

第 1 章 网络安全概述

进入 21 世纪,随着计算机网络的飞速发展,社会的经济、文化和军事越来越多地依赖于计算机网络。然而,计算机网络在给人类生活带来极大便利的同时,网络安全也面临着极大的挑战。敏感信息的泄露、信息的篡改、数据的破坏和计算机病毒的发作都会给社会生活带来难以估量的损失,解决开放式网络环境下的网络安全问题刻不容缓。网络安全作为一门综合的、交叉的学科,涉及数学、物理、管理、信息技术和计算机技术等多个学科领域。

1.1 网络安全的基础知识

“某银行的客户数据被黑客窃取”、“QQ 号码被别人盗用”等是生活中遇到的网络信息安全事件。而网络安全,从本质上讲就是网络上的信息安全,凡涉及网络信息的保密性、完整性、可用性的网络技术和理论都是网络安全研究的领域,是信息安全在当前网络环境下最重要的研究领域。

1.1.1 网络安全的定义

国际标准化组织 ISO7498-2 标准中对安全的定义是:安全就是最大限度地减少数据和资源被攻击的可能性。由于 Internet 的开放性和超越组织与国界等特点,使它在安全性的保障上存在一系列的隐患。那么,什么是网络安全呢?

目前,网络安全并没有公认和统一的定义,现在采用比较多的定义是:网络安全是指利用网络管理控制和技术措施,保证在一个网络环境里信息数据的机密性、完整性、可使用性、真实性和可控性受到保护。

从内容上看,网络安全包括以下 4 个方面:

(1)物理安全:计算机机房的物理条件、环境和设施的安全,计算机硬件、配套设备及网络传输线路的安全。

(2)数据安全:保护数据不被非法存取,确保数据的完整性、一致性和机密性等。

(3)软件安全:保护网络系统不被非法侵入,系统软件与应用软件不被非法复制和篡改,不受病毒的侵害。

(4)安全管理:运行时对突发事件的安全处理等,包括采用计算机安全技术及规范安全管理制度,开展安全审计和进行分析等措施。

1.1.2 网络安全的需求

进入 21 世纪,随着 Internet 的发展,以网络技术及服务为代表的第二次“信息革命”浪潮席卷全球,迅速渗透到政府、企业、经济以及与日常生活息息相关的各个领域。网络上传

息的广泛传播给用户带来了极大的方便,但同时,也给用户带来了安全问题。

过去的网络大多是封闭式的,因而简单的安全性设备就足以确保其安全。然而,当今的网络已发生了变化,确保网络的安全性已成为更加复杂而且必需的任务。用户每次连接到网络上,原有的安全状况就会发生变化。简单的安全措施和设备已对现有的网络安全问题无能为力了。所以,很多网络频繁地成为网络犯罪的牺牲品。

互联网犹如一柄“双刃剑”,为人们工作带来便利的同时,伴随着日趋严重的网络入侵安全问题。一些公用站点,既要访问 Internet 的共享信息资源,又要使 Intranet 的一部分信息对外提供服务,资源共享的同时也带来了安全问题。政府部门内部网关系到很多政府机密、政府形象等敏感信息,网络的安全性更为重要。因此,保护政府内部网络的信息安全,防范来自外部网络的黑客和非法入侵者的攻击,建立强健的网络信息安全防范系统,在某种程度上决定着政府部门信息化建设的成败。

另外,随着企业网上业务的不断扩大和电子商务的发展,对网络的安全服务也提出了新的要求。严防黑客入侵、切实保障网络交易的安全,不仅关系到个人的资金安全和企业的数据安全,还关系到国家的经济安全、国家经济秩序的稳定,网络安全问题已成为亟待解决的问题。

对于个人用户而言,涉及个人隐私的信息在网络上传输时应受到保护,避免其他人利用窃听、冒充、篡改和抵赖等手段侵犯其隐私,同时也避免其他用户的非授权访问和破坏。

由此可见,目前整个 Internet 的安全问题不容乐观。而且网络安全的内涵也发生了根本的变化,它不仅从个别的防卫变成了一种非常普遍的防范,而且还从一种专门的领域变成了无处不在,有网络的地方就有网络安全问题。

1.1.3 网络安全的目标

从技术角度来说,网络信息安全的目标主要表现在以下 6 个方面。

1. 保密性

保密性是网络信息不被泄露给非授权的用户、实体或供其利用的特性。即防止信息泄露给未授权用户或实体,信息只为授权用户使用的特性。破坏信息的保密性是对信息发动攻击的最终目的。

2. 完整性

完整性是网络信息未经授权不能进行改变的特性。即网络信息在存储或传输过程中保持不被偶然或蓄意地删除、修改、伪造、乱序、重放、插入等破坏的特性。完整性是一种面向信息的安全性,它要求保持信息的原样,即信息的正确生成以及正确存储和传输。

保证完整性的目的就是保证计算机系统上的数据和信息处于一种完整和未受损害的状态,这就是说,数据不会因有意或无意的事件而被改变或丢失,信息完整性的丧失直接影响到数据的可用性。完整性与保密性不同,保密性要求信息不被泄露给未授权的用户或实体,而完整性则要求信息不受各种原因的破坏。

3. 可靠性

可靠性是网络信息系统能够在规定条件下和规定的时间内完成规定功能的特性。它是网络系统安全的最基本要求之一,是所有网络信息系统建设和运行的目标。网络系统的可靠性主要体现在以下 3 个方面:

(1)抗毁性是指系统在人为破坏下的可靠性。例如,部分线路或结点失效后,系统是否

仍然能够提供一定程度的服务。增强抗毁性可以有效地避免因各种灾害(战争、地震等)造成的大面积网络瘫痪事件。

(2)生存性是在随机破坏下系统的可靠性,主要反映随机性破坏和网络拓扑结构改变对系统可靠性的影响。

(3)有效性主要反映在网络信息系统部件失效的情况下满足业务性能要求的程度。例如,网络部件失效虽然没有引起连接性故障,但是却造成质量指标下降、平均延时增加、线路阻塞等现象的发生。

4. 可用性

可用性是网络信息可被授权用户或实体访问并按需求使用的特性。即网络信息服务在需要时,允许授权用户或实体使用的特性,或者是网络部分受损或需要降级使用时,仍能为授权用户或实体提供有效服务的特性。可用性是网络信息系统面向用户的安全性能。

网络信息系统最基本的功能是向用户提供服务,而用户的需求是随机的、多方面的,有时还有时间要求。可用性一般用系统正常使用时间和整个工作时间之比来度量。

5. 不可抵赖性

不可抵赖性也称不可否认性。在网络信息系统的信息交互过程中,确信参与者的真实同一性。即所有参与者都不可能否认或抵赖曾经完成的操作和承诺。利用信息源证据可以防止发信方不真实地否认已发送信息,利用递交接收证据可以防止收信方事后否认已经接收信息。

6. 可控性

可控性是对网络信息的传播及内容具有控制能力的特性。

网络信息安全与保密的核心是通过计算机、网络、密码技术和其他安全技术保护在公用网络信息系统中传输、交换和存储消息的保密性、完整性、可靠性、可用性、不可抵赖性和可控性等。

1.2 威胁网络安全的因素

随着 Internet 的发展,人们的生活发生了许多变化。银行存款的通存通兑、网上汇款、上网冲浪、网上购物等无不给人们的生活带来便利。但是,在享受计算机网络带来方便和快捷的同时,网络病毒、网络攻击以及网络犯罪也达到空前猖獗的程度,网络本身所具有的脆弱性也日益显现出来。

1.2.1 网络安全的主要威胁

目前,计算机网络所面临的安全威胁很多,主要有以下几个方面。

1. 对加密算法的攻击

数据加密技术是最基本的网络安全技术,被誉为信息安全的核心。最初主要用于保证数据在存储和传输过程中的保密性。网络中使用的加密算法,从加密的种类上来分主要包括对称加密和非对称加密。公钥密码(非对称加密)体系能适应网络的开放性要求,密钥管理简单,并且可方便地实现数字签名和身份认证等功能,是目前电子商务等技术的核心基础。

密码分析还原技术主要分为密码还原技术和密码猜测技术。对于迭代分组加密算法可以使用差分分析、线性分析和穷举法等攻击方法进行有效破解。1997年1月28日,美国RSA数据安全公司在RSA安全年会上发起了一项“秘密密钥挑战”竞赛,美国科罗拉多州的程序员 Venser 用了96天的时间,在Internet上数万名志愿者的协同工作下,成功地找到了DES的密钥,这一事件表明,破译DES已成为可能。目前,已经有人宣称破解了512位的RSA算法和97位的椭圆曲线公钥算法。

2. 协议漏洞渗透

网络中包含着种类繁多但层次清晰的网络协议规范。这些协议规范是网络运行的基本准则,也是构建在其上的各种应用和服务运行的基础。

在一些已经有一定历史的网络中,网络的核心协议(TCP/IP协议)对安全的考虑有着先天的不足,部分网络协议具有严重的安全漏洞。如传统的以太网主要是以共享的方式来完成数据分组的传送。发往目的结点的分组数据实际上被发送给了其所属网段的每一个结点。目的结点接收这些分组,并与其他结点共享传送带宽,只要可以作为目标网络环境的一个结点,就可以接收到目标网络中流动的所有数据信息。通过对网络标准协议的分析,黑客可以从中总结出针对协议的攻击过程,利用协议的漏洞实现对目标网络的攻击,以会话侦听与劫持技术和地址欺骗技术应用较多。

