接。以上连接往往可以在机房配线架上以跨线方式实现,不需要再使用长距离的铜缆或光纤。

4) 1000Base T

1000Base T 是一种使用 5 类 UTP 作为网络传输媒介的千兆以太网技术,最长有效距离与 100Base TX 一样,可以达到 100m。用户可以采用这种技术在原有的快速以太网系统中实现 100~1000Mbps 的平滑升级。与前文所介绍的其他三种网络介质不同,1000Base T 不支持 8B/10B 编码方案,需要采用专门的更加先进的编码/译码机制。

4. 万兆以太网

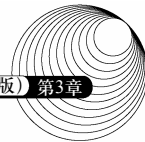
1) 10GE 以太网

2002 年 6 月,IEEE 802.3ae 10G 以太网标准发布,以太网的发展势头又得到了一次增强。确定万兆以太网标准的目的是将 802.3 协议扩展到 10Gbps 的工作速度,并扩展以太网的应用空间,使之能够包括 WAN 链接。万兆以太网与 SONET OC-192 帧结构的融合,可以与 OC-192 电路和 SONET/SDH 设备一起运行,保护了传统基础设施投资,使供应商能够在不同地区通过城域网提供端到端以太网。

物理层: 802.3ae 大体分为两种类型,一种是与传统以太网连接,速率为 10Gbps 的“LAN PHY”;另一种是连接 SDH/SONET,速率为 9.58464Gbps 的“WAN PHY”。每种 PHY 分别可使用 10GBase-S (850nm 短波)、10GBase-L (1310nm 长波)、10GBase-E (1550nm 长波) 三种规格,最大传输距离分别为 300m、10km、40km,其中 LAN PHY 还包括一种可以使用 DWDM 波分复用技术的“10GBASE-LX4”规格。WAN PHY 与 SONET OC-192 帧结构融合,可与 OC-192 电路、SONET/SDH 设备一起运行,保护传统基础投资,使运营商能够在不同地区通过城域网提供端到端以太网。

传输介质层: 802.3ae 目前支持 9 μ m 单模、50 μ m 多模和 62.5 μ m 多模三种光纤,而对电接口的支持规范 10GBASE-CX4 目前正在讨论之中,尚未形成标准。

数据链路层: 802.3ae 继承了 802.3 以太网的帧格式和最大/最小帧长度,支持多层星型连接、点到点连接及其组合,充分兼容已有应用,不影响上层应用,进而降低了升级风险。与传统的以太网不同,802.3ae 仅仅支持全双工方式,而不支持单工和半双工方式,不采用 CSMA/CD 机制。802.3ae 不支持自协商,可简化故障定位,并提供广域网物理层接口。



人们不仅在万兆以太网的技术和性能方面看到了其实质性的提高,也正因如此,以太网正在从局域网逐步延伸至城域网和广域网,在更广阔的范围内发挥其作用。

2) 40GE 以太网

2003 年 5 月 26 日,在以太网技术行将迎来 30 岁诞辰之际,思科高级副总裁 Luca Cafiero 指出,未来两年内,以太网的最高数据传输速率将可望提高至 40Gbps。他称,业内将 40G 而非 100G 确定为以太网下一步发展目标的重要原因在于,与 100G 以太网相比,研发 40G 以太网在技术上面面临的挑战相对较小,更为切实可行。与此同时,Cafiero 还指出,实际上,借助新发布的 Supervisor Engine 720 引擎,思科公司的 Catalyst6500 旗舰级企业交换平台目前已可以为每一接口卡提供 40Gbps 的数据传输速率支持。他还指出,新型以太网技术成功的关键在于能够推动单位数据传输成本的下降。也就是说,新的以太网技术的 1bps 数据传输成本必须低于原有技术才能大获成功。

3.2.3 以太网技术基础

1. IEEE 802.3 帧的结构

媒介访问控制子层(MAC)的功能是以太网的核心技术,它决定了以太网的主要网络性能。MAC 子层通常又分成帧的封装/解封和媒介访问控制两个功能模块。在了解该子层的功能时,首先要了解以太网的帧结构,其帧结构如图 3-7 所示。

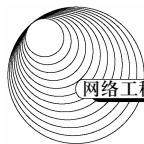
7	1	6	6	2	46~1500	4
前导码	帧首定界符 (SFD)	目的地址 (DA)	源地址 (SA)	长度 (L)	逻辑链路层 协议数据单元 (LLC-PDU)	帧检验序列 (FCS)

图 3-7 IEEE 802.3 帧的结构

(1) 前导码: 包含了 7 个字节的二进制“1”、“0”间隔的代码,即 1010...10 共 56 位。当帧在媒介上传输时,接收方就能建立起位同步,因为在使用曼彻斯特编码的情况下,这种“1”、“0”间隔的传输波形为一周期性方波。

(2) 帧首定界符(SFD): 它是长度为 1 个字节的 10101011 二进制序列,此码一过,表示一帧实际开始,以使接收器对实际帧的第一位定位。也就是说,实际帧由余下的 DA+SA+L+LLCPDU+FCS 组成。

(3) 目的地址(DA): 它说明了帧企图发往的目的站地址,共 6 个字节,可以是单址(代表单个站)、多址(代表一组站)或全地址(代表局域网上的所有站)。当目的地址出现多址



时,即表示该帧被一组站同时接收,称为“组播(multicast)”。当目的地址出现全地址时,即表示该帧被局域网上的所有站同时接收,称为“广播(broadcast)”。通常以 DA 的最高位来判断地址的类型,若最高位为“0”,则表示单址;为“1”表示多址或全地址。全地址时,DA 字段为全“1”代码。

(4) 源地址(SA):它说明发送该帧的站的地址,与 DA 一样占 6 个字节。

(5) 长度(L):共占两个字节,表示 LLC-PDU 的字节数。

(6) 数据链路层协议数据单元(LLC-PDU):它的范围处在 46~1500 字节之间。注意,46 字节最小 LLC-PDU 长度是一个限制,目的是要求局域网上的所有站都能检测到该帧,即保证网络正常工作。如果 LLC-PDU 小于 46 个字节,则发送站的 MAC 子层会自动填充“0”代码补齐。

(7) 帧检验序列(FCS):它处在帧尾,共占 4 字节,是 32 位冗余检验码(CRC),检验除前导码、SFD 和 FCS 以外的所有帧的内容,即从 DA 开始至 DATA 完毕的 CRC 检验结果都反映在 FCS 中。当发送站发出帧时,一边发送,一边逐位进行 CRC 检验。最后形成一个 32 位 CRC 检验和填在帧尾 FCS 位置中一起在媒介上传输。接收站接收帧后,从 DA 开始同样边接收边逐位进行 CRC 检验。最后接收站形成的检验和若与帧的检验和相同,则表示媒介上传输的帧未被破坏。反之,接收站认为帧被破坏,会通过一定的机制要求发送站重发该帧。

一个帧的长度为 $DA+SA+L+LLCPDU+FCS=6+6+2+(46\sim1500)+4=64\sim1518$ 即,当 LLC-PDU 为 46 字节时,帧最小,帧长为 64 字节;当 LLC-PDU 为 1500 字节时,帧最大,帧长为 1518 字节。

2. 以太网的跨距

系统的跨距表示了系统中任意两个站点间的最大距离范围,媒介访问控制方式 CSMA/CD 约束了整个共享型快速以太网系统的跨距。

前文介绍了 CSMA/CD 的重要的参数碰撞槽时间(Slot time),可以认为:

$$\text{Slot time} \approx 2S/0.7C + 2t_{\text{PHY}}$$

如果考虑一段媒介上配置了中继器,且中继器的数量为 N ,设一个中继器的延时为 t_r ,则

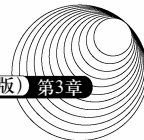
$$\text{Slot time} \approx 2S/0.7C + 2t_{\text{PHY}} + 2Nt_r$$

由于 $\text{Slot time} = L_{\min}/R$, $L_{\min}$ 称为最小帧长度, R 为传输速率,则系统跨距 S 的表达式为:

$$S \approx 0.35C(L_{\min}/R - 2t_{\text{PHY}} - 2Nt_r)$$

通过前面的学习可知, $L_{\min}=64\text{B}=512\text{b}$, $C=3\times10^8\text{m/s}$,所以在 10M 以太网环境中, $R=10\times10^6\text{bps}$;在 100M 以太网环境中, $R=100\times10^6\text{bps}$ 。

如果忽略 $2t_{\text{PHY}}$ 和 $2Nt_r$, 10M 以太网环境中最大跨距为 5376m, 100M 以太网环境中最大跨距为 537.6m。如果在实际应用中忽略中继器,只算上 $2t_{\text{PHY}}$,则 10M 以太网环境中最大跨距



为 5000m 左右, 100M 以太网环境中最大跨距约为 412m。然而, 在实际应用中, 物理层所耗去的时间和中继器所耗去的时间都是不能忽略的, 这也就是有 5-4-3 法则的原因。尤其是在 R 变大时, 跨距成几何级数递减, 当 R 为 1000M 时, 依据这个法则, 根据物理层所耗去的时间的大小, 甚至会出现跨距为负的情况, 则网络变得不可用。为此, 1Gbps 以太网上采用了帧的扩展技术, 目的是为了在半双工模式下扩展碰撞域, 达到增长跨距的目的。

帧扩展技术是在不改变 802.3 标准所规定的最小帧长度的情况下提出的一种解决办法, 把最小帧长一直扩展到 512 字节 (即 4096 位)。若形成的帧小于 512 字节, 则在发送时要在帧的后面添上扩展位, 达到 512 字节再发送到媒介上去。扩展位是一种非“0”、“1”数值的符号, 若形成的帧已大于或等于 512 字节, 则发送时不必添加扩展位。这种解决办法使得在媒介上传输的帧长度最短不会小于 512 字节, 在半双工模式下大大扩展了碰撞域, 媒介的跨距可延伸至 330m。在全双工模式下, 由于不受 CSMA/CD 约束, 无碰撞域概念, 因此在媒介上的帧无必要扩展到 512 字节。

100Base TX/FX 系统的跨距如图 3-8 所示。由于跨距实际上反映了一个碰撞域, 因此图中用两个 DTE 之间的距离来表示, DTE 可以是一个网桥、交换器或路由器, 也可以认为是系统中的两个站点。中继器用 R 表示一般是一个共享型集线器, 它的功能是延伸媒介和连接另一个媒介段。

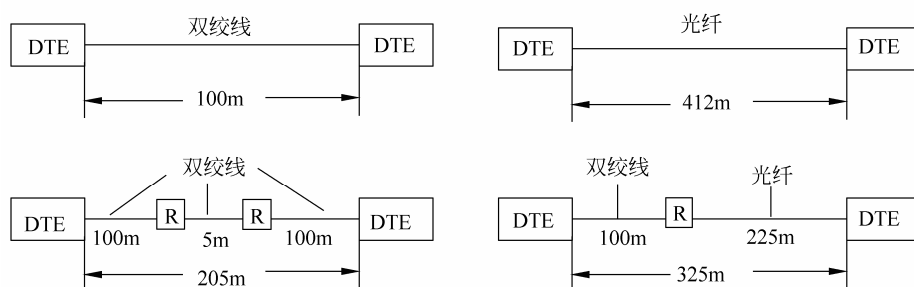
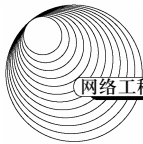


图 3-8 100M 以太网的跨距

在双绞线媒介情况下, 由于最长媒介段距离为 100m, 加一个中继器, 就延伸一个最长媒介段距离, 达到 200m。如果想再延伸距离, 加两个中继器后, 也只能达到 205m, 205m 即为 100Base TX 的跨距。

在光纤媒介情况下, 不使用中继器, 跨距可达到 412m, 即是一个碰撞域范围, 但光纤的最长媒介段 2km 要远远大于 412m。另外, 加一个中继器后, 并不能延伸距离, 由于中继器的延迟时间, 跨距反而变小了, 在加两个中继器时, 跨距几乎和双绞线加两个中继器的跨距相同。因此, 在实际应用中通常采用混合方式, 即中继器一侧采用光纤, 另一侧采用双绞线。双绞线



可直接连接用户终端,跨距可达 100m,光纤可直接连接路由器或主干全双工以太网交换机,跨距可达 225m。

3. 交换型以太网

在交换型以太网出现以前,以太网系统均为共享型以太网系统。在整个系统中,由于受到 CSMA/CD 媒介访问控制方式的制约,所以整个系统处一在一个碰撞域范围中,系统中每个站都可能在往媒介上发送帧,那么每个站要占用媒介的机率就是 $10\text{Mbps}/n$, n 为站数。以太网受到 CSMA/CD 制约后,所有站均在争用媒介而共同分割带宽,称“共享型”以太网。

在 20 世纪 80 年代后期,即 10Base T 出现后不久,就出现了以太网交换型集线器。到了 20 世纪 90 年代,快速以太网的交换技术和产品更是发展迅速,应用广泛。交换型以太网系统中的交换型集线器也称为以太网交换机,以其为核心连接站点或者网段。如图 3-9 所示,交换器的各端口之间在交换器上同时可以形成多个数据通道,图中在交换器上同时存在 4 个数据通道,它们可以存在于站与站、站与网段或者网段与网段之间。网段即是多个站点构成的一个共享媒介的集合,一般是一个共享型集线器连接若干个站点构成一个网段。

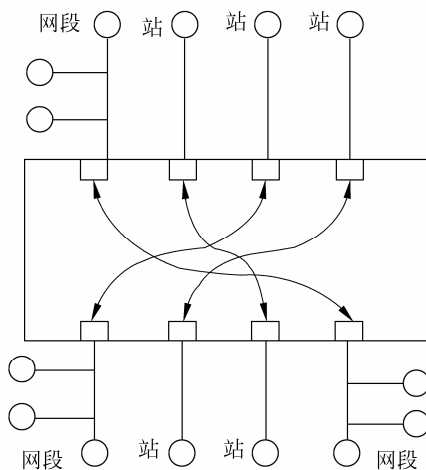
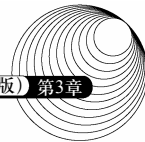


图 3-9 以太网交换机示意图

既然是在交换器上同时存在多个端口间的通道,也就意味着系统同时存在多个碰撞域,每一个碰撞域的一对端口都独占带宽(一个享有发送带宽,另一个享有接收带宽),那么就整个系统的带宽来说,就不再是只有 10Mbps (10Base T 环境)或 100Mbps (100Base T 环境),而是与交换器所具有的端口数有关。可以认为,若每个端口为 10Mbps,则整个系统带宽可达 $10M \cdot n$, 其中 n 为端口数,若 $n=10$,则系统带宽可达 100Mbps。因此,拓宽了整个系统带宽



是交换型以太网系统最明显的特点。

综上所述, 交换型以太网系统与共享型以太网比较有如下优点。

(1) 每个端口上可以连接站点, 也可以连接一个网段。不论站点或网段均独占该端口的带宽 (10Mbps 或 100Mbps)。

(2) 系统的最大带宽可以达到端口带宽的 n 倍, 其中 n 为端口数。 n 越大, 系统的带宽越高。

(3) 交换器连接了多个网段, 每一个网段都是独立、被隔离的。但如果需要, 独立网段之间通过其端口也可以建立暂时的数据通道。

(4) 被交换器隔离的独立网段上数据流信息不会随意广播到其他端口上去, 因此具有一定的数据安全性。

4. 全双工以太网

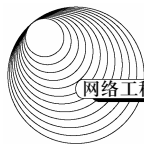
交换器设备工作时不同的逻辑数据通道之间已不再受到 CSMA/CD 的约束, 但每条逻辑数据通道的两个端口之间却仍然受到 CSMA/CD 的约束, 即一条逻辑数据通道就是一个碰撞域。

当交换器以太网技术和应用发展到一定阶段后, 不仅要求整个系统的带宽要达到一定高度, 而且还要求整个系统的跨距也要有一定的保证, 特别在 100Mbps 及 1Gbps 以太网环境中, 使用光纤作为媒介的情况下, 若再使用受到 CSMA/CD 约束的一般半双工技术和产品, 网络覆盖范围的矛盾会尤为突出。为了解决上述问题, 全双工以太网技术和产品问世了, 且在 1997 年由 IEEE 802.3x 标准来说明该技术的规范。

全双工以太网技术是用来说明以太网设备端口的传输技术, 与传统半双工以太网技术的区别在于每个端口和交换机背板之间都存在两条逻辑通路。这样, 每一个端口就可以同时接收和发送帧, 不再受到 CSMA/CD 的约束, 在端口发送帧时不再会发生帧的碰撞, 已无碰撞域的存在。这样一来, 端口之间媒介的长度仅仅受到数字信号在媒介上的传输衰变的影响, 而不像传统以太网半双工传输时还要受到碰撞域的约束。

图 3-10 所示为两个端口之间全双工传输, 端口上设有端口控制功能模块和收发器功能模块, 端口上是全双工还是半双工操作一般可以自适应, 也可以用人工设置。当全双工操作时, 帧的发送和接收可以同时进行, 这样与传统半双工操作方式比较, 传输链路的带宽提高了一倍, 即端口支持 10Mbps 或者 100Mbps 传输率, 而其带宽却分别是 20Mbps 和 200Mbps。在全双工传输帧时, 端口上既无侦听的机制, 链路上又不会多路访问, 也不再需要碰撞检测, 传统半双工方式下的媒介访问控制 CSMA/CD 的约束已不存在。

在 10Mbps 端口传输率情况下, 只有 10BaseT 及 10Base FL 支持全双工操作, 而在 100Mbps 快速以太网情况下, 除了 100Base T4 外, 100BaseTX 和 100Base TX 均支持全双工操作。千兆位以太网 1000Base X 也支持全双工操作。即只有链路上提供独立的发送和接收媒介才能支持全



双工操作。表 3-1 说明了支持全双工操作的各类以太网网段的最长距离, 并与传统半双工操作受碰撞域限定的网段最长距离进行比较。

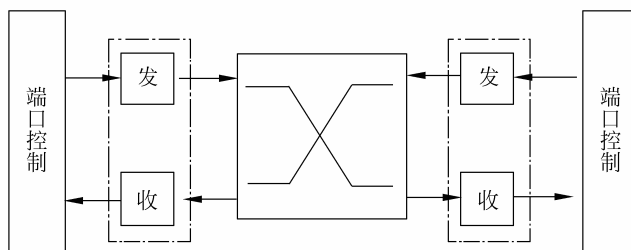
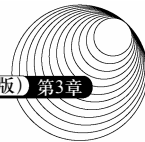


图 3-10 全双工以太网交换机示意图

表 3-1 以太网网段的最长距离

以太网类型	传 输 媒 介	全双工网段最长距离	半双工网段最长距离
10Base T	UTP	100m	100m
10Base F	MMF	2km	2km
100Base T	UTP、STP	100m	100m
100Base F	MMF	2km	412m
1000Base LX	MMF	550m	330m
	SMF	5km	330m
1000Base SX	MMF62.5μm	300m	
	MMF50μm	550m	330m
1000Base CX	STP	25m	25m
1000Base T	UTP	100m	100m

从表 3-1 可知, 使用双绞线媒介, 100m 的距离对于半双工操作来说并非是碰撞域的跨距, 仍是数字信号驱动的最长距离, 因此不论是 10Mbps、100Mbps 还是 1000Mbps 环境, 全双工操作并未占有优势。对于媒介采用光纤来说, 10Base FL 在两种情况下光纤最长距离均为 2km, 这是因为在 10Mbps 传输速率情况下, 由碰撞域决定的半双工网段最长距离要大于 2km, 2km 的光纤仍是由数字信号在光纤上传输的最长距离。在 100Base FX 的以太网中, 全双工网段距离可达 2km, 而传统的半双工操作情况下, 由于受到 CSMA/CD 的约束, 碰撞域的跨距决定了网段最长距离为 412m。在 1000Base LX 的以太网中, 采用单模光纤全双工网段距离可达 5km, 而传统的半双工操作情况下, 由于受到 CSMA/CD 的约束, 碰撞域的跨距决定了网段最长距离为 330m。在 1000Base SX 的以太网中, 因不能使用单模光纤, 多模光纤扩展网络有效距离的效果并不明显。



对于网络中的客户机而言, 由于其访问服务器时, 发送和接收的负载往往是很不均衡的, 因此, 用全双工操作方式连接客户站, 可延伸距离, 有明显的得益。对于服务器, 由于会受到许多客户站的同时访问, 所以发送和接收的负载一般较接近均衡, 所以使用全双工操作方式增加带宽是明显的得益, 但由于系统服务器往往与系统主交换机放置在一起, 因此延伸距离上显得无必要。对于交换机之间的连接来说, 使用全双工操作方式, 在延伸连接距离和拓展带宽上均能得益。

3.2.4 以太网交换机的部署

在应用级的局域网中, 很少存在只使用单台交换机的局域网, 一方面因为单台交换机的端口数量有限, 另一方面, 单台交换机的地理位置使得联网计算机终端的距离受限。通常, 在一个局域网中, 使用几台交换互相连接在一起, 从而达到扩展端口和扩展距离的目的。那么交换机与交换机是如何连接在一起的呢? 目前广泛使用的模式有两种, 一种是级连(uplink)模式, 另一种是堆叠(stack)模式。

1. 级连模式

级连模式是最常规、最直接的一种扩展方式。级连模式通过双绞线或光纤实现, 一般在交换机的前面板上有专门的级连口, 如果没有, 也可以用交叉接法来级连。级连是通过端口进行的, 级连后两台交换机是上下级的关系。

级连模式起源于早期的共享型集线器(Hub), 共享型集线器的物理拓扑结构是星型, 而逻辑拓扑结构还是总线型的, 集线器仅仅相当于一条浓缩了的总线, 在集线器的某一个端口级连另一台集线器, 只是相当于把浓缩的总线又加长了一些, 仍然是一条总线, 所有端口都要在一个碰撞域里受到 CSMA/CD 的约束。但这样相当于把传输媒介加长了, 在加长的传输媒介上又增加了一些端口。但付出的代价是, 在这个碰撞域里又多了一些端口共享整个带宽, 从而导致网络性能低下。当然这种级连方式必须遵循 5-4-3 法则, 也就是级连不能超过 4 层。级连模式的典型结构如图 3-11 所示。

需要特别指出的是, 对于那些没有专用级连端口的集线器之间的级连, 双绞线接头中线对的分布与连接网卡和集线器时有所不同, 必须要用交叉线。而许多集线器为了方便用户, 提供了一个专门用来串接到另一台集线器的端口, 在对此类集线器进行级连时, 双绞线均应为直通线接法。不管采用交叉线还是直通线进行级连, 都没有改变级连的本质。

用户如何判断自己的集线器是否需要交叉线连接呢? 主要方法有以下几种。

(1) 查看说明书。如果该集线器在级连时需要交叉线连接, 一般会在设备说明书中进行说明。

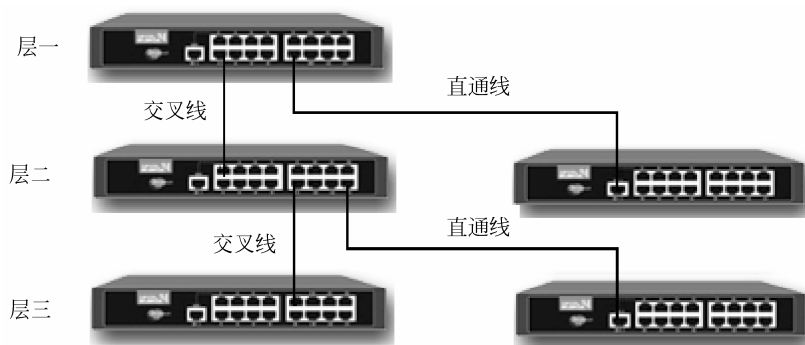
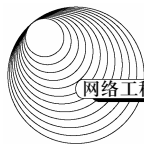


图 3-11 级连模式的以太网交换机

(2) 查看连接端口。如果该集线器在级连时不需要交叉线,大多数情况下都会提供一至两个专用的互连端口,并有相应标注,如“Uplink”、“MDI”、“Out to Hub”,表示使用直通线连接。

(3) 实测。这是最管用的一种方法,可以先制作两条用于测试的双绞线,其中一条是直通线,另一条是交叉线,之后用其中的一条连接两个集线器,这时注意观察连接端口对应的指示灯,如果指示灯亮,表示连接正常,否则换另一条双绞线进行测试。

随着快速以太网技术和交换技术的出现,级连模式又逐渐变成组建大型局域网最理想的扩展方式,成为以太网扩展端口应用中的主流技术。在交换机上进行级连,级连交换机的端口共享的仅仅是被级连交换机中被级连端口的带宽,而不是整个网络的带宽。更何况目前的交换机级连通常都是高速交换机端口级连低速交换机,即 1000M 端口级连 100M 的交换机,100M 端口则级连 10M 的交换机,或者是交换机级连共享型的集线器。由此一来,极大程度地克服了传统集线器级连共享带宽而导致网络性能降低的弊端。虽然交换机的级连在一定程度上仍然受 CSMA/CD 的约束,但其优势却是不可替代的,通常表现在以下几个方面。

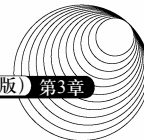
(1) 级连模式可使用通用的以太网端口进行层次间互联,其中包括 100M 端口、1000M 端口以及新兴的 10G 端口。

(2) 级连模式是组建结构化网络的必然选择,级连使用普通的、长度限制并不严格的电缆(光纤),各个级连单元的位置相对较随意,非常有利于综合布线。

(3) 级连模式通常是解决不同品牌交换机之间以及交换机与集线器之间连接的有效手段。

2. 堆叠模式

堆叠模式通常是为了扩充带宽用的,通常用专门的堆叠卡插在交换机的后面,用专门的堆



叠电缆连接几台交换机,堆叠后这几台交换机相当于一台交换机。堆叠是采用交换机背板的叠加,使多个工作组交换机形成一个工作组堆,从而提供高密度的交换机端口,堆叠中的交换机就像一个交换机一样,配制一个 IP 地址即可。

级连是通过交换机的某个端口与其他交换机相连的,而堆叠是通过集线器的背板连接起来的,它是一种建立在芯片级上的连接,如 2 个 24 口交换机堆叠起来的效果就像是一个 48 口的交换机。

堆叠模式的优点如下。

(1) 增加网络端口的同时,还增加了逻辑数据通道,扩充了网络带宽,不同堆叠单元的端口之间可以直接交换,进行快速转发,从而极大地提高了网络性能。

(2) 不受 5-4-3 原则的约束,堆叠单元可以超过 4 个。

(3) 提供简化的本地管理,将一组交换机作为一个对象来管理。

堆叠模式的缺点如下。

(1) 堆叠是一种非标准化技术,各个厂商之间不支持混合堆叠,同一组堆叠交换机必须是同一品牌。

(2) 堆叠模式不支持即插即用,在物理连接完毕之后,还要对交换机进行相应的设置,才能正常运行。

(3) 不存在拓扑管理,一般不能进行分布式布置。

常见的堆叠有菊花链堆叠和矩阵堆叠两种。

所谓菊花链,就是从上到下串起来,形成单一的一个菊花链堆叠总线。菊花链模式是简化的级联模式,主要的优点是提供集中管理的扩展端口,对于多交换机之间的转发效率并没有提升,主要是因为菊花链模式是采用高速端口和软件来实现的。菊花链模式使用堆叠电缆将几台交换机以环路的方式组建成一个堆叠组,然后加一根从上到下起冗余备份作用的堆叠电缆。图 3-12 所示是 2003 年 6 月北电网络推出的 BayStack 5510 菊花链堆叠交换机,一个堆叠中有 8 个交换机,整个堆叠的带宽高达 640Gbps,每台交换机与上下相邻单元间都具有 40Gbps 的全双工带宽。

矩阵堆叠需要提供一个独立的或者集成的高速交换中心(堆叠中心),所有堆叠的交换机通过专用的高速堆叠端口上行到统一的堆叠中心,堆叠中心一般是一个基于专用 ASIC 的硬件交换单元,ASIC 交换容量限制了堆叠的层数。使用高可靠、高性能的 Matrix 芯片是星型堆叠的关键。由于涉及专用总线技术,电缆长度一般不能超过 2m,所以,矩阵堆叠模式下,所有交换机需要局限在一个机架之内。图 3-13 所示是 3COM 公司的 3300 系统交换机连成矩阵堆叠的示意图。