同样,应用层协议也有漏洞。攻击者可以利用简单邮件传输协议SMTP(simple mail transfer protocol)没有提供认证的弱点,伪造邮件冒充某一授权人或对邮件服务器进行电子邮件轰炸,即向目标发送大量邮件使其崩溃。

3. 系统漏洞

任何应用程序都不可避免地存在着一些逻辑漏洞,这在IT行业中已经形成了共识。操作系统也不例外,几乎每天都有人宣布发现了某个操作系统的安全漏洞。而这些安全漏洞也就成为了入侵者的攻击对象。通过对这些安全漏洞的分析,确认漏洞的引发方式以及引发后对系统造成的影响,攻击者可以使用合适的攻击程序引发漏洞的启动。从漏洞类型上划分,主要包括服务流程漏洞和边界条件漏洞。

(1)服务流程漏洞指服务程序在运行处理过程中,由于流程次序的颠倒或对意外条件的处理的随意性,造成用户有可能通过特殊类型的访问绕过安全控制部分或使服务进入异常的运行状态。例如,在对输入不作严格限制的CGI程序中,用户可以输入含有运行代码的请求。如果服务器没有对输入进行合法性的处理,CGI程序就会在执行的过程中启动用户写入的运行代码,造成系统信息的泄露或破坏。

(2)边界条件漏洞则主要针对服务程序中存在的边界处理不严谨的情况。在对服务程序的开发过程中,很多边界条件尤其是对输入信息的合法性处理往往很难作到周全,在正常情况下,对边界条件考虑的不严密并不会造成明显的错误,这种不严谨的处理却会带来严重的安全隐患。在边界漏洞中,以内存溢出错误最为普遍,影响也最为严重。有很多攻击都是利用超长的数据填满数据区并造成溢出错误,利用这种溢出在没有写权限的内存中写入非法数据。

4. 拒绝服务攻击

拒绝服务DoS(denial of service)攻击主要利用网络协议的一些薄弱环节,通过发送大量无效请求数据包造成服务器进程无法短期释放,耗尽系统资源,使得服务器无法对正常的请求进行响应,造成服务的瘫痪。DoS攻击最主要的目的是造成被攻击服务器资源耗尽或

系统崩溃而无法提供服务。凡是导致合法用户不能访问正常网络服务的行为都是拒绝服务攻击。也就是说拒绝服务攻击的目的就是要阻止合法用户对正常网络资源的访问。这样的入侵对于服务器来说可能并不会造成损害,但可以造成人们对被攻击服务器所提供服务的信任度下降,影响公司的声誉以及用户对网络服务的使用。

随着计算机与网络技术的发展,计算机的处理能力迅速增强,内存大大增加,同时也出现了千兆级别的网络,这使得 DoS 攻击的困难程度加大,分布式的拒绝服务 DDoS(distribution denial of service)攻击应运而生。DDoS 是拒绝服务群起进攻的方式。DDoS 最早可追溯到 1996 年初,随着流氓软件、病毒、木马大量充斥网络,获得傀儡主机变得简单而容易,攻击者可以控制众多的傀儡主机对远程计算机发起拒绝服务攻击。

DDoS 攻击与传统的拒绝服务攻击一样,只不过进攻源不只一个。黑客首先进入成百上千没有安全防护系统的计算机,在计算机内安装攻击程序,使之成为傀儡主机,之后控制众多傀儡计算机同时向目标发起进攻。目标机即刻受到来自多个地方的攻击,使传统的防范措施失去作用,最终死机。在传统方式的拒绝服务攻击中,作为受害者的计算机可能会察觉攻击源,并关闭这些连接;在分布式拒绝服务攻击中,计算机应关闭除它信任的连接之外的所有连接,这在公共 Internet 站点上根本无法实现。因此,迄今为止,对分布式拒绝服务攻击还没有通用的防护手段。

5. 计算机病毒和恶意代码的威胁

计算机病毒(computer virus)在《中华人民共和国计算机信息系统安全保护条例》中被明确定义为“编制或者在计算机程序中插入的破坏计算机功能或者破坏数据,影响计算机使用并且能够自我复制的一组计算机指令或者程序代码”。根据现有的病毒资料,计算机病毒具有寄生性、传染性、潜伏性、隐蔽性、破坏性、可触发性等特点。

从早期具有代表性的“小球”和“石头”病毒到“幽灵”病毒,病毒种类一直在不断地增长。根据瑞星公司发布的 2009 年互联网安全报告的统计,截至 2009 年 10 月底,该公司截获病毒样本 15 306 914 个,比 2008 年同期增加 64.5%。2004 年 5 月 1 日,“震荡波”开始在互联网上传播,到 5 月 3 日全球约有 1 800 万台计算机报告感染了这一病毒,计算机反复出现自动关机的提示框,然后数秒钟后机器反复重启,网络资源被程序消耗,最终导致网络瘫痪。近年来发生在我国的病毒案件也有很多,2006 年年底爆发的“熊猫烧香”病毒的影响最大。“熊猫烧香”对被感染系统中的 *.exe、*.com 等文件添加病毒网址,用户一旦打开这些文件,IE 就会自动链接到指定的病毒网址中下载病毒。它能够终止大量的杀毒软件和防火墙程序的运行,让杀毒软件和防火墙程序起不了任何作用。

恶意代码编写者一般利用软件漏洞、用户本身或者两者的混合来传播恶意代码。恶意代码中有些表现了高度的心理操纵能力,如 Love Letter;有些利用商品软件缺陷的恶意代码,如 Code Red、Nimda 等;有些恶意代码是自启动的蠕虫和嵌入脚本,本身就是软件;一些像特洛伊木马、电子邮件蠕虫等恶意代码,利用受害者的心理操纵他们执行不安全的代码;还有一些是哄骗用户关闭保护措施来安装恶意代码。宏病毒 Concept,用了 3 个月的时间才流行开来,Love Letter 用了大约一天,而 Code Red 用了大约 90 分钟,Nimda 用了不到 30 分钟。恶意代码的传播方式在迅速地演化,从引导区传播到某种类型文件传播、宏病毒传播、邮件传播到网络传播,发作和流行的时间也越来越短。

计算机病毒与恶意代码随着网络的普及,增长速度越来越快,变种越来越多,发生的频率也呈逐年上升的趋势,对网络安全造成的威胁也越来越大,带来了巨大的安全隐患。

1.2.2 威胁网络安全的因素

威胁网络安全的因素归纳起来主要包括自然因素、人为因素和系统本身因素。

1. 自然因素

自然因素是指自然环境对计算机网络设备与设施的影响,一般来自于各种自然灾害、恶劣的场地环境、电磁干扰等。这些事件一般具有突发性和不可抗拒的特点,会直接威胁网络安全,影响信息的存储媒体,造成的影响范围通常较大,损坏程度也较严重。

计算机硬件系统对运行环境有特定的要求。计算机硬件系统是一种精密仪器型的电子设备,其中有的部位电磁信号强度很弱,抗电磁干扰能力很差,当电场强度过大时,小信号电路就不易正常工作;当磁感应强度较大时,磁记录设备的信息难免不被破坏;辐射会使大规模集成电路中的某些部位产生较大的光电流而损伤;元部件间的接插方式易受灰尘、湿度、有害气体的锈蚀,或因振动、冲击而松动,影响牢靠的电气连接;温度、湿度的偏高以及灰尘会使元部件的电特性变差甚至不能工作;密封防尘、防湿与改善机内温度难以两全;不少元部件和系统要求供电电源的电压、频率稳定,供电不能突然中断,否则,元部件会突然损坏,或者系统遭受无法挽回的损失。例如,2008 年年初我国南方遇到的雪灾导致通信中断,造成很大损失。

2. 人为因素

人为因素主要包括两类:非恶意威胁和恶意威胁。

1) 非恶意威胁

非恶意威胁主要来自一些人为的误操作或一些无意的行为。例如,文件的误删除、输入错误的数据、操作员安全配置不当、用户口令选择不慎、用户将自己的账号随意转借他人或与别人共享等,这些无意的行为都可能给信息的安全带来威胁。

2) 恶意威胁

恶意威胁是计算机网络系统面临的最大威胁。恶意威胁网络安全的人群主要有 3 种:故意破坏者、不遵守规则者和刺探秘密者。故意破坏者企图通过各种手段去破坏网络资源与信息,例如,涂抹别人的主页、修改系统配置、造成系统瘫痪;不遵守规则者企图访问不允许访问的系统,他可能仅仅是到网络中看看,找些资料,也可能想盗用别人的计算机资源;刺探秘密者的企图非常明确,即通过非法手段侵入他人系统,以窃取商业秘密与个人资料。这 3 种人多采用主动攻击的手段。

攻击者对网络进行攻击的手段从不同的角度可分成不同的种类。从主动和被动的角度考虑可分为主动攻击和被动攻击。

(1) 主动攻击是一种破坏力极大的攻击手段。它是指避开或打破安全防护,引入恶意代码,破坏数据和系统的完整性。例如,篡改网络中的信息(修改数据内容,删除其中的部分内容,用一条虚假的数据替代原始数据,将某些额外数据插入其中等);否认自己曾经发布过的信息、伪造对方来信、修改来信等;制造和传播计算机病毒,这是一种破坏力极大的攻击手段。

(2) 被动攻击是指监视公共媒体上信息的传送。例如,采用口令嗅探、窃听等手段获取正在传输的信息,然后对获取的报文内容和通信流进行分析。被动攻击本身具有隐蔽性,并不涉及数据的任何改变,所以对被动攻击的检测很困难。这种攻击表面上不对系统造成什么破坏,但实际上是为进行进一步的破坏做前期准备工作。