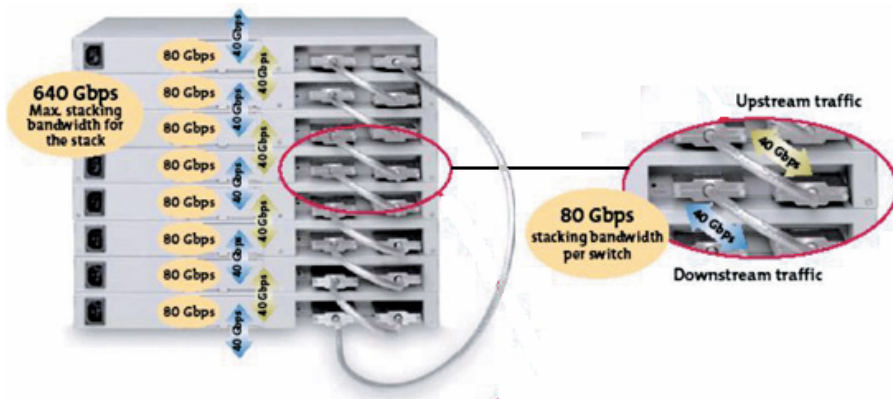
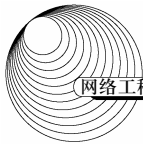


图 3-12 菊花链堆叠交换机

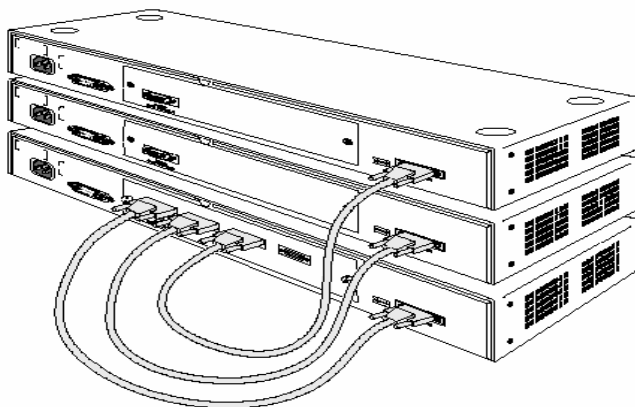


图 3-13 矩阵堆叠交换机

3. 混合模式

通过前文的介绍不难得出结论,堆叠模式是一种集中管理的端口扩展技术,不能提供拓扑管理,没有国际标准,且兼容性较差。但是,对于那些对带宽要求较高并需要大量端口的单节点局域网,堆叠模式可以提供比较优秀的转发性能和方便的管理特性。级连模式是组建网络的基础,可以灵活利用各种拓扑、冗余技术,对于那些对带宽要求不高且级连层次很少的网络,级连方式可以提供最优化的性能。可见,级连模式和堆叠模式的优点和缺点都十分鲜明,单纯地运用任何一种模式,都不会最大限度地优化网络。在实际的应用中,由于网络的复杂性。用户需求的多重性,通常同时使用两种模式进行交换机的部署,称其为混合模式。图 3-14 所示就

是考虑了半双工以太网最大跨距约束的一个典型的混合模式交换机部署方案。

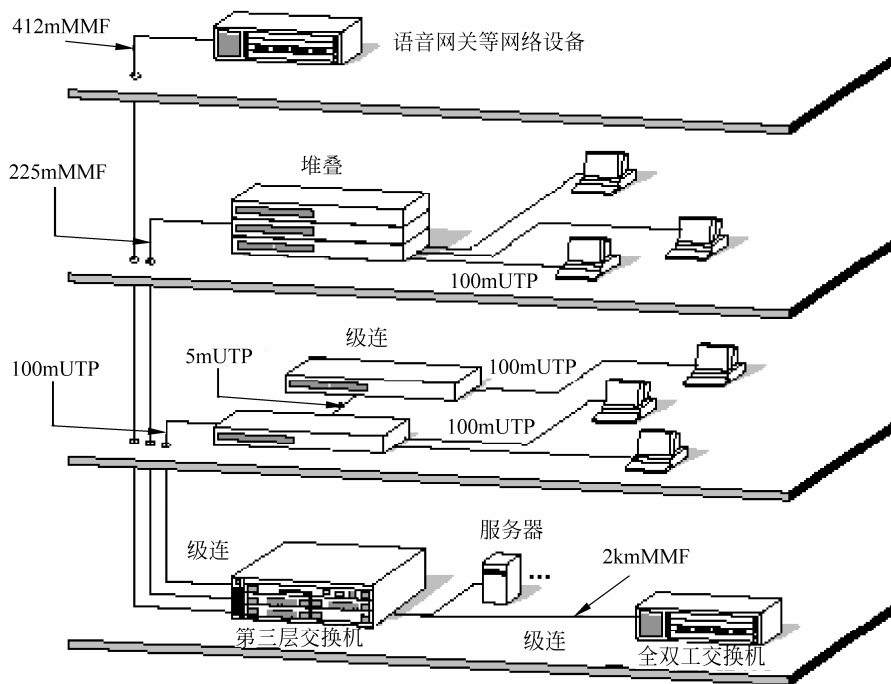


图 3-14 混合模式部署交换机

3.3 交换机与路由器的基本配置

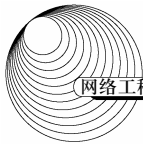
3.3.1 交换机的基本配置

不同厂家生产的不同型号的交换机，其具体的配置命令和方法是有差别的。不过配置的原理基本都是相同的，下面主要以 Cisco Crystal 2950 系列交换机为例讲解交换机配置的基本技术和技能。

1. 电缆连接及终端配置

如图 3-15 所示，接好 PC 和交换机各自的电源线，在未开机的条件下，把 PC 的串口 1 (COM1) 通过控制台电缆与交换机的 Console 端口相连，即可完成设备的连接工作。

交换机 Console 端口的默认参数如下。



- 端口速率: 9600bps。
- 数据位: 8。
- 奇偶校验: 无。
- 停止位: 1。
- 流控: 无。

在配置 PC 的超级终端时只需将端口属性的配置和上述参数相匹配, 就可以成功地访问到交换机。如图 3-16 所示为以 Windows 2000 中文专业版为例配置 COM 端口属性窗口。

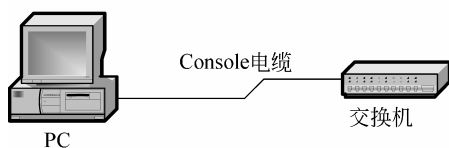


图 3-15 仿真终端与交换机的连接

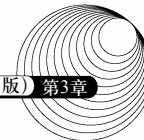


图 3-16 仿真终端端口参数配置

2. 交换机的启动

在连接好线路, 配置好超级终端仿真软件后, 就可以打开交换机, 此时超级终端窗口就会显示交换机的启动信息, 清单如下。

```
C2950 Boot Loader (C2950-HBOOT-M) Version 12.1(11r)EA1, RELEASE SOFTWARE (fc1)
Compiled Mon 22-Jul-02 18:57 by miwang
Cisco WS-C2950T-24 (RC32300) processor (revision C0) with 21039K bytes of memory.
2950T-24 starting...
Base ethernet MAC Address: 00E0.8FB5.032C
Xmodem file system is available.
Initializing Flash...
flashfs[0]: 1 files, 0 directories
flashfs[0]: 0 orphaned files, 0 orphaned directories
flashfs[0]: Total bytes: 64016384
flashfs[0]: Bytes used: 3058048
flashfs[0]: Bytes available: 60958336
```



flashfs[0]: flashfs fsck took 1 seconds.
...done Initializing Flash.

Boot Sector Filesystem (bs:) installed, fsid: 3
Parameter Block Filesystem (pb:) installed, fsid: 4

Loading "flash:/c2950-i6q4l2-mz.121-22.EA4.bin"...

[OK]

Restricted Rights Legend

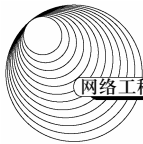
Use, duplication, or disclosure by the Government is
subject to restrictions as set forth in subparagraph
(c) of the Commercial Computer Software - Restricted
Rights clause at FAR sec. 52.227-19 and subparagraph
(c) (1) (ii) of the Rights in Technical Data and Computer
Software clause at DFARS sec. 252.227-7013.

cisco Systems, Inc.
170 West Tasman Drive
San Jose, California 95134-1706

Cisco Internetwork Operating System Software
IOS (tm) C2950 Software (C2950-I6Q4L2-M), Version 12.1(22)EA4, RELEASE SOFTWARE(fc1)
Copyright (c) 1986-2005 by cisco Systems, Inc.
Compiled Wed 18-May-05 22:31 by jharirba

Cisco WS-C2950T-24 (RC32300) processor (revision C0) with 21039K bytes of memory.
Processor board ID FHK0610Z0WC
Running Standard Image
24 FastEthernet/IEEE 802.3 interface(s)
2 Gigabit Ethernet/IEEE 802.3 interface(s)

32K bytes of flash-simulated non-volatile configuration memory.
Base ethernet MAC Address: 00E0.8FB5.032C
Motherboard assembly number: 73-5781-09
Power supply part number: 34-0965-01



Motherboard serial number: FOC061004SZ
Power supply serial number: DAB0609127D
Model revision number: C0
Motherboard revision number: A0
Model number: WS-C2950T-24
System serial number: FHK0610Z0WC

Cisco Internetwork Operating System Software
IOS (tm) C2950 Software (C2950-I6Q4L2-M), Version 12.1(22)EA4, RELEASE SOFTWARE(fc1)
Copyright (c) 1986-2005 by cisco Systems, Inc.
Compiled Wed 18-May-05 22:31 by jharirba

Press RETURN to get started!

以上启动过程为用户提供了丰富的信息,利用这些信息,可以对交换机的硬件结构和软件加载过程有直观的认识。这些信息对于了解该路由器以及对它做相应的配置很有帮助。另外,部件号、序列号、版本号等信息在产品验货时都是非常重要的信息。

3. 交换机的配置模式

在进行交换机的配置之前,需要了解交换机的基本配置模式。常见的交换机配置模式有普通用户模式、特权模式、全局配置模式和局部配置模式。在这些配置模式下,用户对交换机所具有的权限是不同的。在普通用户模式下,用户只能够对交换机进行简单的操作,如查询操作系统版本和系统时间,使用很少的几个命令;在特权模式下,用户可以使用较多的命令对交换机进行查看、配置等操作;在全局配置模式下,主要完成对交换机的配置,如虚拟局域网的配置、访问控制列表的配置等;在局部配置模式下,用户可以对某个具体端口进行配置,这几种配置模式是递进的关系。

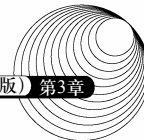
(1) 用户模式。在交换机正常启动后,用户使用超级终端仿真软件或 Telnet 登录交换机,可自动进入用户配置模式,其命令状态如下。

```
switch>
```

(2) 特权模式。在用户模式下,输入以下命令可以进入特权模式。

```
switch>enable  
switch#
```

(3) 全局配置模式。在特权模式下,输入以下命令可以进入全局配置模式。



```
switch>config terminal
switch(config)#
```

(4) 局部配置模式。局部配置模式包括端口配置模式和线路配置模式，在全局配置模式下，输入以下命令可以进入局部配置模式。

```
Switch(config)#interface fastEthernet 0/1
Switch(config-if)#    (端口配置模式)
switch(config)#line console 0
switch(config-line)#  (线路配置模式)
```

4. 交换机的基本配置

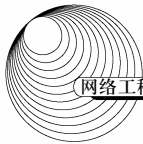
在默认配置下，所有接口处于可用状态，并且都属于 VLAN1，这种情况下交换机就可以正常工作了。但为了方便管理和使用，首先应对交换机做基本的配置。最基本的配置可以通过启动时的对话框配置模式完成，也可以在交换机启动后再进行配置。

(1) 配置 enable 口令和主机名。在交换机中可以配置使能口令 (enable password) 和使能密码 (enable secret)，一般情况下只需配置一个，当两者同时配置时，后者生效。这两者的区别是使能口令以明文显示，而使能密码以密文形式显示。

Switch>	(用户执行模式提示符)
Switch >enable	(进入特权模式)
Switch #	(特权模式提示符)
Switch #config terminal	(进入全局配置模式)
Switch(config)#	(全局配置模式提示符)
Switch(config)#enable password cisco	(设置 enable password 为 cisco)
Switch(config)#enable secret cisco1	(设置 enable secret 为 cisco1)
Switch(config)#hostname C2950	(设置主机名为 C2950)
C2950(config)#end	(退回到特权模式)
C2950#	

(2) 配置交换机 IP 地址、默认网关、域名、域名服务器。应该注意的是，这里所设置的 IP 地址、网关、域名等信息，是为交换机本身用来管理交换机所设置的，和连接在该交换机上的计算机或其他网络设备无关。也就是说，所有与交换机连接的主机都应该设置自身的域名、网关等信息。

C2950(config)#ip address 192.168.1.1 255.255.255.0	(设置交换机 IP 地址)
C2950(config)#ip default-gateway 192.168.1.254	(设置默认网关)
C2950(config)#ip domain-name cisco.com	(设置域名)



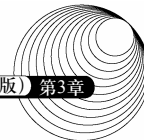
```
C2950(config)#ip name-server 200.0.0.1      (设置域名服务器)
C2950(config)#end
```

(3) 配置交换机的端口属性。交换机的端口属性默认支持一般网络环境下的正常工作，一般情况下是不需要对其端口进行设置的。在某些情况下需要对其端口属性进行配置时，配置的对象主要有速率、双工和端口描述等信息。

```
C2950(config)#interface FastEthernet0/1      (进入接口 0/1 的配置模式)
C2950(config-if)#speed ?                     (查看 speed 命令的子命令)
  10 Force 10 Mbps operation                  (显示结果)
  100 Force 100 Mbps operation
  auto Enable AUTO speed configuration
C2950(config-if)#speed 100                   (设置该端口速率为 100Mbps)
C2950(config-if)#duplex ?                    (查看 duplex 命令的子命令)
  auto Enable AUTO duplex configuration
  full Force full duplex operation
  half Force half-duplex operation
C2950(config-if)#duplex full                  (设置该端口为全双工)
C2950(config-if)#description TO_PC1          (设置该端口描述为 TO_PC1)
C2950(config-if)#^Z                          (返回特权模式，同 end)
C2950#show interface FastEthernet0/1         (查看端口 0/1 的配置结果，结果略)
C2950#show interface FastEthernet0/1 status  (查看端口 0/1 的状态，结果略)
```