从攻击来源的角度考虑,可分为外部攻击和内部攻击。

由网络外部因素引起的安全问题都是外部攻击,如外部入侵者对网络的威胁等。而由网络内部因素引起的安全问题就是内部威胁,如网络系统内部入侵者对网络进行有意或无意的攻击,系统本身的问题等。实际上,内部人员往往是利用偶然发现的系统弱点或预谋突破网络系统安全进行攻击。由于内部人员更了解网络结构,因此他们的非法行为对网络威胁更大。据统计资料显示,来自内部人员的攻击在所有受到的攻击事件中占到 80% 以上。

3. 系统本身因素

网络面临的威胁并不仅仅来自于自然或人为,很多时候还来自于系统本身。如系统本身的电磁辐射或硬件故障、软件“后门”、软件自身的漏洞等。

(1) 计算机硬件系统的故障指计算机的硬件物理损坏或机械故障。计算机故障发生的原因有很多,如常见的集成电路本身引起的故障、静电感应引起的故障以及环境问题引起的故障等。

(2) 软件的“后门”常常是软件公司程序设计人员为了自身方便而在开发时为调试程序预留的。这样就为软件调试、进一步开发或远程维护提供了方便,但同时也为非法入侵提供了通道。这些“后门”一般不为外人所知,“后门”一旦被打开,其造成的后果将不堪设想。

(3) 软件也不可能是百分之百的正确,一定会存在漏洞或缺陷。这些缺陷或漏洞就成为攻击者攻击的首选目标。例如,常用的操作系统,无论是 Windows 还是 UNIX 几乎都存在或多或少的安全漏洞;众多的服务器、浏览器、一些桌面软件等都被发现过存在安全隐患。任何一个软件系统都可能会因为程序员的一个疏忽、设计中的一个缺陷等而存在漏洞,从而威胁到网络的安全运行。

除上述 3 个主要因素外,还有管理不善、规章制度不健全,或有章不循、安全管理水平低、人员技术素质差和操作失误等,都会对网络安全造成威胁。

1.3 网络安全防范体系

为了能够准确了解用户的安全需求,选择各种安全产品和策略,必须建立一些系统的方法来进行网络安全防范。网络安全防范体系的科学性、可行性是其可顺利实施的保障。

1.3.1 网络安全防范的层次

根据网络的现状和网络的结构,可将安全防范的层次划分为物理安全、系统安全、网络安全、应用安全和安全管理,如图 1-1 所示。其中不同的层次反映了不同的安全问题。

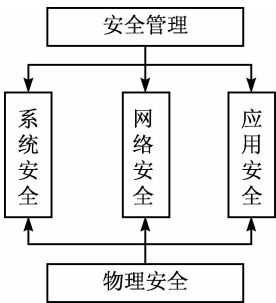


图 1-1 安全防范层次

1. 物理安全

保证计算机网络各种设备的物理安全是整个网络安全的前提。物理安全是指保护计算机通信线路、网络设备、设施等免遭地震、水灾、火灾等环境事故带来的破坏以及减少由于人为操作失误或错误和各种计算机犯罪等行为导致的破坏。此外,除了在网络规划时要满足场地、环境等要求之外,还要防止信息在空间的扩散。

物理层的安全主要体现在通信线路的可靠性(线路备份、网管软件、传输介质)、软硬件设备的安全性(替换设备、拆卸设备、增加设备)、设备的备份、防灾害能力、防干扰能力和设备的运行环境(温度、湿度、烟尘)、不间断电源保障等。

2. 系统安全

作为计算机系统安全功能的管理者和执行者,操作系统负责对计算机系统的各种资源、操作、运算、用户进行管理和控制。该层次的安全问题来自网络内使用的操作系统的安全,如 Windows NT、Windows XP、UNIX 等操作系统的安全。主要表现在:一是操作系统本身的缺陷带来的不安全因素,包括身份认证、访问控制、系统漏洞等;二是对操作系统的安全配置问题;三是病毒对操作系统的威胁。针对以上不安全因素,系统管理员应该利用系统风险评估工具,找出不应该安装或应该缩小权限的程序,并使用实时入侵检测系统对用户的活动进行跟踪,当检测到威胁系统安全的结果或行为时,及时采取有效的防范措施加以阻止。

3. 网络安全

随着网络上信息传输的日益频繁,网络的传输量大大增加,而且越来越多的用户选择网络作为机密信息的传输工具。为了应对可能出现的安全问题,网络层提供了访问控制、安全传输和连接服务,在实际应用中主要采用防火墙技术和 VPN 技术。通过防火墙技术控制具有特定 IP 地址的用户进入网络,形成内外网之间的一道屏障,将不安全的用户和数据阻挡在网络之外。VPN 技术能够在嘈杂的外部信道上为用户建立一条虚拟的专用链路,避免高昂的租用专线的费用,为用户实现安全的连接服务。

4. 应用安全

应用层的安全主要是针对用户身份进行认证,有效地解决诸如电子邮件、Web 服务、DNS 等特定应用的安全问题。通过提供包括身份认证、不可否认、数据保密、数据完整性检查乃至访问控制等功能,使合法用户能够对特定的数据进行可控的操作。

5. 安全管理

安全管理包括安全技术和设备的管理、安全管理制度、部门与人员的组织规则等,管理的制度化在很大程度上影响着整个网络的安全。严格的安全管理制度、明确的部门安全职责划分、合理的人员角色配置都可以在很大程度上减少其他层次的安全漏洞。

1.3.2 网络安全体系结构

为了保证以上 5 个层次的安全,同时也为了适应网络技术的发展,国际标准化组织(ISO)根据开放系统互连(OSI)参考模型制定了一个网络安全体系结构,在这个体系结构中规定了 5 种安全服务和 9 种安全机制,来解决网络中的信息安全与保密问题。

1. 安全服务

网络安全体系结构中规定的 5 种安全服务包括认证服务、访问控制服务、数据保密性服务、数据完整性服务以及防抵赖服务。

1) 认证服务

对象认证服务是防止主动攻击的重要措施,这种安全服务提供对通信中的对等实体和数据来源的鉴别,它对开放系统环境中的各种信息安全有着重要的作用。认证就是识别和证实。识别是辨别一个对象的身份,证实是证明该对象的身份就是其声明的身份。OSI 参考模型可提供对等实体认证的安全服务和数据源认证的安全服务。

2) 访问控制服务

这种服务可以防止未经授权的用户非法使用系统资源。访问控制可以分为自主访问控制和强制访问控制两类。访问控制服务主要位于应用层、传输层和网络层。它可以放在通信源、通信目标或两者之间的某一部分。这种服务不仅可以提供给单个用户,也可以提供给封闭的用户组中的所有用户。

3) 数据保密性服务

数据保密性服务是针对信息泄露、窃听等威胁的防御措施,目的是保护网络中各系统之间交换的数据,防止因数据被截获而造成的泄密。这种服务又分为信息保密、选择段保密和业务流保密。信息保密是保护通信系统中的信息或网络数据库的数据;选择段保密是保护信息中被选择的部分数据段;业务流保密是防止攻击者通过观察业务流,如信源、信宿、传送时间、频率和路由等来得到敏感的信息。

4) 数据完整性服务

这种服务用来防止非法用户的主动攻击,以保证数据接收方收到的信息与发送方发送的信息一致。数据完整性服务又分为连接完整性服务、选择段有连接完整性服务、无连接完整性服务以及选择段无连接完整性服务。

连接完整性服务为一个连接上的所有信息提供完整性,具体方法是探测是否对信息进行了非法的插入、删除或篡改;选择段有连接完整性服务为一个连接所传送信息中选择的信息段提供完整性,方法是探测对选择的信息段是否进行了非法的插入、删除或篡改。无连接完整性服务为无连接的各个信息提供完整性,方法是鉴别所收到的信息是否被非法篡改过;选择段无连接完整性服务为在各个无连接的信息中所选择的信息段提供完整性,方法是鉴别所选择的信息段是否被非法篡改过。

5) 防抵赖服务

防抵赖服务又称不可否认性服务,主要是用来防止发送数据方发送数据后否认自己发送过数据,或接收数据方收到数据后否认自己收到过数据。这种服务又可细分为不得否认发送、不得否认接收和依靠第三方。不得否认发送服务向数据的接收者提供数据来源的证据,从而可防止发送者否认发送过这些数据或否认这些数据的内容;不得否认接收服务向数据的发送者提供数据交付证据,从而防止了数据接收者事后否认收到过这些数据或否认它的内容;依靠第三方服务是在通信双方互不信任,但对第三方(公证方)则绝对信任的情况下,依靠第三方来证实已发生的操作。

2. 安全机制

安全服务依赖于安全机制的支持。安全机制是利用密码算法对重要而敏感的信息进行处理。OSI 参考模型提供的安全机制主要有加密机制、数字签名机制、访问控制机制、数据完整性机制、认证机制、业务流量填充机制、路由控制机制、公证机制和安全审计跟踪机制共 9 种。

1) 加密机制

加密机制主要用来加密存储数据,是保护数据最常用的方法。加密机制既可以单独使用,也可以与其他机制结合起来使用。加密机制通过密钥的管理机制来实现,需要在相应的加密体制下完成。

2) 数字签名机制

数据加密是保护数据最常用的方法,但这种方法只能防止第三者获得真实数据,而无法防止通信双方在通信时否认发送或接收过数据,发生伪造数据、篡改数据、假冒发送者或接收者等问题,解决这些问题的最好方法是使用数字签名机制。数字签名机制由两个过程组成:对信息进行签字的过程和对已签字信息进行证实的过程。数字签名机制必须保证签名只能由签名者的私有信息产生。