(4) 配置和查看 MAC 地址表。有关 MAC 地址表的配置有三个方面的，即超时时间、永久地址和限制性地址。交换机学习到的动态 MAC 地址的超时时间默认为 300s，可以通过命令来修改这个值。设置了静态 MAC 地址，这个地址永久存在于 MAC 地址表中，不会超时，所有端口均可以转发以太网帧给该端口。限制性静态 (restricted static) 地址是在永久地址的基础上，同时限制了源端口，其安全性更高。

```
C2950(config)#mac-address-table ?            (查看 mac-address-table 的子命令)
  aging-time Aging time of dynamic addresses
  permanent Configure a permanent address
  restricted Configure a restricted static address
C2950(config)#mac-address-table aging-time 100 (设置超时时间为 100s)
C2950(config)#mac-address-table permanent 0000.0c01.bbcc f0/3 (加入永久地址)
C2950(config)#mac-address-table restricted static 0000.0c02.bbcc f0/6 f0/7 (加入静态地址)
C2950(config)#end
C2950#show mac-address-table                 (查看整个 MAC 地址表)
```



Number of permanent address: 1

Number of restricted static address: 1

Number of dynamic addresses: 0

Address	Dest Interface	Type	Source Interface List
0000.0C01.BBCC	FastEthernet 0/3	Permanent	All
0000.0C02.BBCC	FastEthernet 0/6	Static	F0/7

可以看到永久地址有 1 个, 设置在 F0/3 端口上; 限制性地址有 1 个, 设置目标端口为 F0/6, 源端口为 F0/7。如果交换机上连接有其他计算机, 则每个连接的端口都会产生动态 MAC 地址表项。用 `clear` 命令可以清除 MAC 地址表的某项设置, 比如如下命令可以清除限制性地址。

```
C2950#clear mac-address-table restricted static
```

在使用配置命令时可以使用缩写形式, 缩写的程度以不引起命令混淆为前提, 比如 `Router>enable`, 在 `Router>` 模式下以 `en` 开头的命令只有 `enable` 一个, 所以该命令可以缩写为 `en`、`ena`、`enab` 或 `enabl`。再如 `Router#` 模式下以 `con` 开头的命令只有 `config`, 所以 `config` 可以缩写为 `con`、`conf` 或 `confi` 等。因此 `Router#config Terminal` 就可以缩写为 `Router#con t` 等。

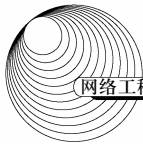
交换机的另外一个重要配置就是配置虚拟局域网 (VLAN), 本书将在 3.2.6 小节集中讨论有关 VLAN 的配置。

3.3.2 配置和管理 VLAN

VLAN 技术是交换技术的重要组成部分, 也是交换机的重要进步之一。它用以把物理上直接相连的网络从逻辑上划分为多个子网。每一个 VLAN 对应着一个广播域, 处于不同 VLAN 上的主机不能进行通信, 不同 VLAN 之间的通信要引入第三层交换技术才可以解决。对虚拟局域网的配置和管理主要涉及 VLAN 中继、VTP 协议和 VLAN 的配置。

VLAN 中继 (VLAN Trunk) 也称为 VLAN 主干, 是指在交换机与交换机或交换机与路由器之间连接的情况下, 在互相连接的端口上配置中继模式, 使得属于不同 VLAN 的数据帧都可以通过这条中继链路进行传输。

VLAN 中继协议 (即 VTP 协议) 可以帮助交换机设置 VLAN。VTP 协议可以维护 VLAN 信息全网的一致性。VTP 有三种工作模式, 即服务器模式、客户模式和透明模式, 其中服务器模式下可以设置 VLAN 信息, 服务器会自动将这些信息广播到网上的其他交换机以统一配置; 客户模式下交换机不能配置 VLAN 信息, 只能被动接受服务器的 VLAN 配置; 而透明模式下是独立配置, 可以配置 VLAN 信息, 但是不广播自己的 VLAN 信息, 同时它接收到服务器发



来的 VLAN 信息后并不使用，而是直接转发给别的交换机。

交换机的初始状态是工作在透明模式，有一个默认的 VLAN，所有端口都属于这个 VLAN 内。

1. 划分 VLAN 的方法

虚拟局域网是交换机的重要功能，通常虚拟局域网的实现形式有三种，即静态端口分配、动态虚拟网和多虚拟网端口配置。

静态虚拟网的划分通常是网管人员使用网管软件或直接设置交换机的端口，使其直接从属某个虚拟网。这些端口一直保持这些从属性，除非网管人员重新设置。这种方法虽然比较麻烦，但比较安全，容易配置和维护。

支持动态虚拟网的端口可以借助智能管理软件自动确定它们的从属。端口借助信息包的 MAC 地址、逻辑地址或协议类型来确定虚拟网的从属。当一网络节点刚连接入网时，交换机端口还未分配，于是交换机通过读取网络节点的 MAC 地址动态地将该端口划入某个虚拟网。这样一旦网管人员配置好后，用户的计算机可以灵活地改变交换机端口，而不会改变该用户的虚拟网从属性。

多虚拟网端口配置支持一用户或一端口可以同时访问多个虚拟网。这样可以将一台网络服务器配置成多个业务部门（每种业务设置成一个虚拟网）都可同时访问，也可以同时访问多个虚拟网的资源，还可让多个虚拟网间的连接只需一个路由端口即可完成。但这样会带来安全方面的隐患。

静态虚拟网是使用最普遍的一种划分 VLAN 的方法，这里以该方法为例介绍 VLAN 配置的知识。

2. 配置 VTP 协议

为了让读者清楚地了解 VTP（VLAN Trunking Protocol）协议的工作情况以及如何来配置 VTP 协议，这里结合一个综合实例拓扑结构（如图 3-17 所示）来配置 VTP 协议及跨交换机的 VLAN。这里用交叉双绞线把 2950A 交换机的 FastEthernet0/24 端口和 2950B 交换机的 FastEthernet0/24 端口连接起来，作为两交换机间的 Trunk 线路。

配置 2950A 交换机为服务器模式，操作如下。

```
Switch>enable                (进入特权模式)
Switch#config terminal        (进入配置子模式)
Switch(config)#hostname 2950A (修改主机名为 2950A)
2950A(config)#end
```

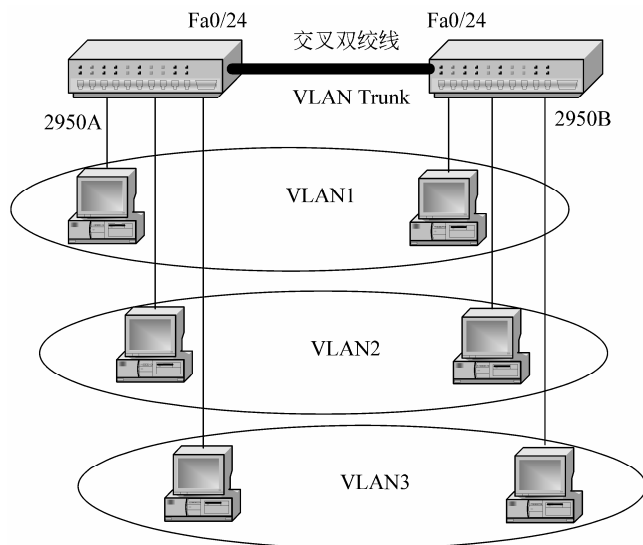
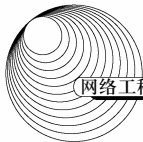


图 3-17 VLAN 拓扑结构图

```

2950A#
2950A #vlan database                (进入 VLAN 配置子模式)
2950A (vlan)#vtp ?                  (查看和 VTP 配合使用的命令)
2950A (vlan)#vtp server              (设置本交换机为 Server 模式)
Setting device to VTP SERVER mode.
2950A (vlan)#vtp domain vtpserver   (设置域名)
Changing VTP domain name from NULL to vtpserver
2950A (vlan)#vtp pruning             (启动修剪功能)
Pruning switched ON
2950A (vlan)#exit                    (退出 VLAN 配置模式)
APPLY completed.
Exiting....
2950A #show vtp status               (查看 VTP 设置信息)
VTP Version                        : 2
Configuration Revision              : 0
Maximum VLANs supported locally     : 64
Number of existing VLANs           : 1
VTP Operating Mode                  : Server
VTP Domain Name                     : vtpserver
VTP Pruning Mode                    : Enable

```



```
VTP V2 Mode : Disabled
VTP Traps Generation : Disabled
MD5 digest : 0x82 0x6B 0xFB 0x94 0x41 0xEF 0x92 0x30
Configuration last modified by 0.0.0.0 at 3-1-93 00:02:51
2950A #
```

配置 2950B 交换机为客户机模式，则它会从服务器（2950A）那里学习到 VTP 的其他信息及 VLAN 信息。

```
Switch#config terminal (进入配置子模式)
Switch(config)#hostname 2950B (修改主机名为 2950B)
2950B(config)#end
2950B#vlan database
2950B(vlan)#vtp client
Setting device to VTP CLIENT mode.
2950B(vlan)#exit
```

3. 配置 VLAN Trunk 端口

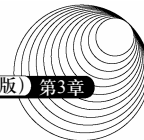
跨交换机的同一 VLAN 内的数据经过 Trunk 线路进行交换，默认情况下，trunk 允许所有 VLAN 通过。使用 `switchport trunk allowed vlan remove vlan-list` 可以去掉某一 VLAN，在交换机 2950A 和 2950B 上做如下相同的配置操作。

```
Switch#config
Configuring from terminal, memory, or network [terminal]?
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)#interface f0/24 (进入端口 24 配置模式)
Switch(config-if)#switchport mode trunk (设置当前端口为 Trunk 模式)
Switch(config-if)# switchport trunk allowed vlan all (设置允许从该端口交换数据的 VLAN)
Switch(config-if)#exit
Switch(config)#exit
Switch#
```

4. 创建 VLAN

VLAN 信息可以在服务器模式或透明模式交换机上创建。这里在 2950A 交换机上创建两个 VLAN。

```
2950A#vlan database
2950A (vlan)#vlan 2 (创建一个 VLAN2)
```



```
VLAN 2 added:
  Name: VLAN0002                (系统自动命名)
2950A (vlan)#vlan 3 name vlan3  (创建一个 VLAN3, 并命名为 vlan3)
VLAN 3 added:
  Name: vlan3
2950A (vlan)#exit
```

5. 将端口加入某个 VLAN

配置完 VTP 协议及 VLAN Trunk 端口后就可以设置将端口归属于哪个 VLAN。在交换机 2950A 和 2950B 上做如下相同的配置操作, 则 vlan2 中包含两个交换机的 fa0/9 端口, vlan3 中包含两个交换机的 f0/10 端口, 其余端口可以做类似设置。除了加入 vlan2 和 vlan3 的端口外, 其余各端口均属于 vlan1 (交换机默认的 VLAN)。

```
Switch#config terminal
Enter configuration commands, one per line.  End with CNTL/Z.
Switch(config)#interface f0/9                (进入端口 9 的配置模式)
Switch(config-if)#switchport mode access     (设置端口为静态 VLAN 访问模式)
Switch(config-if)#switchport access vlan 2   (把端口 9 分配给相信的 vlan2)
Switch(config-if)#exit
Switch(config)#interface f0/10
Switch(config-if)#switchport mode access
Switch(config-if)#switchport access vlan 3
Switch(config-if)#exit
Switch(config)#exit
Switch#show vlan                            (查看 VLAN 配置信息)
(结果省略)
Switch#
```

3.3.3 路由器

1. 路由器概述

计算机网络中有非常多的主机, 这些主机分布在不同的局域网中, 如果没有能够连接不同局域网的设备, 那么这些主机之间就不可能进行通信。路由器就是这样一种能够将多个局域网相连接的网络设备, 如图 3-18 所示。

互联网络中有大量路由器, 用来连接各个不同的局域网。路由器可以学习和传播各种路由信息, 并根据这些路由信息将网络中的分组转发到正确的网络中。

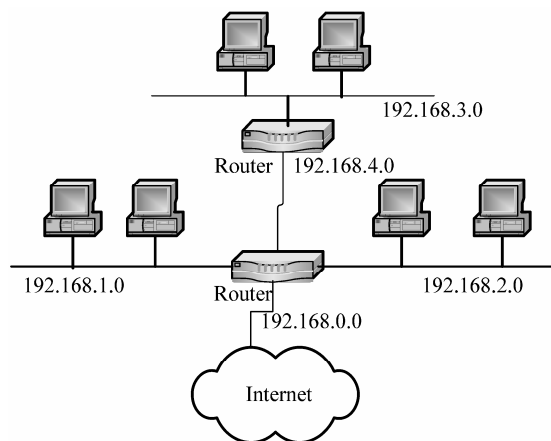
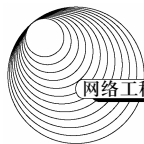


图 3-18 路由器

路由器是工作在 OSI 七层模型第三层(网络层)的设备,其具有局域网和广域网两种接口。它可以作为企业内部网络和 Internet 骨干网络的连接设备来使用。路由器通过路由表为进入路由器的数据分组选择最佳的路径,并将分组传输到适当的出口。

2. 路由器的功能

路由器主要有三种功能,即网络互联、网络隔离和流量控制。

1) 网络互联

路由器的主要功能是实现网络互联,它主要采用以下技术来实现不同网络之间的数据包传输。

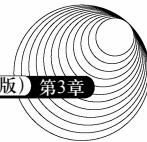
地址映射:地址映射技术可以完成逻辑地址(IP地址)与物理地址(MAC地址)之间的转换,从而完成数据在同一网段内的传输。

路由选择:每个路由器都会保持着一个独立的路由表,该路由表根据数据包中的目的IP地址判断该数据包应该送往的下一个路由器的地址。路由表分为静态和动态两种,建立和维护更新路由表是路由器完成路由选择的关键。

协议转换:路由器可以连接不同结构的局域网。不同结构的局域网要进行连接,需要连接设备能够实现协议的转换(如IP协议向IPX协议之间的转换)。

2) 网络隔离

路由器一方面用来连接各个局域网,保证各个局域网之间的通信,另一方面路由器可以根据数据包的源地址、目的地址、数据包类型等对数据包能否被转发作出适当的判断,从而隔离各个局域网之间不需要传输的数据包。这种隔离能够将各个局域网中的广播风暴隔离在每个局



域网之内,防止局域网中的广播风暴影响到整个网络的性能;同时能够保证网络的安全,将不必要的数据流量隔离,以保证网络的安全。

3) 流量控制

路由器具有非常好的流量控制能力,它可以利用相应的路由算法来均衡网络负载,从而有效控制网络拥塞,避免因拥塞而导致网络性能下降。

3. 路由表

路由器的主要工作就是为经过路由器的每个数据帧寻找一条最佳传输路径,并将该数据有效地传送到目的站点。由此可见,选择最佳路径的策略即路由算法是路由器的关键所在。

为了完成这项工作,路由器中保存着各种传输路径的相关数据——路由表(Routing Table),供路由选择时使用。打个比方,路由表就像平时使用的地图一样,标识着各种路线,路由表中保存着子网的标志信息、网上路由器的个数和下一个路由器的名字等内容。路由表可以由系统管理员固定设置好的,也可以由系统动态修改;可以由路由器自动调整,也可以由主机控制。

1) 静态路由表

由系统管理员事先设置好的固定的路由表称为静态(static)路由表,一般是在系统安装时就根据网络的配置情况预先设定的,它不会随网络结构的改变而改变。

2) 动态路由表

动态(Dynamic)路由表是路由器根据网络系统的运行情况而自动生成的路由表。路由器根据路由选择协议(Routing Protocol)提供的功能,自动学习和记忆网络运行情况,在需要时自动计算数据传输的最佳路径。

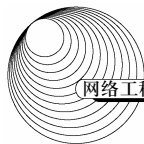
4. 路由选择协议

路由选择协议是一种网络层协议,它通过提供一种共享路由选择信息的机制,允许路由器与其他路由器通信以更新和维护自己的路由表,并确定最佳的路由选择路径。通过路由选择协议,路由器可以了解未直接连接的网络的状态,当网络发生变化时,路由表中的信息可以随时更新,以保证网络上的路由选择路径处于可用状态。

路由表由路由协议生成。路由协议根据其生成路由表的方式,可以分为静态路由协议和动态路由协议两种。静态路由协议下的路由信息完全由管理员手动完成,在完成路由表后,除非管理员再次调整路由信息,否则不会发生变化。而动态路由协议可以根据网络的状态变化而不断地调整路由器中的路由表,以使路由表能够随时保持最新的状态,保证信息包能够顺利转发。

1) 静态路由协议

在静态路由协议下,路由信息由管理员配置而成,它适用于小型的局域网络(拥有5台以



下的路由器)。静态路由协议具有运行速度快、占用资源少、配置方法简单的特点,但是由于静态路由需要管理员手动配置,如果在网络中的状态发生变化,需要修改路由信息,那么对于网络管理员来说,将会是非常大的工作量,同时管理和配置的难度也较大。所以,静态路由协议在小型的网络中能够工作得很好,但在较大规模的网络中并不能够很好地运行和维护。

2) 动态路由协议

动态路由协议根据路由信息更新方式的不同,可以分为距离矢量路由协议和链路状态路由协议两种。

距离矢量(Distance-vector)路由协议采用距离矢量路由选择算法,确定到网络中任一链路的方向(向量)与距离,如RIP协议。

链路状态(Link-state)路由协议创建整个网络的准确拓扑,以计算路由器到其他路由器的最短路径,如OSPF、IS-IS等。

3.3.4 路由器的配置

1. 路由器的基本配置

与交换机的配置类似,路由器的配置操作模式有普通用户模式、特权模式和配置模式。在用户模式下,用户只能发出有限的命令,这些命令对路由器的正常工作没有影响;在特权模式下,用户可以发出丰富的命令,以便更好地控制和使用路由器;在配置模式下,用户可以创建和更改路由器的配置,对路由器的管理和配置主要工作在配置模式下。

其中,配置模式又分为全局配置模式和接口配置模式、路由协议配置模式、线路配置模式等子模式。在不同的工作模式下,路由器有不同的命令提示状态。

配置路由器的连接方式如图3-19所示,使用专用的配置线缆将路由器的Consol端口(配置端口)与计算机的串行口(RS232接口)相连,然后打开计算机中的超级终端进行连接。超级终端的配置和参数的设置方式可参见3.5.2小节中所述。在路由器中同样可以配置使能口令(enable password)和使能密码(enable secret),一般情况下只需配置一个,当两者同时配置时,后者生效。这两者的区别是使能口令以明文显示,而使能密码以密文形式显示。主机名及路由器口令的设置和3.3.1小节对交换机配置的主机名及口令相同,这里不再复述。

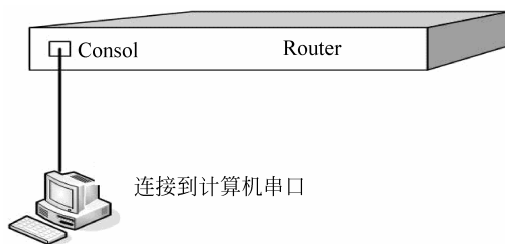
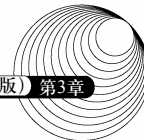


图 3-19 路由器配置连接方式



配置路由器以太网接口, 路由器一般提供 1 个或多个以太网接口槽, 每个槽上会有 1 个以上以太网接口。以太网接口也因此而命名为 {Ethernet 槽位/端口} 或 {FastEthernet 槽位/端口}。比如 FastEthernet0/0、FastEthernet1/1, 也可缩写为 F0/0、F1/1。

以 Cisco2600 系列交换机为例, 连接好仿真终端到路由器的 Console 电缆线, 就可以对路由器进行初始的配置工作。以太网接口配置如下。

```
Router>enable                                (进入特权执行模式)
Router #config t                              (进入全局配置模式)
Enter configuration commands, one per line.  End with CNTL/Z.
Router (config)#interface FastEthernet0/1    (进入接口 F0/1 配置模式)
Router (config-if)#ip address 192.168.1.11 255.255.255.0 (设置接口 IP 地址)
Router (config-if)#no shutdown               (激活接口)
10:05:01:%LINK-3-UPDOWN:Interface FastEthernet0/0, changed state to up
Router (config-if)#end                       (退回特权模式)
Router#show running-config                  (检查配置结果)
```

2. 静态路由的配置

通过配置静态路由, 用户可以人为地指定对某一网络访问时所要经过的路径, 网络结构比较简单, 且一般到达某一网络所经过的路径唯一的情况下采用静态路由。下面通过一个实例介绍设置静态路由、查看路由表, 理解路由原理及概念。

如图 3-20 所示设计拓扑结构, 3 台路由器分别命名为 R1、R2、R3, 所使用的接口和相应的 IP 地址分配如图 3-20 所示, 其中 “/24” 表示子网掩码为 24 位, 即 255.255.255.0。

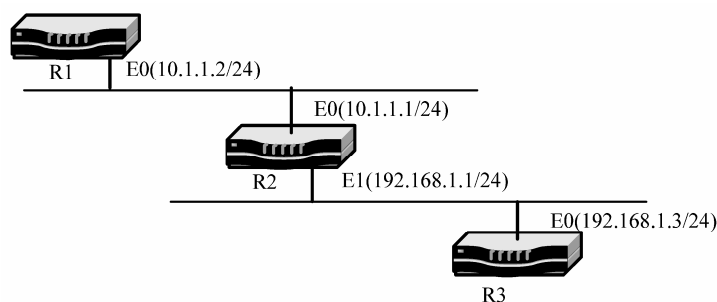
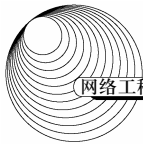


图 3-20 静态路由实例图

这里要在路由器中配置静态路由以实现路由器 R2 到 R3 在 IP 层的连通性, 也就是要求从 R2 可以 ping 通 R3, 从 R3 可以 ping 通 R1。



首先根据拓扑结构图配置各路由器的以太网接口,配置好以太网接口后测试路由器间基本的连通性。首先从 R1 路由器 ping 路由器 R2 和 R3。

```
R1#ping 10.1.1.2          (ping R2, 结果连通)
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 10.1.1.2, time out is 2 seconds:
!!!!!!
Success rate is 100 percent(5/5), round-trip min/avg/max=4/4/4 ms
R1#ping 192.168.1.3       (ping R3, 结果连通)
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 192.168.1.3, time out is 2 seconds:
!!!!!!
Success rate is 100 percent(5/5), round-trip min/avg/max=4/4/4 ms
```

然后从 R2 路由器 ping 路由器 R1 的 E1 接口。

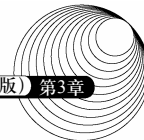
```
R2#ping 192.168.1.1       (ping R1 的 E1 接口, 结果不连通)
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 192.168.1.1, time out is 2 seconds:
!!!!!!
Success rate is 0 percent(0/5)
R2#show ip route          (查看路由表)
(省略解释信息)
Gateway of last resort is not set
  10.0.0.0/24 is subnetted, 1 subnets
C       10.1.1.0 is directly connected, Ethernet0
```

可见从 R2 到 R1 的 E1 接口是跨网段的 ping, 没有 ping 通, 查看路由表发现路由表中显示只有直接相连的网段 10.1.1.0/24 在其路由表内, 标志为“C”表示连接 (connected), 为此可以在 R2 路由表中加入静态路由。

```
R2#config t
R2(config)#ip route 192.168.1.0 255.255.255.0 10.1.1.1    (加入静态路由)
R2(config)#end
```

这条静态路由信息表示从该路由器出发发往 192.168.1.0 255.255.255.0 网段的数据包下一跳点 (Next Hop) 的地址是 10.1.1.1 (即通过 R1 的 E0 接口地址)。接下来再查看路由表信息, 发现路由表中多了一条路由信息, 标记“S”表示静态路由 (static)。

```
  192.168.0.0/24 is subnetted, 1 subnets
S       192.168.1.0 [1/0] via 10.1.1.1
  10.0.0.0/24 is subnetted, 1 subnets
```



```
C 10.1.1.0 is directly connected, Ethernet0
```

这时再从 R2 ping R1 的 E1 接口, 发现可以 ping 通了。

```
R2#ping 192.168.1.1          (ping R1 的 E1 接口, 结果连通)
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 192.168.1.1, time out is 2 seconds:
!!!!!!
Success rate is 100 percent(5/5), round-trip min/avg/max=4/4/4 ms
```

但是此时从 R2 ping R3 的 E0 接口却不能 ping 通。原因是, 从 R2 发出的数据包可以经过 R1 的 E0 端口转发而让 R3 收到, 但是 R3 此时不能发送响应数据包 ICMP Echo Reply 到 10.1.1.2, 因为在 R3 上还没有发往 R2 路由器的路由信息。这可以通过 debug ip packet 和 debug ip icmp 命令来监视 IP 包和 ICMP 包的传输情况而知。

```
R2#ping 192.168.1.3          (ping R3 的 E0 接口, 结果不连通)
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 192.168.1.1, time out is 2 seconds:
!!!!!!
Success rate is 0 percent(0/5)
```

此时需要在 R3 上加入发往 R2 网段数据包的路由信息。

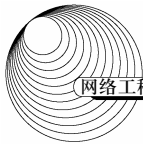
```
R3#config t
R3(config)#ip route 10.1.1.0 255.255.255.0 192.168.1.1    (加入静态路由)
R3(config)#end
```

这样一来就可以实现路由器 R2 到 R3 之间的连通性, 互相可以 ping 通对方接口的 IP 地址。应该注意的是, 在有些路由器上默认情况是不启动 IP 路由的, 这时可以用 ip routing 和 no ip routing 来启动和关闭 IP 路由。

3.3.5 配置路由协议

本节主要讲述对路由协议的配置。IP 路由选择协议用有效、无循环的路由信息填充路由表, 从而为数据包在网络之间传递提供了可靠的路径信息。路由选择协议又分为距离矢量、链路状态和平衡混合三种。

距离矢量 (Distance Vector) 路由协议计算网络中所有链路的矢量和距离, 并以此为依据确认最佳路径。使用距离矢量路由协议的路由器定期向其相邻的路由器发送全部或部分路由表。典型的距离矢量路由协议有 RIP 和 IGRP。



链路状态（Link State）路由协议使用为每个路由器创建的拓扑数据库来创建路由表，每个路由器通过此数据库建立一个整个网络的拓扑图。在拓扑图的基础上通过相应的路由算法计算出通往各目标网段的最佳路径，并最终形成路由表。典型的链路状态路由协议是 OSPF（Open Shortest Path First，开放最短路径优先）路由协议。

平衡混合（Balanced Hybrid）路由协议结合了链路状态和距离矢量两种协议的优点，此类协议的代是 EIGRP，即增强型内部网关协议。

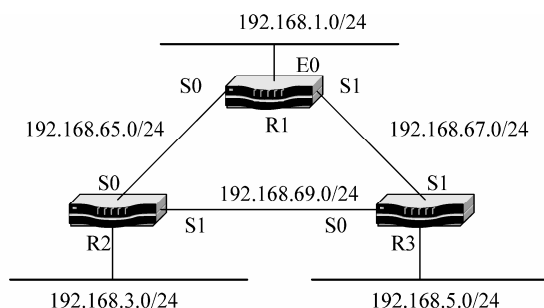
下面将分别讨论如何在路由器中配置 RIP（路由选择信息协议）和 OSPF 动态路由协议。

1. 配置 RIP 协议

RIP 是距离矢量路由选择协议的一种。路由器收集所有可到达目的地的不同路径，并且保存有关到达每个目的地的最少站点数的路径信息，除到达目的地的最佳路径外，任何其他信息均予以丢弃。同时，路由器也把所收集的路由信息用 RIP 协议通知相邻的其他路由器。这样，正确的路由信息逐渐扩散到了全网。

RIP 使用非常广泛，它简单、可靠，便于配置。RIP 版本 2 还支持无类域间路由（Classless Inter-Domain Routing, CIDR）、可变长子网掩码（Variable Length Subnetwork Mask, VLSM）和不连续的子网，并且使用组播地址发送路由信息。但是 RIP 只适用于小型的同构网络，因为它允许的最大跳数为 15，任何超过 15 个站点的目的地均被标记为不可达。RIP 每隔 30s 广播一次路由信息。

假设有图 3-21 所示的网络拓扑结构，试通过配置 RIP 协议使全网连通。在配置之前先按照拓扑结构连接好网络设备，其中串口之间需要用 DTE（数据终端设备）和 DCE（数据通信设备）电缆对接，或者用 DCE 转 DTE 电缆连接。



图中各接口 IP 地址分配如下：

R1: E0 192.168.1.1
R1: S0 192.168.65.1
R1: S1 192.168.67.1

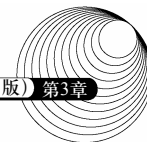
R2: E0 192.168.3.1
R2: S0 192.168.65.2
R2: S1 192.168.69.1