3) 访问控制机制

访问控制机制是从计算机系统的处理能力方面对信息提供保护。它根据实体的身份及其信息,来决定该实体的访问权限。访问控制按照事先确定的规则决定主体对客体的访问是否合法。当某一主体试图非法使用一个未经授权的资源时,访问控制将拒绝这一企图,并将这一事件报告给审计跟踪系统,审计跟踪系统将给出报警并记录日志档案。

4) 数据完整性机制

在网络中传输或存储数据可能会因为一些因素,使数据的完整性受到破坏。例如,有时数据在信道中传输时受到信道干扰,有时数据在传输和存储过程中被非法入侵者篡改,有时传输或存储的程序和数据感染上了计算机病毒。要避免这样的问题,应使用数据完整性机制。一般所说的数据完整性包括两种形式:数据单元的完整性和数据单元序列的完整性。保证数据完整性的一般方法是:发送实体在一个数据单元上加一个标记,这个标记是数据本身的函数,如一个分组校验或密码校验函数,它是经过加密的。接收实体产生一个对应的标记,并将所产生的标记与接收的标记相比较,以确定在传输过程中数据是否被修改过。

5) 认证机制

在计算机网络中认证主要有站点认证、报文认证、用户和进程的认证等。多数认证过程采用加密技术和数字签名技术。随着科学技术的发展,用户生理特性认证技术将得到越来越多的应用。

6) 业务流量填充机制

攻击者攻击的方法之一就是流量分析。攻击者通常通过分析网络中某一路径上的业务流量和流向来判断某些事件的发生。应付这种攻击,可以在无信息传输时,连续发送伪随机数据进行填充,使攻击者不知道哪些是有用信息哪些是无用信息,从而挫败业务流量分析攻击。但填充的信息只有经过加密保护才有效。

7) 路由控制机制

在大型计算机网络中,数据从源结点到达目的结点可能存在多条路径,其中有些路径是安全的,而有些路径是不安全的。路由控制机制可根据信息发送者的申请选择安全路径,以确保数据安全。为了使用安全的子网、中继站和链路,既可预先安排网络中的路由,也可对其动态地进行选择。

8) 公证机制

在大型计算机网络中可能会存在一些安全问题。例如,因为系统故障等原因使网络中的数据丢失。再如,网络中的有些用户不诚实、不可信。为了解决这些问题,需要有一

个各方都信任的第三方,它就像是一个国家设立的公证机构,来提供公证服务,仲裁出现的问题。引入公证机制后,通信双方进行数据通信时必须经过这个机构来交换,以确保公证机构能得到必要的信息,供以后仲裁时使用。仲裁数字签名技术就是这种公证机制的一种技术支持。

9) 安全审计跟踪机制

审计跟踪机制能创建和维护受保护客体的访问审计跟踪记录,并能阻止非授权的用户对它进行访问或破坏。审计跟踪机制提供了一种不可忽视的安全机制,它潜在的价值在于经事后的安全审计可以检测和调查网络中的安全漏洞。

安全服务与安全机制有着密切的联系,安全服务由安全机制来实现,它体现了网络安全模型的功能。一个安全服务可以由一个或几个安全机制来实现,同样,一个安全机制也可用于不同的安全服务。

1.3.3 网络安全策略

由于计算机网络具有连接形式多样性、终端分布不均匀性和网络的开放性、互连性等特征,致使网络易受黑客、恶意软件的攻击,所以网上信息的安全是一个至关重要的问题。因此,网络必须有足够强的安全策略。

1. 物理安全策略

制定物理安全策略的目的是保护计算机系统、网络服务器、打印机等硬件实体和通信链路免受自然灾害、人为破坏和搭线攻击;验证用户的身份和使用权限,防止用户越权操作;确保计算机系统有一个良好的电磁兼容工作环境;建立完备的安全管理制度,防止非法进入计算机控制室和各种偷窃、破坏活动的发生。

抑制和防止电磁泄漏是物理安全策略的一个主要问题。目前,主要防护措施有两类:一类是对传导发射的防护,主要采取对电源线和信号线加装性能良好的滤波器;另一类是对辐射的防护。对辐射的防护措施又可分为两种:一是采用各种电磁屏蔽措施,如对设备的金属屏蔽和各种接插件的屏蔽;二是干扰的防护措施,即在计算机系统工作的同时,利用干扰装置产生一种与计算机系统辐射相关的伪噪声向空间辐射来掩盖计算机系统的工作频率和信息特征。

2. 访问控制策略

访问控制是网络安全防范和保护的主要策略,它的主要任务是保证网络资源不被非法使用和非正常访问,也是维护网络系统安全、保护网络资源的重要手段。访问控制策略可以说是保证网络安全最重要的核心策略之一。下面介绍两种访问控制策略。

1) 入网访问控制

入网访问控制为网络访问提供了第一层访问控制。它控制哪些用户能够登录到服务器并获取网络资源,控制准许用户入网的时间和在哪台工作站入网。用户的入网访问控制可分为用户名的识别与验证、用户口令的识别与验证、用户账号的缺省限制检查。用户和用户组被赋予一定的权限,网络控制用户和用户组可以访问哪些目录、子目录、文件和其他资源,可以指定用户对这些文件、目录、设备能够执行哪些操作。

2) 目录级安全访问控制

网络应允许控制用户对目录、文件、设备的访问。用户在目录一级指定的权限对所有文件和子目录有效,用户还可进一步指定对目录下的子目录和文件的权限。对目录和文件的

访问权限一般有 8 种:超级用户(supervisor)权限、读(read)权限、写(write)权限、创建(create)权限、删除(erase)权限、修改(modify)权限、文件查找(file scan)权限、存取控制(access control)权限。网络系统管理员应当为用户指定适当的访问权限,这些访问权限控制着用户对服务器的访问,用户对网络资源的访问权限对应一张访问控制表,用以表明用户对网络资源的访问能力。访问权限的有效组合可以让用户有效地完成工作,同时又能有效地控制用户对服务器资源的访问,从而加强了网络和服务器的安全性。

3. 加密策略

信息加密的目的是保护网内的数据、文件、口令和控制信息,保护网上传输的数据。信息加密过程由多种加密算法来具体实施。在多数情况下,信息加密是保证信息机密性的方法。据不完全统计,到目前为止,密码技术是网络安全最有效的技术之一,已经公开发表的加密算法多达数百种。一个加密网络,不但可以防止非授权用户的搭线窃听和入网,而且也是对付恶意软件的有效方法之一。

在网络安全中,除了综合运用上述策略之外,加强网络的安全管理,制定有关规章制度,对于确保网络安全、可靠地运行,也将起到十分有效的作用。

1.4 网络安全评估标准

前面章节介绍了网络中存在的安全问题,那么,怎样评价一个网络的安全程度呢?在 20 世纪 70 年代,David Bell 和 Leonard La Padula 开发了一个安全计算机的操作模型。该模型是基于政府机构的各种级别分类信息(一般、秘密、机密、绝密)和各种许可级别。如果主体的许可级别高于客体(文件)的分类级别,则主体能访问客体。如果主体的许可级别低于客体(文件)的分类级别,则主体不能访问客体。这个模型的概念进一步发展,就成了可信任计算机系统评估标准。

1.4.1 可信任计算机系统评估标准

1983 年,美国国防部制定了 5200.28 标准——可信任计算系统评估标准 TCSEC(the trusted computing system evaluation criteria),即桔皮书(orange book)。桔皮书自从 1985 年成为美国国防部的标准以来,一直是评估多用户主机和小型操作系统的主要方法,其他子系统,如数据库和网络,也是通过桔皮书的解释来评估的。

桔皮书将安全分为不同的等级,这些等级描述了不同类型的物理安全、用户身份验证、操作系统软件的可信任性和用户应用程序。这些标准也限制了什么类型的系统可以连接到用户系统,并给出了一套标准来定义满足特定安全等级所需的安全功能及其保证的程度。

TCSEC 中定义了系统安全的 5 个要素:

- (1)系统的安全策略。
- (2)系统的可审计机制。
- (3)系统安全的可操作性。
- (4)系统安全的生命期保证。
- (5)针对以上系统安全要素而建立并维护的相关文件。

同时,TCSEC 还定义了系统安全等级来描述以上所有要素的安全特性:

- (1)D:最低保护(minimal protection)。
- (2)C:被动的自主访问策略(discretionary access policy enforced)。
- (3)B:被动的强制访问策略(mandatory access policy enforced)。
- (4)A:形式化证明的安全(formally proven security)。

每个等级之内还可以细分,具体介绍如下。

1. D1 级

D1 级是可用的最低的安全等级。该标准说明整个系统都是不可信任的。对于硬件来说,没有任何保护可言,操作系统容易受到损害;对于用户来说,当他们对存储在计算机上的信息进行访问时,没有身份验证要求。

该安全级别系统,典型的如 MS-DOS、MS-Windows 和 Apple 公司的 Macintosh System 等。这些操作系统不区分用户,没有定义方法来决定谁在敲击键盘,也没有控制计算机硬盘上的什么信息是可以访问的。

2. C1 级

C 级有两个安全子级别:C1 级和 C2 级。C1 级也称自选安全保护(discretionary security protection)系统,它描述了一个典型的 UNIX 系统上可用的安全等级。C1 级计算机系统要求硬件有一定的安全机制,用户必须注册并用注册口令登录,通过系统的认证并判断其访问权限。用户的访问权限是指文件和目录许可权限(permission)。自选访问控制(discretionary access control)使文件和目录的拥有者或者系统管理员,能够阻止某个人或几组人访问某些程序和信息。但是,这并没有阻止系统管理账户执行活动,不审慎的系统管理员可能损害系统安全。