R3: E0 192.168.5.1
R3: S0 192.168.69.2
R3: S1 192.168.67.2

图 3-21 RIP 协议配置拓扑图

首先根据图中要求配置各路由器各接口地址。

R1#config t



```

R1 (config)#no logging console
R1 (config)#interface FastEthernet0/1
R1 (config-if)#ip address 192.168.1.1 255.255.255.0
R1 (config-if)#no shutdown
R1 (config-if)#exit
R1 (config)#interface serial 0
R1 (config-if)#ip address 192.168.65.1 255.255.255.0
R1 (config-if)#no shutdown
R1 (config-if)#exit
R1 (config)#interface serial 1
R1 (config-if)#ip address 192.168.67.1 255.255.255.0
R1 (config-if)#no shutdown

```

在全局配置模式下使用 `no logging console` 配置命令, 可以防止大量的端口状态变化信息和报警信息对配置过程的影响。为了查明串行接口所连接的电缆类型, 从而正确配置串行接口, 可以使用 `show controllers serial` 命令来查看相应的控制器。注意在配置端口时使用 `no shutdown` 命令, 因为默认情况下各物理接口是处于关闭状态的, 配置完成需要对端口进行激活。

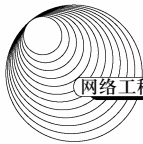
使用类似配置 R1 各接口地址的方法, 可以配置好路由器 R2 和 R3 的各接口地址。此时路由表中只有和路由器直接相连的各网段的路由信息, 即每个路由器只可以 ping 通和它直接相连的路由器的接口。此时可以用 `show ip route` 命令查看路由表信息。

```

R1#show ip route
Codes: C – connected, S – static, I – IGRP, R – RIP, M – mobile, B – BGP
       D – EIGRP, EX – EIGRP external, O – OSPF, IA – OSPF inter area
       N1 – OSPF NSSA external type 1, N2 – OSPF NSSA external type 2
       E1 – OSPF external type 1, E2 – OSPF external type 2, E – EGP
       I – IS-IS, L1 – IS-IS level-1, L2 – IS-IS level-2, ia – IS-IS inter area
       * – candidate default U – per-user static route, o – ODR
       P – periodic downloaded static route
Gateway of last resort is not set
  192.168.0.0/24 is subnetted, 3 subnets
C       192.168.1.0 is directly connected, Ethernet0
C       192.168.65.0 is directly connected, Serial0
C       192.168.67.0 is directly connected, Serial1

```

配置完接口地址后就可以进行 RIP 协议配置, RIP 协议配置非常简单。首先使用 `ip routing` 允许路由选择协议, 在有些路由器上默认情况是关闭的。用 `router rip` 命令进入 RIP 协议配置模



式, 然后使用 **network** 语句声明进入 RIP 进程的网络就可以了。

配置路由器 R1 如下。

```
R1 (config)#ip routing
R1 (config)#router rip                (进入 RIP 协议配置子模式)
R1 (config-router)#network 192.168.1.0 (声明网络 192.168.1.0/24)
R1 (config-router)#network 192.168.65.0
R1 (config-router)#network 192.168.67.0
R1 (config-router)#version 2          (设置 RIP 协议版本 2)
R1 (config-router)#exit
```

配置路由器 R2 如下。

```
R1 (config)#ip routing
R1 (config)#router rip                (进入 RIP 协议配置子模式)
R1 (config-router)#network 192.168.3.0 (声明网络 192.168.3.0/24)
R1 (config-router)#network 192.168.65.0
R1 (config-router)#network 192.168.69.0
R1 (config-router)#version 2          (设置 RIP 协议版本 2)
R1 (config-router)#exit
```

配置路由器 R3 如下。

```
R1 (config)#ip routing
R1 (config)#router rip                (进入 RIP 协议配置子模式)
R1 (config-router)#network 192.168.5.0 (声明网络 192.168.5.0/24)
R1 (config-router)#network 192.168.67.0
R1 (config-router)#network 192.168.69.0
R1 (config-router)#version 2          (设置 RIP 协议版本 2)
R1 (config-router)#exit
```

配置完 RIP 协议后, RIP 协议的路由器广播自己的路由信息到周边路由器, 此时各路由器就可以学习到其他路由器的路由信息。此时再查看路由信息则有所不同, 如下查看 R3 上的路由表。

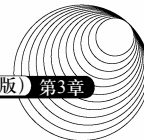
```
R3#show ip route
```

Codes: C – connected, S – static, I – IGRP, R – RIP, M – mobile, B – BGP

D – EIGRP, EX – EIGRP external, O – OSPF, IA – OSPF inter area

N1 – OSPF NSSA external type 1, N2 – OSPF NSSA external type 2

E1 – OSPF external type 1, E2 – OSPF external type 2, E – EGP



I – IS-IS, L1 – IS-IS level-1, L2 – IS-IS level-2, ia – IS-IS inter area

* - candidate default U – per-user static route, o - ODR

P – periodic downloaded static route

Gateway of last resort is not set

192.168.0.0/24 is subnetted, 6 subnets

```
C      192.168.1.0 is directly connected, Ethernet0
C      192.168.65.0 is directly connected, Serial0
C      192.168.67.0 is directly connected, Serial1
R      192.168.65.0 [120/1] via 192.168.67.1, 00:00:15, Serial
      [120/1] via 192.168.69.1, 00:00:24, Serial0
R      192.168.1.0 [120/1] via 192.168.67.1, 00:00:15, Serial
R      192.168.3.0 [120/1] via 192.168.69.1, 00:00:24, Serial0
```

路由表中的项目解释如下。

R: 192.168.3.0 [120/1] via 192.168.69.1, 00:00:24, Serial0

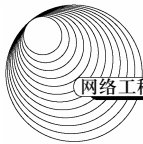
- R: 表示此项路由是由 RIP 协议获取的, 另外 C 代表直接相连的网段。
- 192.168.3.0: 表示目标网段。
- [120/1]: 120 表示 RIP 协议的管理距离默认为 120, 1 是该路由的度量值, 即跳数。
- via: 经由的意思。
- 192.168.69.1: 表示从当前路由器出发到达目标网的下一跳点的 IP 地址。
- 00:00:24: 表示该条路由产生的时间。
- Serial0: 表示该条路由使用的接口。

从路由表中可以看出多了 3 条 RIP 路由信息, 这 3 条路由信息分别可以访问到网络 192.168.65.0/24、192.168.1.0/24、192.168.3.0/24, 其中访问到 192.168.65.0/24 的路由有 2 条。之所以保存两条路由信息, 是因为到达目的网段需要经过的跳数相同(都为 1)。访问另外两个网段也可以经过另外两个路由器来转发, 但是因为那样要经过 2 跳, 而 RIP 协议是选择跳数作为唯一度量路由选择的标准, 所以它只将跳数最少路径保留在路由表中, 而其余路径都被放弃。

另外, 因为 RIP 版本 2 支持不连续子网和可变长子网掩码, 所以各网段的 IP 地址可以是任何形式的合法 IP。通过对各路由器配置 RIP 协议, 可以达到全网连通的目的。

2. 配置 OSPF 协议

开放最短路径优先(Open Shortest Path First, OSPF)协议是重要的路由选择协议。它是一种链路状态路由选择协议, 是由 Internet 工程任务组开发的内部网关(Interior Gateway Protocol, IGP)路由协议, 用于在单一自治系统(Autonomous System, AS)内决策路由。



链路是路由器接口的另一种说法,因此 OSPF 也称为接口状态路由协议。OSPF 通过路由器之间通告网络接口的状态来建立链路状态数据库,生成最短路径树,每个 OSPF 路由器使用这些最短路径构造路由表。下面分别介绍 OSPF 协议的相关要点。

- 自治系统。自治系统包括一个单独管理实体下所控制的一组路由器,OSPF 是内部网关路由协议,工作于自治系统内部。
- 链路状态。所谓链路状态,是指路由器接口的状态,如 Up、Down、IP 地址、网络类型以及路由器和它邻接路由器间的关系。链路状态信息通过链路状态通告(Link State Advertisement, LSA)扩散到网上的每台路由器。每台路由器根据 LSA 信息建立一个关于网络的拓扑数据库。
- 最短路径优先算法。OSPF 协议使用最短路径优先算法,利用从 LSA 通告得来的信息计算每一个目标网络的最短路径,以自身为根生成一棵树,包含了到达每个目的网络的完整路径。
- 路由标示。OSPF 的路由标示是一个 32 位的数字,它在自治系统中被用来唯一识别路由器。默认使用最高回送地址,若回送地址没有被配置,则使用物理接口上最高的 IP 地址作为路由器标示。
- 邻居和邻接。OSPF 在相邻路由器间建立邻接关系,使它们交换路由信息。邻居是指共享同一网络的路由器,并使用 Hello 包来建立和维护邻居路由器间的关系。
- 区域。OSPF 网络中使用区域(Area)来为自治系统分段。OSPF 是一种层次化的路由选择协议,区域 0 是一个 OSPF 网络中必须具有的区域,也称为主干区域,其他所有区域要求通过区域 0 互连到一起。

设计如图 3-22 所示的网络拓扑结构图来配置 OSPF 协议。

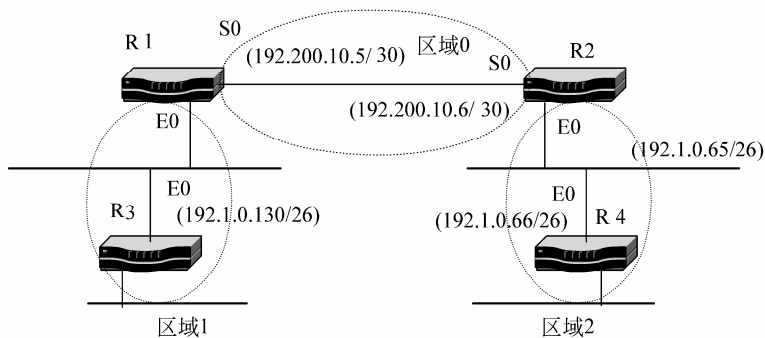
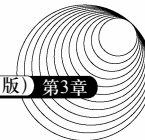


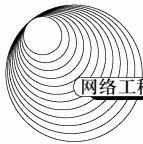
图 3-22 OSPF 协议配置实例图

R1:
interface ethernet 0



```
ip address 192.1.0.129 255.255.255.192
!
interface serial 0
ip address 192.200.10.5 255.255.255.252
!
router ospf 100
network 192.200.10.4 0.0.0.3 area 0
network 192.1.0.128 0.0.0.63 area 1
!
R2:
interface ethernet 0
ip address 192.1.0.65 255.255.255.192
!
interface serial 0
ip address 192.200.10.6 255.255.255.252
!
router ospf 200
network 192.200.10.4 0.0.0.3 area 0
network 192.1.0.64 0.0.0.63 area 2
!
R3:
interface ethernet 0
ip address 192.1.0.130 255.255.255.192
!
router ospf 300
network 192.1.0.128 0.0.0.63 area 1
!
R4:
interface ethernet 0
ip address 192.1.0.66 255.255.255.192
!
router ospf 400
network 192.1.0.64 0.0.0.63 area 1
!
```

在上述配置中，首先对每台路由器接口进行配置。接口配置完成后使用 `router ospf 100` 命令启动一个 OSPF 路由选择协议进程，其中“100”为进程号，不同于 IGRP 的是，这里每台路由器的进程号并不需要一致。最后使用 `network` 将相应的网段加入 OSPF 路由进程，则此接口



所对应的网段就加入了 OSPF 进程中。

随后就可以和配置其他协议一样,用如下命令来调试或查看配置信息和路由信息。

```
debug ip ospf events
debug ip ospf packet
show ip ospf
show ip ospf database
show ip ospf interface
show ip ospf neighbor
show ip route
```

3.4 综合布线

3.4.1 综合布线系统概述

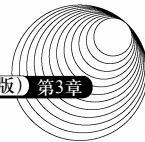
1. 什么是综合布线系统

综合布线系统 (Premises Distribution System, PDS) 又称结构化综合布线系统 (Structured Cabling Systems, SCS)。综合布线系统是为通信与计算机网络而设计的,它可以满足各种通信与计算机信息传输的要求,是为具有综合业务需求的计算机数据网开发的。综合布线系统具体的应用对象主要是通信和数据交换,即话音、数据、传真、图影像信号。综合布线系统是一套综合系统,它可以使用相同的线缆、配线端子板、相同的插头及模块插孔,解决传统布线存在的兼容性问题。综合布线系统是建筑智能化大厦工程的重要组成部分,是智能化大厦传送信息的神经中枢。

2. 综合布线系统的特点

综合布线系统是信息技术和信息产业大规模高速发展的产物,是布线系统的一项重大革新,它和传统布线系统比较,具有明显的优越性,具体表现在以下 6 个方面。

(1) 兼容性。其设备可以用于多种系统。沿用传统的布线方式,会使各个系统的布线互不相容,管线拥挤不堪,规格不同,配线插接头型号各异,所构成的网络内的管线与插接件彼此不同而不能互相兼容,一旦要改变终端机或话音设备位置,势必重新敷设新的管线和插接件。而综合布线系统不存在上述问题,它将语音、数据信号的配线统一设计规划,采用统一的传输线、信息插接件等,把不同信号综合到一套标准布线系统中。同时,该系统相比于传统布线大为简化,不存在重复投资,可以节约大量资金。



(2) 开放性。对于传统布线,一旦选定了某种设备,也就选定了布线方式和传输介质,如要更换一种设备,原有布线将全部更换,这样极为麻烦,又增加大量资金。综合布线系统采用开放式体系结构,符合国际标准,对现有著名厂商的硬件设备均是开放的,对通信协议也同样是开放的。

(3) 灵活性。传统布线各系统是封闭的,体系结构是固定的,若增减设备将十分困难。而综合布线系统,所有传递信息线路均为通用的,即每条线路均可传送语音、传真和数据,所用系统内的设备(计算机、终端、网络集散器、HUB或MAU、电话、传真)的开通及变动无须改变布线,只要在设备间或管理间作相应的跳线操作即可。

(4) 可靠性。传统布线各系统互不兼容,因此在一个建筑物内存在多种布线方式,形成各系统交叉干扰,这样各个系统可靠性降低,势必影响整个建筑系统的可靠性。综合布线系统布线采用高品质的材料和组合压接方式构成一套高标准的信息网络,所有线缆与器件均通过国际上的各种标准,保证了综合布线系统的电气性能。综合布线系统全部使用物理星型拓扑结构,任何一条线路有故障都不会影响其他线路,从而提高了可靠性,各系统采用同一传输介质,互为备用,又提高了备用冗余。

(5) 经济性。综合布线系统设计信息点时要求按规划容量,留有适当的发展容量,因此,就整体布线系统而言,按规划设计所做的经济分析表明,综合布线系统会比传统的价格性能比更优,后期运行维护及管理费也会下降。

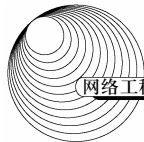
(6) 先进性。随着信息时代快速发展,数据传递和话音传送并驾齐驱,多媒介技术的迅速崛起,如仍采用传统布线,在技术上太落后。综合布线系统采用双绞线与光纤混合布置方式是比较科学和经济的方式。

3. 综合布线标准

综合布线的标准很多,但在实际工程项目中,并不需要涉及所有标准和规范,而应根据布线项目性质和涉及的相关技术工程情况适当引用标准规范。通常来说,布线方案设计应遵循布线系统性能和系统设计标准,布线施工工程应遵循布线测试、安装、管理标准及防火、机房及防雷接地标准。

例如一个典型的办公网络的布线系统,集成方案中通常采用如下标准。

- 国家标准《建筑与建筑群综合布线系统工程设计规范》GB 30511—2000。
- 国家标准《建筑与建筑群综合布线系统工程施工和验收规范》GB 30512—2000。
- 《大楼通信综合布线系统第一部分总规范》YD/T 926.1-2001。
- 《大楼通信综合布线系统第二部分综合布线用电缆光纤技术要求》YD/T 926.2—2001。
- 《大楼通信综合布线系统第三部分综合布线用连接硬件技术要求》YD/T 926.3—2001。



- 北美标准 ANSI/TIA/EIA 568B 《商用建筑通信布线标准》。
- 国际标准 ISO/IEC 11801 《信息技术——用户通用布线系统》（第2版）。
- 《国际电子电气工程师协会：CSMA/CD 接口方法》 IEEE 802.3。

4. 综合布线系统的构成

综合布线系统由6个子系统组成，即建筑群子系统、设备间子系统、垂直子系统、管理子系统、水平子系统和工作区子系统。大型布线系统需要用铜介质和光纤介质部件将6个子系统集成在一起。综合布线系统的6个子系统的构成如图3-23所示。

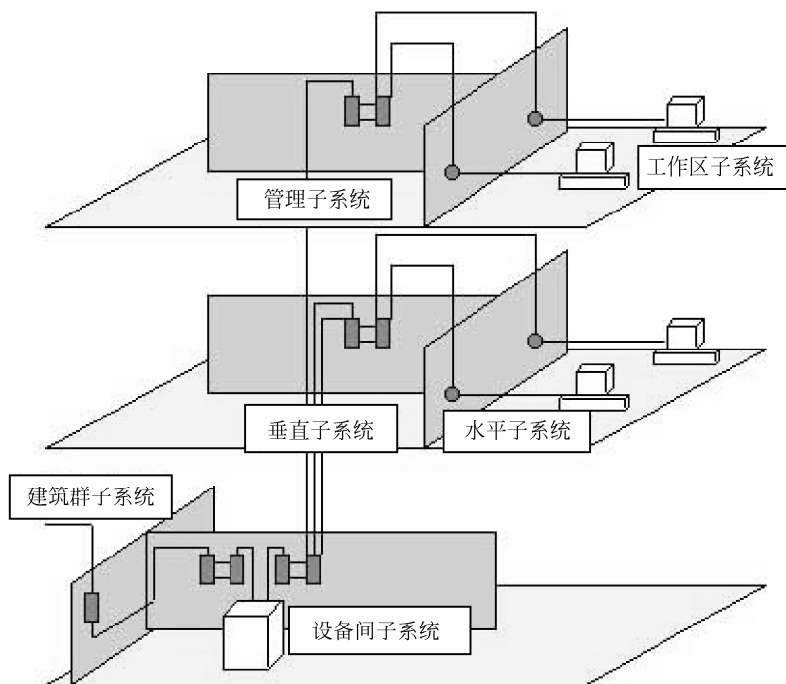
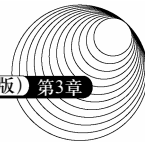


图 3-23 综合布线系统的构成

- 水平子系统 (Horizontal Subsystem)：由信息插座、配线电缆或光纤、配线设备和跳线等组成。国内称之为配线子系统。
- 垂直子系统 (Backbone Subsystem)：由配线设备、干线电缆或光纤、跳线等组成。国内称之为干线子系统。
- 工作区子系统 (Work Area Subsystem)：为需要设置终端设立的独立区域。



- 管理子系统 (Administration Subsystem)：是针对设备间、交接间、工作区的配线设备、缆线、信息插座等设施进行管理的系统。
- 设备间子系统 (Equipment room Subsystem)：是安装各种设备的场所，对综合布线而言，还包括安装的配线设备。
- 建筑群子系统 (Campus Subsystem)：由配线设备、建筑物之间的干线电缆或光纤、跳线等组成。

3.4.2 综合布线系统设计

1. 系统设计原则

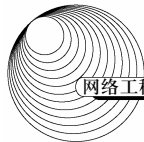
与其他系统设计一样，设计者首先要进行用户需求分析，然后根据需求分析进行方案设计。但需要指出的是，综合布线系统理论上可以容纳话音，包括电话、传真、音响（广播）；数据包括计算机信号、公共数据信息；图像包括各种电视信号、监视信号；控制包括温度、压力、流量、水位以及烟雾等各类控制信号。但实际工程中，至少在目前技术条件和工程实际需要中多为话音和数据，原因是多方面的，其中值得注意的是，话音的末端装置和计算机网络的终端用户装置往往是要变动的，有的是经常变动的，因此采用综合布线系统及其跳选功能，很容易在不改动原有敷线条件的情况下满足用户的需求。此外，本来可用同轴电缆可靠地传输电视信号，若改用综合布线，则要增设昂贵的转换器。对消防报警信号，用普通双绞线已达到要求，若改用综合布线，经过配线架再次终接，也无此必要。因此集成化的要求应视实际需要来定。

在进行综合布线系统设计时通常应遵循以下原则。

- (1) 采用模块化设计，易于在配线上扩充和重新组合。
- (2) 采用星型拓扑结构，使系统扩充和故障分析变得十分简易。
- (3) 应满足通信自动化与办公自动化的需要，即满足话音与数据网络的广泛要求。
- (4) 确保任何插座互连主网络，尽量提供多个冗余互连信息点插座。
- (5) 适应各种符合标准的品牌设备互连入网，满足当前和将来网络的要求。
- (6) 电缆的敷设与管理应符合综合布线系统设计要求。

2. 工作区子系统设计

根据综合布线设计规范的工程经验，并结合用户的实际建筑情况，除去走廊、过道等因素，考虑建筑面积的 70% 为实际办公面积，办公区每 $8\sim 10\text{m}^2$ 一个双孔信息出口，可配一部电话、一部计算机。信息插座通常可有如下所述三种安装形式。



(1) 信息插座安装于地面上。要求安装于地面的金属底盒应当是密封的,防水、防尘,并可带有升降的功能。此方法设计安装造价较高,并且由于事先无法预知工作人员的办公位置,也不知分隔板的确切位置,因此灵活性不是很好。

(2) 信息插座安装于分隔板上。此方法适于分隔板位置确定的情况下,安装造价较为便宜。

(3) 信息插座安装于墙上。此方法在分隔板位置未确定情况下,可沿大开间四周的墙面每隔一定距离均匀地安装 RJ45 埋入式插座。此方法和前两种方式相比,无论在系统造价、移动分隔板的方便性、整洁度方面,还是在安装和维护方面都是很好的。

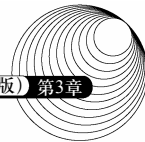
标准信息插座型号为 RJ45,采用 8 芯接线,全部按标准制造,符合 ISDN 标准。通常数据和话音均采用 MDVO(多媒介信息)模块式超五类信息插座。在 RJ45 插座内不仅可以插入数据通信用的 RJ45 接头,也可以插入电话机专用的 RJ12 插头。

3. 水平子系统设计

水平子系统将垂直子系统线路延伸到了用户工作区,由工作区的信息插座、信息插座至楼层配线设备(FD)的配线电缆或光纤、楼层配线设备和跳线等组成。该系统从各个子配架子系统出发连向各个工作区的信息插座。水平子系统要求走廊的吊顶上应安装有金属线槽,进入房间时,从线槽引出金属管以埋入方式由墙壁而下到各个信息点。通常水平子系统采用双绞线,在需要时也可采用光纤;根据整个综合布线系统的要求,应在交换间或设备间的配线设备上连接。如果采用双绞线,长度不应超过 90m。在保证链路性能的情况下,水平光纤距离可适当加长。信息插座采用 8 位模块式通用插座或光纤插座。配线设备交叉连接的跳线应选用综合布线专用的软跳线,在电话应用时也可选用双芯跳线。

双绞线作为水平子系统的主要组成部分,通常采用管线敷设,一般应使用 20 年左右,这也对双绞线的性能和质量提出了更高的要求。所以,应根据具体网络工程合理选择双绞线。比较好的办法是从实际应用出发,考虑未来发展的余地和投资费用,确保安装质量。从实际出发是指要考虑目前用户对网络应用的要求有多高,100M 以太网是否够用。因为网络的布线系统是一次性长期投资,考虑未来发展是指要考虑到网络的应用是否在一段时期内会有对千兆以太网或未来更高速的网络的需求。

就目前而言,进行一个新工程的永久性的综合布线,通常需要在超 5 类和 6 类之间选择。超 5 类系统可以支持千兆以太网的运行,而且不同厂商的超 5 类系统之间可以互用。6 类价格较之超 5 类更昂贵,但其带宽却由 200MHz 扩大到 250MHz,提高了 25%,显示了传输速率的增强。目前,6 类双绞线已经在少数工程中超前采用。需要注意的是,6 类系统是专用的,元件的指标仍在研究之中。各个厂商的元器件都有独特的设计和性能指标,互通的可能性很小。



4. 垂直子系统设计

垂直子系统主要用于连接各层配线室,并连接主配线室。垂直子系统要求建筑物竖井中应立有金属线槽,且每隔两米焊一根粗钢筋,以安装和固定垂直子系统的电缆。竖井中的线槽应和各层配线室之间有金属线槽连通。

垂直子系统实现计算机设备、程控交换机(PBX)、控制中心与各管理子系统间的连接,常用介质是大对数双绞线电缆、光纤。垂直干线部分提供了建筑物中主配线架与分配线架连接的路由,通常采用 IBDNPLUS 型、ATMM 或 Dflex 型 62.5/125 多模光纤来实现这种连接。

ATMM、Dflex 属于 3 类双绞线,通常被用做电话及广播信号等低速率的主干传输线缆。IBDN PLUS 型 62.5/125 多模光纤属于超 5 类的传输介质,其特性与水平子系统所用的同类线材的物理特性相同,被用做计算机、视频图像等高速数据应用的主干传输线缆。

多模光纤的优点为光耦合率高、纤芯对准要求相对较宽松。当弯曲半径大于其直径 20 倍时不影响信号的传输,是符合 IEEE 802.5 FDDI 和 EIA/TIA568 标准的光传输介质。用于计算机数据传输距离超过 100m 时的应用,其传输距离可达 2km。在保密性要求高的场合,建议也采用光纤传输。对于距离强电磁干扰源较近的情况,亦需要利用光纤的抗干扰性好的优点。充分考虑到投资的回报率和性能价格比,一般情况下,语音干缆采用符合 EIA/TIA568 标准的大对数 3 类双绞线,数据干缆采用 NTF-CMGR-06 多模光纤。

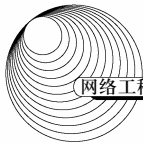
5. 管理子系统设计

管理子系统由交连、互连配线架组成,为连接其他子系统提供连接手段。交连和互连允许将通信线路定位或重定位到建筑物的不同部分,以便能更容易地管理通信线路,并且在移动终端设备时能方便地进行插拔。

分配线间是各管理子系统的安装场所,分配线间可位于大楼的某一层或以多层共用一个配线间的方式分布,用于将连接至工作区的水平线缆与自设备间引出的垂直线缆相连接。

对于信息点不是很多,使用功能又近似的楼层,为便于管理,可共用一个子配线间;对于信息点较多的楼层,应在该层设立配线室。配线室的位置可选在弱电竖井附近的房间内,用于安装配线架和安装计算机网络通信设备。

通常管理子系统使用墙装式光纤接续装置(光纤配线架),置于各层的配线间内。其上嵌 1 块 6ST 耦合器面板,ST 接头由陶瓷材料制成,最大信号衰减量小于 0.2dB。光纤接续装置将自设备间引出的光纤引入,通过光纤跳线与网络设备相连,由网络设备上的 UTP 端口经 UTP



跳线与配线架（置于 19 英寸机柜中）相连。

6. 设备间子系统设计

设备间子系统（主配线间）由设备间中的电缆、连接器和相关支撑硬件组成，它把公共系统设备的各种不同设备互连起来。该子系统将中继线交叉连接处和布线交叉处与公共系统设备（如 PBX）连接起来。

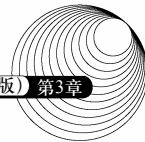
通常主配线架设置在程控机房内，用于垂直干缆和 PABX 的连接，建议采用 QCBIX 系列配线架，可充分满足话音通信的要求。通常计算机网络主配线架设在网管中心，使用光纤配线架，端接来自各分配线间的光纤，并通过光纤跳线和计算机网络中心交换机相连。光纤配线架采用 24/48 口配线箱，适用于光纤数量多密度大的场合，可直接安装在标准的 19 英寸机柜内，用于主干光纤和网络设备的连接，十分易于管理。