另外,许多日常系统管理任务只能由以 root 注册的用户来执行。随着现在计算机系统的分散化,随便走进一个组织,都会发现两个以上的人知道根口令。由于无法区分具体是谁对系统作过改变,所以这本身就是一个问题。

3. C2 级

第二个子级别 C2 级可用来帮助解决上述问题。除包含 C1 级的特征外,C2 级还包含受控的访问环境(controlled access environment)的安全特征。该环境具有进一步限制用户执行某些命令或访问某些文件的能力,这不仅给予许可权限,而且给予身份验证级别。另外,这种安全级别要求对系统加以审核,包括为系统中发生的每个事件编写一个审核记录。审核用来跟踪记录所有有关安全的事件,例如,由系统管理员执行的活动。审核还要求身份认证,以此来确定实际执行命令的人。其缺点在于它需要额外的处理器和磁盘子系统资源。

对于一个 C2 级系统的用户来说,使用附加身份验证,在没有根口令的情况下,有权执行系统管理任务是可能的。即单独的用户执行了工作而不是系统管理员,这使得追踪与系统管理有关的任务发生了变化。但是,附加身份验证不能与可应用于程序的 SGID 和 SUID 许可权限相混淆。这些附加身份验证还是允许用户执行特定命令或访问某些核心表的特定身份验证。例如,无权浏览进程表的用户,当执行 PS 命令时,只能看到用户自己的进程。

4. B1 级

B 级安全包含 3 个子级别。B1 级也称标志安全保护(labeled security protection),是支持多级安全(如秘密和绝密)的第一个级别。这个级说明一个处于强制性访问控制之下的对

象,系统不允许文件的拥有者改变其许可权限。

5. B2 级

B2 级也称结构保护(structured protection),要求计算机系统中所有对象都加标签,而且给设备(磁盘、磁带或终端)分配单个或多个安全级别。这是提出较高安全级别的对象与另一个较低级别的对象通信的第一个级别。

6. B3 级

B3 也称安全域级别(security domain),使用安装硬件的办法来加强域,例如,内存管理硬件用于保护安全域免遭无权访问或其他安全域对象的修改。该级别也要求用户的终端通过一条可信任的途径连接到系统上。

7. A 级

A 级也称验证设计(verified design),是当前桔皮书中的最高安全级别,它包含了一个严格的设计、控制和验证过程。与前面提到的各级别一样,这一级别包含了较低级别的所有特性。设计必须是从数学上经过验证的,而且必须进行对秘密通道和可信任分布的分析。可信任分布(trusted distribution)的含义是,硬件和软件在传输过程中受到保护,以防止破坏安全系统。

以上这些标准能够被用来衡量计算机平台(如操作系统及其基于的硬件)的安全性。如标准的 UNIX(只有 login 口令、文件保护等安全措施)被定为 C1 级,DOS 被定为 D1 级。目前,很少有操作系统能够符合 B 级标准。

1.4.2 国际安全标准

在借鉴 TCSEC 成功经验的基础上,20 世纪 90 年代初,西欧四国(英国、法国、荷兰、德国)共同组成的欧洲委员会统一了他们的信息技术安全评估准则 ITSEC(information technology security evaluation criteria)。

在 ITSEC 中,一个基本观点是:应当分别衡量安全的功能和安全的保证,而不应像 TCSEC 那样混合考虑安全的功能和安全的保证。因此,ITSEC 对每个系统赋予两种等级:F 即安全功能等级,E 即安全保证等级。这样,一个系统可能有最高等级所需的所有安全功能(F6),但由于某些功能不能保证到最高等级,从而使该系统的安全保证等级较低(E4),此系统的安全等级将是 F6/E4。

在 ITSEC 中,另一个基本观点是:被评估的应是整个系统(如硬件、操作系统、数据库管理系统、应用软件),而不只是计算平台。因为一个系统的安全等级可能比其每个组成部分的安全等级都高。另外,某个等级所需的总体安全功能可能分布在系统的不同组成中,而不是所有组成都要重复这些安全功能。

1993 年,加拿大发布了加拿大可信计算机产品评价准则 CTCPEC。

20 世纪 90 年代中期,美国联合西欧四国和加拿大,并会同国际化标准组织提出了制订通用安全评价准则(CC)的计划,全称是 common criteria for IT security evaluation。它定义了 IT 安全评价和描述模型的一般概念和原则,提出了选择、定义和说明产品以及系统 IT 安全客体的明确的安全要求。

1.4.3 我国安全标准

计算机信息安全越来越成为关系到国计民生的大事,而国内还缺乏一套全面而专业化

的信息系统安全建设和管理的标准规范。我国是国际标准化组织的成员国,我国的信息安全标准化工作也取得了一系列的成果。从 20 世纪 80 年代中期开始,自主制定和采用了一批相应的信息安全标准。1999 年 9 月 13 日,由国家质量技术监督局对外公布了国家标准《计算机信息系统安全保护等级划分准则》。该准则结合我国信息安全系统建设的实际,不同用户可以根据自己的安全需求,对应不同的等级标准来制订适合于自己的信息安全系统解决方案。该准则将计算机信息系统的安全等级划分为 5 级。

- 第 1 级:用户自主保护级。
- 第 2 级:系统审计保护级。
- 第 3 级:安全标记保护级。
- 第 4 级:结构化保护级。
- 第 5 级:访问验证保护级。

2001 年 1 月 1 日,该准则开始实施。它的实施一方面为计算机信息系统安全法规的制定提供了有力支持,另一方面为计算机信息系统安全产品研发提供了较详细的功能框架,还为安全系统的建设和管理提供了明确的技术指导。

1.5 实践项目

1.5.1 虚拟机环境搭建

现在很多人都拥有个人计算机,但是组建一个自己的局域网或者是做小规模实验,一台计算机是不够的,最少也要两台。而虚拟机就可以解决这个问题。虚拟机可以在一台计算机上虚拟出很多的主机,并可以在该平台上安装多个操作系统。

VMware 是 VMware 公司出品的一个虚拟机软件。通过安装 VMware 可以在一台计算机上将硬盘和内存的一部分拿出来虚拟出若干台机器,每台机器可以运行单独的操作系统而互不干扰,这些“新”机器各自拥有自己独立的 CMOS、硬盘和操作系统,可以对它们进行分区、格式化、安装系统和应用软件等操作,所有的这些操作都是一个虚拟的过程不会对真实的主机造成影响。在此平台上将虚拟出来的几个操作系统连成一个网络,从而在单机环境下实现多机的网络实验。

VMware 的硬件要求如下:

- 处理器:400 MHz 或者更快(建议 1 GHz),单个或者多个处理器。
- 内存:最小 128 MB(建议 512 MB)。
- 磁盘驱动器:基本安装需要 100 MB 空闲磁盘空间。客户操作系统和应用程序建议至少 500 MB 空闲磁盘空间。

运行虚拟机安装程序 VMware Workstation V6.0.2,根据提示安装即可,具体的过程这里不作详述。安装过程如图 1-2 所示,安装完毕后如图 1-3 所示,单击 Finish 按钮即可。