按照标准的设计要求，设备间，尤其是要集中放置设备的设备间，应尽量满足如下要求。

- (1) 将服务电梯安排在设备间附近，以便装运笨重的设备。
- (2) 室温应保持在 $18^{\circ}\text{C} \sim 27^{\circ}\text{C}$ 之间，相对湿度保持在 $30\% \sim 55\%$ 。
- (3) 保持室内无尘或少尘，通风良好，亮度至少达 30lx 。
- (4) 安装合适的消防系统（如采用湿型消防系统，不要把喷头直接对准电气设备）。使用防火门，使用至少能耐火 1 小时的防火墙和阻燃漆。
- (5) 提供合适的门锁，至少要有一扇窗口留做安全出口。
- (6) 尽量远离存放危险物品的场所和电磁干扰源（如发射机和电动机）。
- (7) 设备间的地板负重能力至少应为 500kg/m^2 。
- (8) 标准的天花板高度为 240cm ，门的大小至少为 $210\text{cm} \times 150\text{cm}$ ，向外开。
- (9) 在设备间尽量将设备机柜放在靠近竖井的位置，在柜子上方应装有通风口用于设备通风。
- (10) 在配线间内应至少留有两个专用的 $220\text{V}/10\text{A}$ 单相三极电源插座。如果需要在配线间内放置网络设备，则还应根据放置设备的供电需求配有另外的 $220\text{V}/10\text{A}$ 专用线路，此线路不应与其他大型设备并联，并且最好先连接到 UPS，以确保对设备的供电及电源的质量。

7. 建筑群子系统设计

建筑群子系统由连接各建筑物之间的综合布线缆线、建筑群配线设备（CD）和跳线等组成。建筑物之间的缆线宜采用地下管道或电缆沟的敷设方式。建筑物群干线电缆、光纤、公用网和专用网电缆、光纤（包括天线馈线）进入建筑物时，都应设置引入设备，并在适当位



置终端转换为室内电缆、光纤。引入设备还包括必要的保护装置。引入设备宜单独设置房间,如条件合适也可与 BD 或 CD 合设。建筑群和建筑物的干线电缆、主干光纤布线的交接不应多于两次。从楼层配线架(FD)到建筑群配线架(CD)之间只应通过一个建筑物配线架(BD)。

8. 管线设计

在综合布线系统中,管线设计通常有两种方案,一种是用于墙上型信息出口的,采用走吊顶的装配式槽形电缆桥架的方案,这种方式适用于大型建筑物,为水平子系统提供机械保护和支撑;另一种是用于地面型信息出口的地面线槽走线方式,这种方式适用于大开间的办公间,有大量地面型信息出口的情况。

1) 装配式槽形电缆桥架

装配式槽形电缆桥架是一种闭合式的金属托架,安装在吊顶内,从弱电井引向各个设有信息点的房间,再由预埋在墙内的不同规格的铁管将线路引到墙上的暗装铁盒内。

线槽的材料为冷轧合金板,表面可进行相应处理,如镀锌、喷塑、烤漆等。线槽可以根据情况选用不同的规格。根据本项目的需要,选择的是容积为 $50(B) \times 100(H) \text{mm}^2$ 、长度为 2m、重量为 3.67kg/m 的槽体配以上盖板宽为 100mm、长为 2m、重量为 2.20kg/m 的线槽和容积为 $50(B) \times 50(H) \text{mm}^2$ 、长度为 2m、重量为 1.91kg/m 的槽体配以上盖板宽为 50mm、长度为 2m 重量为 0.87kg/m 的两种规格的线槽。为保证线缆的转弯半径,线槽须配以相应规格的分支辅件,以提供线路路由的弯转自如。

同时为确保线路的安全,应使槽体有良好的接地端。金属线槽、金属软管、电缆桥架及各分配线箱均需整体连接,然后接地。如果不能确定信息出口的准确位置,拉线时可先将线缆盘在吊顶内的出线口,待具体位置确定后,再引到各信息出口。

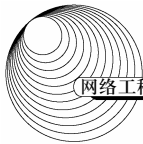
2) 地面线槽走线

地面线槽走线方式通常先在地面垫层中预埋金属线槽,主线槽从弱电井引出,沿走廊引向各方向,到达设有信息点的各房间时,再用支线槽引向房间内的各信息点出线口。强电线路可以与弱电线路平行配置,但需分隔于不同的线槽中。这样可以向每一个用户提供一个包括数据、语音、不间断电源、照明电源出口的集成面板,真正做到在一个清洁的环境下实现办公室自动化。

由于地面垫层中可能会有消防等其他系统的线路,所以必须与由建筑设计单位和建筑施工单位一起,综合各系统的实际情况,完成地面线槽路由部分的设计。另外,地面线槽也需整体连接,然后接地。

按照标准的线槽设计方法,应根据水平线的外径来确定线槽的横截面积,即:

$$\text{线槽的横截面积} = \text{水平线截面积之和} \times 3$$



9. 电气防护、接地及防火设计

综合布线系统应根据环境条件选用相应的缆线和配线设备,或采取防护措施,并应符合下列规定。

(1) 当综合布线区域内存在干扰或用户对电磁兼容性有较高要求时,宜采用屏蔽缆线和屏蔽配线设备进行布线,也可采用光纤系统。采用屏蔽布线系统时,所有屏蔽层应保持连续性。

(2) 综合布线系统采用屏蔽措施时,必须有良好的接地系统。保护地线的接地电阻值,单独设置接地体时,不应大于 4Ω ; 采用接地体时,不应大于 1Ω 。采用屏蔽布线系统时,屏蔽层的配线设备(FD或BD)端必须良好接地,用户(终端设备)端视具体情况接地,两端的接地应连接至同一接地体。若接地系统中存在两个不同的接地体,其接地电位差不应大于 $1V_{r.m.s.}$ 。每一楼层的配线柜都应采用适当截面的铜导线单独布线至接地体,也可采用竖井内集中用铜排或粗铜线引到接地体,导线或铜导体的截面应符合标准。接地导线应接成树状结构的接地网,避免构成直流环路。

(3) 当电缆从建筑物外面进入建筑物时,电缆的金属护套或光纤的金属件均应有良好的接地,同时要采用过压、过流保护措施,并符合相关规定。

(4) 根据建筑物的防火等级和对材料的耐火要求,综合布线应采取相应的措施。在易燃的区域和大楼竖井内布放电缆或光纤,应采用阻燃的电缆和光纤;在大型公共场所宜采用阻燃、低燃、低毒的电缆或光纤;相邻的设备间或交换间应采用阻燃型配线设备。

(5) 当综合布线路由上存在干扰源,且不能满足最小净距要求时,宜采用金属管线进行屏蔽。综合布线电缆与附近可能产生高频电磁干扰的电动机、电力变压器等电气设备之间应保持必要的间距。综合布线电缆与电力电缆的间距应符合表 3-2 的规定。墙上敷设的综合布线电缆、光纤及管线与其他管线的间距应符合表 3-3 的规定。

表 3-2 综合布线电缆与电力电缆的间距

类 别	与综合布线接近状况	最小净距 (mm)
380V 电力 电缆<2kVA	与缆线平行敷设	130
	有一方在接地的金属线槽或钢管中	70
	双方都在接地的金属线槽或钢管中	10
380V 电力 电缆 2-5kVA	与缆线平行敷设	300
	有一方在接地的金属线槽或钢管中	150
	双方都在接地的金属线槽或钢管中	80
380V 电力 电缆>5kVA	与缆线平行敷设	600
	有一方在接地的金属线槽或钢管中	300
	双方都在接地的金属线槽钢管中	150

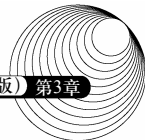


表 3-3 墙上敷设的综合布线电缆、光纤及管线与其他管线的间距

其 他 管 线	最小平行净距 (mm)	最小交叉净距 (mm)
	电缆、光纤或管线	电缆、光纤或管线
避雷引下线	1000	300
保护地线	50	20
给水管	150	20
压缩空气管	150	20
热力管 (不包封)	500	500
热力管 (包封)	300	300
煤气管	300	20

3.4.3 综合布线系统的性能指标及测试

综合布线作为网络中最基本、最重要的组成部分,它是连接每一台服务器和工作站的纽带,作为传输高速数据的介质,综合布线系统对线缆的要求较严格,一旦线缆产生故障,严重时可导致整个网络系统瘫痪。一个布线系统的传输性能是由多种因素决定的,包括线缆特性、连接硬件、跳线、整体回路连接数目以及设计和安装质量。即使线缆和连接硬件都符合国际标准,由于在布线系统的设计和安装过程中加入了许多人为因素,所以必须对整个布线系统进行全面测试,以证明布线系统的安装是合格的。

1. 双绞线系统的测试元素及标准

通常,双绞线系统的测试指标主要集中在链路传输的最大衰减值和近端串音衰减等参数上。链路传输的最大衰减值是由于集肤效应、绝缘损耗、阻抗不匹配、连接电阻等因素,造成信号沿链路传输损失的能量。电磁波从一个传输回路(主串回路)串入另一个传输回路(被串回路)的现象称为串音,能量从主串回路串入回路时的衰减程度称为串音衰减。在 UTP 布线系统中,近端串音为主要的影响因素。

下面给出双绞线系统的几个主要测试元素及标准。需要指出的是,表中数值为通道回路总长度为 100m 以内、基本回路总长度为 94m 以内、测试温度为 20℃ 下的标准值。

1) 链路传输的最大衰减限值

综合布线系统链路传输的最大衰减限值,包括配线电缆和两端的连接硬件、跳线在内,应符合表 3-4 的规定。

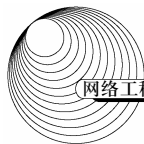


表 3-4 链路传输的最大衰减限值

频率 (MHz)	最大衰减值 (dB)			
	A 级	B 级	C 级	D 级
0.1	16	5.5	—	—
1.0	—	5.8	3.7	2.5
4.0	—	—	6.6	4.8
10.0	—	—	10.7	7.5
16.0	—	—	14.0	9.4
20.0	—	—	—	10.5
31.25	—	—	—	13.1
62.5	—	—	—	18.4
100.0	—	—	—	23.2

2) 近端串音 (NEXT) 衰减限值

综合布线系统任意两线之间的近端串音衰减限值, 包括配线电缆和两端的连接硬件、跳线、设备和工作区连接电缆在内 (但不包括设备连接器), 应符合表 3-5 的规定。

表 3-5 线对间最小近端串音衰减限值

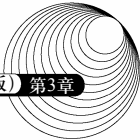
频率 (MHz)	最大衰减值 (dB)			
	A 级	B 级	C 级	D 级
0.1	27	40	—	—
1.0	—	25	39	54
4.0	—	—	29	45
10.0	—	—	23	39
16.0	—	—	19	36
20.0	—	—	—	35
31.25	—	—	—	32
62.5	—	—	—	27
100.0	—	—	—	24

3) 回波损耗限值

综合布线系统中任一电缆接口处的回波损耗限值应符合表 3-6 的规定。

表 3-6 电缆接口处最小回波损耗限值

频率 (MHz)	最小回波损耗值	
	C 级	D 级
$10 \leq f < 16$	15	15
$16 \leq f < 20$		15
$20 \leq f < 100$		10



2. 光纤布线系统的测试元素及标准

在光纤系统的实施过程中涉及光纤的铺设，光纤的弯曲半径，光纤的熔接、跳线，设计方法及物理布线结构的不同会导致两网络设备间的光纤路径上光信号的传输衰减有很大不同。虽然光纤的种类较多，但光纤及其传输系统的基本测试方法大体相同，所使用的测试仪器也基本相同。对磨接后的光纤或光纤传输系统，必须进行光纤特性测试，使之符合光纤传输通道测试标准。基本的测试内容如下。

1) 波长窗口参数

综合布线系统光纤波长窗口的各项参数应符合表 3-7 的规定。

表 3-7 光纤波长窗口参数

光 纤 模 式	波 长 下 限	波 长 上 限	基准试验波长	谱线最大宽度
多模	790nm	910 nm	850 nm	50 nm
多模	1285 nm	1330 nm	1300 nm	150 nm
单模	1288 nm	1339 nm	1310 nm	10 nm
单模	1525 nm	1575 nm	1550 nm	10 nm

2) 光纤布线链路的最大衰减限值

综合布线系统的光纤布线链路的衰减限值应符合表 3-8 的规定。

表 3-8 光纤布线链路的最大衰减限值

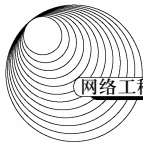
应 用 类 别	链路长度 (m)	多模衰减值 (dB)		单模衰减值 (dB)	
		850 (nm)	1300 (nm)	1310 (nm)	1550 (nm)
水平子系统	100	2.5	2.2	2.2	2.2
垂直子系统	500	3.9	2.6	2.7	2.7
建筑群子系统	1500	7.4	3.6	3.6	3.6

3) 光回波损耗限值

综合布线系统光纤布线链路任一接口的光回波损耗限值应符合表 3-9 的规定。

表 3-9 最小光回波损耗限值

光纤模式、标称波长 (nm)	最小的光回波损耗限值 (dB)
多模 850	20
多模 1300	20
单模 1310	26
单模 1550	26



3. 测试环境

为了保证布线系统测试数据准确可靠,对测试环境有着严格的规定。

1) 测试条件

综合布线最小模式带宽测试现场应无产生严重电火花的电焊、电钻和产生强磁干扰的设备作业,被测综合布线系统必须是无源网络、无源通信设备。

2) 测试温度

综合布线测试现场温度在 $20^{\circ}\text{C}\sim 30^{\circ}\text{C}$ 之间,湿度宜在 $30\%\sim 80\%$ 之间,由于衰减指标的测试受测试环境温度影响较大,当测试环境温度超出上述范围时,需要按照有关规定对测试标准和测试数据进行修正。

3) 测试仪表

按时域原理设计的测试仪均可用于综合布线现场测试,但测试仪的测量扫描步长要满足近端串扰指标测量精度的基本保证,能够在 $0\sim 250\text{MHz}$ 频率范围内提供各测试参数的标称值和阈值曲线,具有自动、连续、单项选择测试的功能。每测试一条链路,时间不应大于 25s ,且每条链路应具有一定的故障定位诊断能力。

4. 测试流程

在开始测试之前,应该认真了解布线系统的特点、用途以及信息点的分布情况,确定测试标准。选定测试仪后按以下程序进行测试。

- (1) 测试仪测试前自检,确认仪表是正常的。
- (2) 选择测试了解方式。
- (3) 选择设置线缆类型及测试标准。
- (4) NVP 值核准,核准 NVP 使用缆长不短于 15m 。
- (5) 设置测试环境湿度。
- (6) 根据要求选择“自动测试”或“单项测试”。
- (7) 测试后存储数据并打印。
- (8) 发生问题修复后复测。
- (9) 测试中出现“失败”查找故障。