图 1-2 安装 VMware Workstation



图 1-3 VMware Workstation 安装完毕

1.5.2 虚拟机设置

虚拟机的设置步骤如下：

(1) VMware Workstation 安装完毕后，双击桌面的虚拟机图标，首次启动虚拟机，如图 1-4 所示。



图 1-4 首次启动虚拟机

(2)单击“新建虚拟机”按钮，打开“新建虚拟机向导”对话框，如图 1-5 所示。

(3)单击“下一步”按钮，打开“选择合适的配置”对话框，这里选择“典型”配置，如图 1-6 所示。



图 1-5 “新建虚拟机向导”对话框



图 1-6 虚拟机配置

(4)单击“下一步”按钮，打开“选择一个客户机操作系统”对话框，这里选择 Microsoft Windows，如图 1-7 所示。

(5)单击“下一步”按钮，打开“虚拟机名称”对话框，设置虚拟机名称和位置，如图 1-8 所示。

(6)单击“下一步”按钮，打开“网络类型”对话框，选择网络连接的类型，如图 1-9 所示。

(7)单击“下一步”按钮，打开“指定磁盘容量”对话框，为虚拟机分配磁盘空间，如图 1-10 所示。这里建议分配 1 GB 以上的空间。



图 1-7 选择客户机操作系统



图 1-8 虚拟机名称和存储路径

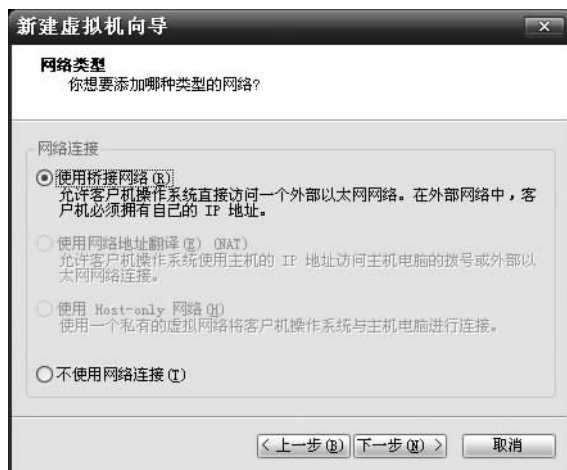


图 1-9 网络连接的类型

(8)单击“完成”按钮,打开如图 1-11 所示的对话框,单击 Close 按钮,即完成新建虚拟机的工作。

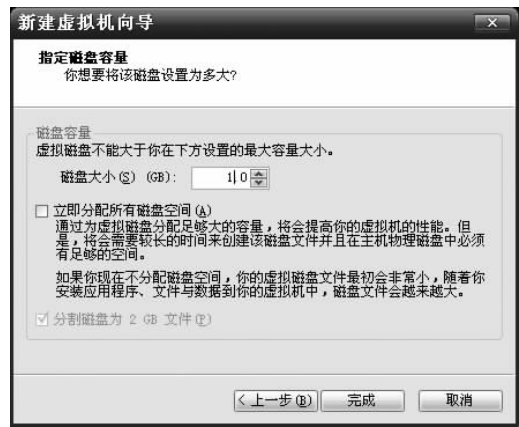


图 1-10 虚拟机分配磁盘空间



图 1-11 完成创建虚拟机

安装完 Windows XP 操作系统的界面如图 1-12 所示。单击“启动该虚拟机”命令,即可在虚拟机中打开 Windows XP 操作系统。



图 1-12 创建完成显示界面

VMware 允许操作系统和应用程序在一台虚拟机内部运行。在 VMware 中,可以在一个窗口中加载一台虚拟机并运行自己的操作系统和应用程序,也可以在运行于桌面上的多台虚拟机之间切换,通过一个网络共享虚拟机(如一个公司局域网)挂起和恢复虚拟机以及退出虚拟机,这一切不会影响主机操作和任何操作系统或者它正在运行的应用程序。例如,一个需要在 Windows XP 或 Linux 中进行测试的操作,也可以使用 VMware Workstation 来搭建平台。

在后面章节的例子中,有些需要在虚拟机上进行操作,配置如本节所述,具体过程将不再赘述。

本章小结

本章主要介绍了网络安全相关基础理论。重点介绍了网络安全的基础知识,强调了威胁网络安全的主要因素;介绍了网络安全防范的层次、体系结构以及安全策略;然后对当前主流的网络安全评估标准,特别是可信任计算机系统评估标准作了详细的介绍;最后,作为实践内容,介绍了虚拟机环境的搭建和配置方法及简单应用。

本章对于网络安全的学习是非常重要的。其网络安全的基础知识可以使初学者对于网络安全有基本的了解,为今后的学习打下理论基础。

习 题 1

1. 网络安全的概念是什么? 为什么要重视网络安全问题?
2. 计算机网络安全面临的主要威胁有哪些?
3. 哪些人对网络安全构成威胁? 网络安全威胁分为哪些类型?
4. 为什么要重视来自内部的或外部有组织的威胁?
5. 从层次上分析,网络安全防范可以分成哪几层? 每层有什么特点?
6. 网络安全桔皮书指的是什么? 包括哪些内容?
7. 国内和国际上对于网络安全方面有哪些立法?
8. 简述网络安全管理的意义和主要内容。
9. 分析计算机网络安全技术的主要应用和发展趋势。
10. 分析计算机网络的安全需求。

第 2 章 网络入侵技术

随着网络技术的广泛应用,网络遭受各种攻击的概率也越来越大,如何有效地应对黑客的入侵和攻击已成为当前网络安全领域的重要课题。有效地防范黑客的入侵和破坏,不仅需要掌握网络安全技术,还应该了解一些常见的黑客入侵手段,做到知己知彼,从而针对各种攻击有的放矢地采取不同的措施,将网络攻击造成的损失降到最小。

2.1 黑客技术

2.1.1 黑客的由来

黑客最早源自英文 hacker,原指热衷于计算机程序的设计者和精通网络、系统、外围设备及软硬件技术的人。这些人具有操作系统和编程语言方面的知识,通过分析知道系统中的漏洞及其原因所在,并公开他们的发现,与其他人分享。他们以改进的目的编写程序去检查远程机器的安全体系。

黑客大体上分为白帽(white hat)、黑帽(black hat)两类。白帽黑客依靠自己掌握的知识帮助系统管理员找出系统中的漏洞并加以完善,而黑帽黑客则是通过各种黑客技能对系统进行攻击、入侵或者做一些其他有害于网络的事情,所以称其为骇客(cracker),而非黑客。黑客和骇客有着根本的区别,黑客发现和修补漏洞,而骇客则利用漏洞进行破坏。但他们最初掌握的基本技能是一样的,所做的事情也差不多,不同的是出发点和目的。

现今,黑客一词已被用于泛指那些专门利用计算机网络搞破坏或恶作剧的人,所以人们常常把从事网络攻击和破坏的人统称为“黑客”。

2.1.2 黑客攻击的动机

2009年8月6日,美国两大社交网站 Facebook 和 Twitter 同时遭遇黑客攻击,用户在登录两大网站时均显示“拒绝服务”。Facebook 网站性能严重受损,而 Twitter 网站则被中断长达数小时,数百万网民无法登录这两大网站。Twitter 网站创始人比兹·斯通在该公司的官方博客中写道:“6日上午, Twitter 网站受到拒绝服务的攻击,这样的攻击完全是恶意攻击。”当天,谷歌网站也是拒绝服务攻击的目标, Facebook 和 Twitter 两家网站已经联手谷歌,展开对黑客身份的调查。

通过对大量的案例进行分析,发现黑客主要有以下几种犯罪动机。

1. 恶作剧

有些黑客往往是以验证自己的能力为目的,专门利用网络漏洞入侵主机,扮演恶作剧的角色,如中美黑客联盟大战。这种黑客在早期出现较多,以青少年居多。

2. 窃取信息

部分人往往会偷窥他人在计算机中的隐私和机密资料,以盗取隐私和相关资料为目的,甚至对该资料进行删改。如“蜜蜂大盗”(偷拍对方照片)。

3. 金钱目的

有些黑客为达到牟利的目的而盗取或篡改机密资料。这种黑客近年来有所发展,并逐步形成了完整产业链,如图 2-1 所示。如盗取网民网上银行账号及密码的“网银大盗”、“证券大盗”等。

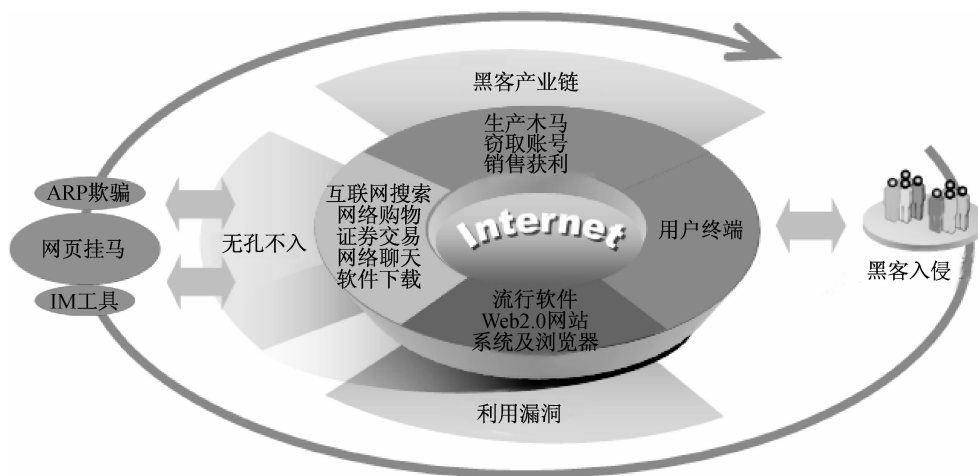


图 2-1 黑客产业链(资料来源于瑞星网站)

4. 政治目的

如敌对国家之间利用网络进行破坏,或是由于个人对政府不满而产生的破坏活动都属于这个类型。有些黑客专门以破坏政府部门保密系统为乐,如白宫主页就曾遭黑客攻击。

5. 报复

由于个人对组织、公司不满或是对他人的怨恨,进而采取破坏网络、盗取数据等行为对其进行报复,或是借此引起别人的注意。

6. 提升个人声望

通过破坏具有高价值的目标,或者具有高技术难度的破坏行为引起他人的关注。借此向他人炫耀技术,并证明自己的能力。

2.1.3 黑客入侵攻击的一般过程

黑客的攻击步骤一般可以分踩点(foot printing)、扫描(scanning)、查点(enumeration)、获取权限(gaining access)、提升权限(escalating)、窃取(pilfering)信息、隐藏痕迹(covering track)、创建后门(creating back door),如图 2-2 所示。

1. 踩点

攻击任何一个网络,第一步就是要搞清楚要攻击的对象,获取目标网络的“足迹”。通过结合使用工具和技巧,攻击者能够得到尽可能多的主机上的信息。踩点主要的目的是获得以下几方面的信息:

- 主机名称

- 主机上暴露的应用程序
- 操作系统和应用程序的版本信息
- 联系人姓名和电子邮件地址
- 指示所用安全机制的类型的隐私和机密保障策略
- 与其相关联的 Web 服务器超链接
- 网络地址范围
- 主机和应用程序的补丁状态
- 应用程序和后端服务器的结构

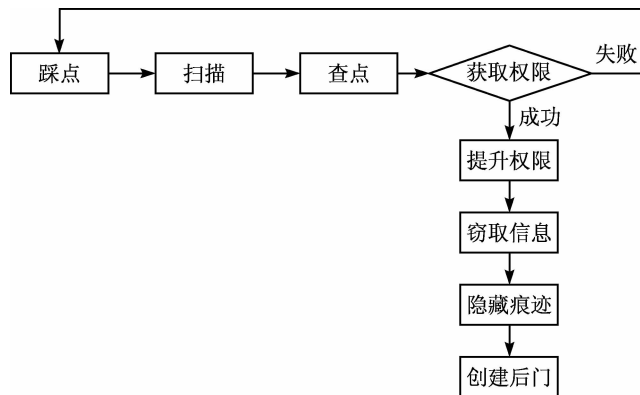


图 2-2 黑客入侵过程

为了获取以上信息,黑客经常采用以下技术:

(1)阅读 HTML 源代码。在网页中右击,在弹出的快捷菜单中选择“查看源文件”。通过这种查看页面源代码的方式往往能够找到网站的建站程序、程序架构体系及程序的漏洞等信息。例如,某网站的源代码 Keywords 关键词后面包含了 DedeCMS,此信息说明该网站是依靠 DedeCMS 程序建设的。CMS 是网站内容管理系统的意义,通常也就是管理员建设网站时选用的网站程序。在了解了网站的建站程序后,黑客可以通过搜索引擎搜索相关程序的漏洞,即便没有可以利用的漏洞信息,也可以等待相应程序出现新漏洞时在第一时间入侵。

在搜索引擎中可以很快搜索到,DedeCMS 是一套基于 PHP+MySQL 架构体系,支持多网站动静混合发布的网站内容管理系统。而通过在搜索引擎中搜索“DedeCMS+漏洞”这样的关键词,可以查找 DedeCMS 在近期是否出现过漏洞信息以及如何利用,这也是黑客入侵前需要整理的重要资料。

(2)who is 查询。who is 是目标主机 Internet 域名注册数据库。作为一个机构,要上互联网就必须去登记,所以登记的地方就必然有它的一些资料。可用的 who is 数据库很多,如 network solutions,all who is 等。通过此查询可以得到的信息如下:

- 注册人
- 域名
- 管理方面联系人
- 记录创建时间和更新时间
- 域名系统(DNS)服务器

who is 查询完成后,用搜索引擎二次查询,搜索 E-mail 地址、电话号码等相关信息,就可以在网络中查询到该网站注册人的一些信息,这些信息有时能够让黑客充分掌握网站管理员的个人信息及在网络中的行踪,甚至管理员所常用的密码就是搜索出的电话号码或者生日等。通过搜索到的域名注册信息,可以追踪到是哪一家域名空间服务商为网站提供的服务,由此可以去服务商的网站查询出服务器的类型、配置、操作系统等详细信息。

2. 扫描

踩点完成后,常用扫描器对目标网络进行扫描。扫描器不仅可以很快地发现本地主机系统配置和软件上存在的安全漏洞,而且还可以不留痕迹地发现一台主机的安全性漏洞,这种自动检测功能快速而准确。扫描的主要目的是使攻击者对目标网络所提供的各种服务进行评估,以便找到最有希望成功的攻击方式。

黑客正是利用扫描器找出目标主机上各种各样的安全漏洞,虽然这时还不能直接攻击网络,但是却为接下来的攻击确定了最有效的攻击方式。

通过扫描还可以了解主机开放了哪些端口、端口所用的协议和服务类型等。下面是一个运用 nmap 的扫描范例。该例中用隐身扫描 SYN 方法扫描,开始一次 SYN 的半开扫描,针对的目标是 target.example.com 所在的 C 类子网。可以用自己在网络上的名称代替其中的 addresses/names。

```
[japleak@root] $ nmap -sS 192.168.1.1
Starting nmap v. 2.12 by Fyodor (fyodor@dhp.com,www.insecure.org/nmap/)
Interesting ports on (192.168.1.1)

Port      State    Protocol  Service
21        open    tcp       ftp
25        open    tcp       smtp
42        open    tcp       nameserver
53        open    tcp       domain
79        open    tcp       finger
80        open    tcp       http
81        open    tcp       hosts2-ns
106       open    tcp       pop3pw
110       open    tcp       pop-3
135       open    tcp       loc-srv
139       open    tcp       netbios-ssn
443       open    tcp       https
```

从以上程序运行的结果中可以看到,135 号和 139 号端口同时打开。Windows NT/XP 通常同时监听 135 号和 139 号端口,目标操作系统是 Windows NT/XP 的可能性很大,因为 Windows 9x 一般只监听 139 号端口。

3. 查点

通过扫描,入侵者掌握了目标主机所使用的操作系统及系统中使用哪个端口等信息。查点可以使入侵者搜索到特定系统上用户的用户组名、路由表、SNMP 信息、共享状况、服务程序等信息。一般情况下,对不同的操作系统使用的查点技术也不一样。

Windows 系统上主要采用 NetBIOS、空会话(null session)、SNMP 代理等。例如:

```
C:\>net view/domain
```

```
Domain
```

```
CORLEONE
```

```
BARZINI_DOMAIN
```

```
TATAGGLIA_DOMAIN
```

```
BARZZI
```

```
The command completed successfully
```

Windows NT/XP 中存在一个漏洞,就是可以建立空会话的链接,允许匿名登录主机,也就存在着暴露共享资源的危险,这样无疑给攻击者提供了一个很容易获得的攻击点。例如:

```
net use\\192.168.202.33\IPC$ ""/user:""
```

这样建立空会话后,就可以利用 NetBIOS 共享中的 net view 查看主机上存在的共享资源。程序如下:

```
C:\>net view\\vito
```

```
Shared resources at\\192.168.7.45
```

```
VITO
```

```
Share name Type Used as Comment
```

```
NETLOGON Disk Logon server share
```

```
Test Disk Public access
```

```
The command completed successfully
```

4. 获取权限

获得了目标主机的信息之后,还需获取足够的访问权限,黑客才能对系统进行攻击。

目前,绝大多数系统都是将口令作为数据访问的第一道屏障。针对流行的操作系统特别是 Windows 系统常见的获得口令的方式有密码猜测、窃听认证散列、IISWeb 服务器漏洞和远程溢出等。当然,对攻击者来说,攻破了这道屏障,就意味着获得了进入系统第一道大门的资格,所以口令攻击是攻击者最常用的攻击手段。

缓冲区溢出也是常见的入侵方法。缓冲区是程序运行时在内存中为保存给定类型的数据而开辟的一个连续空间,这个空间是有限的。当程序运行过程中要放入缓冲区的数据太多时,就会产生缓冲区溢出。攻击者向程序的缓冲区写超出其长度的内容,造成缓冲区的溢出,从而破坏程序的堆栈,使程序转而执行其他命令,以达到攻击的目的。

5. 提升权限

攻击者在获得普通用户的访问权限后,会试图提升权限直至获得超级用户权限,从而完成对系统的绝对控制。提升权限的方法主要是获得密码文件,进而使用破解工具破解口令。例如,Windows NT/XP 中对用户账户的安全管理使用了安全账号管理器(SAM)的机制,SAM 文件是 Windows NT/XP 的用户账户数据库,所有用户的登录名以及口令等相关信息都会保存在这个文件中。使用 LC4 工具可以打开待破解的 SAM 文件,此时 LC4 会自动分析此文件,并显示出文件中的用户名,之后开始破解密码。如果密码不是很复杂,很短的时间内就会得到结果。即使密码比较复杂,经过一定时间后也可能被破解出来。使用字典的方法破解可以缩短破解的时间。

6. 窃取信息

攻击者下一步的行动可能是窃取主机上的各种敏感信息,如软件资料、客户名单、财务报表、信用卡卡号等;也可能是什么都不动,只是把该系统作为他存放资料的仓库;也可能会利用这台已经入侵的主机继续下一步的攻击,如继续入侵内部网络或者利用这台主机发动 DDoS 攻击使网络瘫痪等。2008 年,一个全球性的黑客组织利用 ATM 欺诈程序在一夜之间盗走了 900 万美元。黑客攻破的是一种名为 RBS WorldPay 的银行系统,利用团伙作案的方式从世界 49 个城市总计超过 130 台 ATM 上提取了 900 万美元。

7. 隐藏痕迹

黑客利用各种手段进入目标主机系统并进行破坏活动之后,为了能长时间地保留和巩固他对系统的控制权,并且不被管理员发现,一般会清除入侵和操作记录。日志往往会记录一些黑客攻击的蛛丝马迹,黑客为了不留下这些破坏证据,会将这些记录删除或用假日志覆盖原有记录。

8. 创建后门

为了日后不被觉察地再次进入系统,黑客还会更改某些系统设置,在系统中置入特洛伊木马或其他一些远程操作程序,这种行为被称为创建后门。创建后门的方法常见的有安装批处理、使用木马程序替换系统程序、安装监控机、感染启动文件等。

在许多入侵发生后,黑客会发起 DDoS 攻击。发起攻击的计算机主机常常是宽带接入互联网的已受到病毒或特洛伊木马感染的个人计算机,这些计算机因受到感染而被攻击者远程控制并发起攻击。只要有足够多这样的“僵尸”计算机,大型网站也不能逃脱被阻断服务的厄运。

新的攻击总是层出不穷的。现有的攻击工具就有近百种且新的工具极易生成,所以造成网络入侵事件频繁发生,产生的危害性也越来越大。系统维护人员需要将黑客攻击和信息安全技术紧密结合,利用常见的攻击手段对系统进行检测,并对相关的漏洞采取措施。

2.2 网络扫描

网络扫描就是对计算机系统或其他网络设备进行相关的安全检测,以便发现安全隐患和可被黑客利用的漏洞。而非授权用户入侵一台目标主机,要先测试出网络上的哪些主机是活动的,哪些端口在监听等,这些工作是由扫描器来实现的。

现在网络上有很多扫描器,它们在功能上都设计得非常强大,并且综合了各种扫描需要,将各种功能集成于一身。但有经验的管理员通常使用自己编写开发的扫描工具,这样其在功能上将会完全符合个人意图,而且可以针对新漏洞及时对扫描器进行修改,在第一时间获得最宝贵的目标资料。

扫描器也是网络管理员的得力助手,网络管理员可以通过及时了解自己系统的运行状态和可能存在的漏洞,在黑客攻击之前将系统中的隐患清除,保证服务器的安全稳定。

2.2.1 地址扫描

地址扫描就是判断某个 IP 地址上是否有活动主机以及某台主机是否在线。这是信息收集的初级阶段,其效果直接影响到后续的扫描。ping 就是最原始的主机存活扫描技术,利用

ICMP 的 Echo 字段,用 ping 命令去连接某台主机,发出的请求如果收到回应则代表主机是活动的。

传统的地址扫描手段有:

(1)ICMP Echo 扫描:精度相对较高,向目标主机发送 ICMP Echo Request 数据包,并等待回复的 ICMP Echo Reply 数据包。

(2)ICMP Sweep 扫描:ICMP 进行扫射式的扫描,就是并发性扫描,使用 ICMP Echo Request 一次探测多个目标主机。通常这种探测包会并行发送,以提高探测效率,适用于大范围的评估。

(3)Broadcast ICMP 扫描:广播型 ICMP 扫描,利用了一些主机在 ICMP 实现上的差异,设置 ICMP 请求包的目标地址为广播地址或网络地址,则可以探测广播域或整个网络范围内的主机,子网内所有活动主机都会给以回应。但这种情况只适合于 UNIX/Linux 系统。

(4)Non-Echo ICMP 扫描:在 ICMP 协议中不仅 ICMP Echo 的 ICMP 查询信息类型,ICMP 扫描技术中也用到 Non-Echo ICMP 技术,它利用了 ICMP 的多种服务类型,如 Timestamp 和 Timestamp Reply、Information Request 和 Information Reply、Address Mask Request 和 Address Mask Reply。它既能探测主机,也可以探测网络设备(如路由)。

2.2.2 端口扫描

在完成主机活动性判断之后,就应该去判定主机开放信道的状态,端口就是在主机上面开放的信道,0~1 024 为知名端口,端口总数是 65 535。端口实际上就是从网络层映射到进程的通道。端口是目标系统向外提供服务的窗口,是一个潜在的通信通道,当然也就能够成为一个入侵的通道。黑客常常采用对端口进行扫描的方法攻击。

对目标主机进行端口扫描,能得到许多有用的信息,能够发现系统的安全漏洞,还可以掌握什么样的进程使用了什么样的通信。通过进程取得的信息,为查找后门、了解系统状态提供了有力的支撑。

如果按照端口对应的协议来划分,端口扫描分为 TCP 端口和 UDP 端口。像 21 号端口、23 号端口、80 号端口等这些都是 TCP 协议使用的端口。一般黑客知道了要攻击目标主机的 IP 地址后,还要知道通信的端口号,只要扫描到相应的端口被打开着,就能知道目标主机上运行着什么样的服务,进而有针对性地采取攻击手段。例如,测试到目标主机的端口号是 21,那么主机上正在提供的是 FTP 服务;测试到目标主机的端口号是 80,那么主机提供的就是 HTTP 服务等。

端口扫描的目的是:判断目标主机中开放了哪些服务和判断目标主机的操作系统。如果要想了解端口的开放情况,必须知道端口是如何被扫描的。在端口扫描的具体实现中,扫描软件将尝试与目标主机的某些端口建立连接,如果目标主机的该端口有回复,则说明该端口开放。

常见的端口扫描方法有以下几种:

(1)TCP connect()扫描。这是最基本的 TCP 扫描技术,很容易实现。如果目标主机能够进行 connect 操作,就说明有一个相应的端口打开。否则,这个端口是不能用的,即没有提供服务。使用这种方法可以检测到目标主机开放了哪些端口。在执行这种扫描方式时,不需要对目标主机拥有任何权限。不过,这种扫描是最原始和最先被防护工具拒绝的一种方法。因为在扫描时往往被远程系统记入日志,非常容易被发觉并被过滤掉。

(2)TCP SYN 扫描。为了克服 TCP connect()扫描缺陷,便产生了 TCP SYN 扫描。这

种扫描是向远程主机某端口发送一个只有 SYN 标志位的 TCP 数据包,也就是发出一个连接请求。如果收到了远程目标主机的 SYN/ACK 数据包,那么说明远程主机的该端口是打开的;若没有收到远程目标主机的 SYN/ACK 数据包,而收到的是 RST 数据包,则说明远程主机的该端口没有打开。这对扫描要获得的信息已经足够了,这种扫描的特点是不会在目标主机的日志中留下记录。

(3)TCP FIN 扫描。由于上述扫描的广泛应用,使得防火墙和路由器都采取了相应的措施,会对端口扫描进行完全记录。有些入侵扫描系统也能检测到 TCP SYN 扫描,许多过滤设备能过滤 SYN 数据包。于是端口扫描开始采用 FIN 扫描。

FIN 是中断连接的数据包,很多日志不记录这类数据包。TCP FIN 扫描的原理是:向目标端口发送 FIN 数据包,如果收到了 RST 的回复,表明该端口没有开放;如果没有收到 RST 的回复,表明该端口是开放的,因为打开的端口往往忽略对 FIN 的回复。这种方法还可以用来区别操作系统是 Windows,还是 UNIX。

但是,有的系统不管端口打开与否,一律回复 RST。这时,FIN 扫描就不适用了。

(4)IP 段扫描。这种方法不直接发送 TCP 探测数据包,而是预先将数据包分成几个较小的 IP 数据包传送给目标主机,目标主机收到这些 IP 数据包后,会把它们组合还原为原先的 TCP 探测数据包。将数据包分片的目的是使这些数据包能够通过防火墙和包过滤器而到达目标主机。

(5)UDP ICMP 扫描。这种方法与上面几种方法的不同之处在于,此方法使用的是 UDP 协议。由于防火墙设备的流行,TCP 端口的管理状态越来越严格,不会轻易开放,并且通信监视严格。为了避免这种监视,达到评估的目的,就出现了 UDP 扫描。

UDP 的扫描方法比较单一,基本原理是:当发送一个报文给 UDP 端口,该端口是关闭状态时,端口会返回一个 ICMP 信息;如果端口是打开的,什么信息都不发。这种扫描方式的特点是利用 UDP 端口关闭时返回的 ICMP 信息,不包含标准的 TCP 三次握手协议的任何部分,隐蔽性好,但这种扫描使用的数据包在通过网络时容易被丢弃从而产生错误的探测信息。在使用 UDP 扫描时需要注意:

- UDP 状态、精度比较差,因为 UDP 不是面向连接的,所以整个精度会比较低。
- UDP 扫描速度比较慢,TCP 扫描开放 1 s 的延时,在 UDP 里可能就需要 2 s,这是由于不同操作系统在实现 ICMP 协议的时候为了避免广播风暴都会有峰值速率的限制。ICMP 信息本身并不是传输载荷信息,所以不会有人用它去传输一些有价值信息。为了避免产生广播风暴,操作系统对 ICMP 报文规定了峰值速率,不同操作系统规定的速率不同。

2.2.3 漏洞扫描

漏洞扫描能自动检测远程或本地主机的安全性弱点及存在的安全缺陷。其原理是采用模拟攻击的形式,对工作站、服务器、交换机、数据库等可能存在已知安全漏洞的目标对象进行逐项检查,根据扫描结果形成安全性分析报告。

按常规标准,可以将漏洞扫描分为两种类型:主机漏洞扫描(host scanner)和网络漏洞扫描(network scanner)。主机漏洞扫描是指在系统本地运行检测系统漏洞的程序,该类型的软件有著名的 COPS、tripwire、tiger 等。网络漏洞扫描是指基于 Internet 远程检测目标网络和主机系统漏洞的程序,该类型的软件有 Satan、ISS Internet Scanner 等。

通过漏洞扫描,系统管理员能够发现所维护的 Web 服务器的各种 TCP 端口的分配、提供的服务、Web 服务软件版本和这些服务及软件呈现在 Internet 上的安全漏洞。从而在计算机网络系统安全保卫中做到有的放矢,及时修补漏洞。

漏洞扫描的结果实际上就是系统安全性能的一个评估,它指出了哪些攻击是可能的。它是安全漏洞扫描方案的一个重要组成部分。

2.2.4 常用的扫描软件

无论是地址扫描、端口扫描还是漏洞扫描,都离不开扫描软件。有些扫描软件只能做一种扫描,有些扫描软件功能非常强大,集多种扫描功能于一身。常用的扫描软件如下:

(1)GetNTuser:扫描 NT 主机上存在的用户名、自动猜测空密码和与用户名相同的密码,可以使用指定密码字典猜测密码,也可以使用指定字符来穷举猜测密码等。

(2)PortScan:可以得到对方计算机开放的端口,该工具软件可以将所有端口的开放情况做一个测试,通过端口扫描,可以知道对方开放了哪些网络服务,从而根据某些服务的漏洞进行攻击。

(3)Shed:通过该工具软件来扫描对方主机,就可以知道对方计算机哪些目录是共享的。

(4)X-Scan(X-Scan-v2.3):是国内最著名的综合扫描器之一,它主要进行漏洞检测。扫描内容包括远程操作系统类型及版本,标准端口状态、IIS 漏洞、CGI 漏洞、SQL-SERVER、FTP-SERVER、注册表信息等。该软件可以查看一台主机的端口开放情况,同时也可以查看一段连续 IP 地址的端口开放情况,这对分析整个网络的端口开放情况很有用处。

(5)SuperScan:这是非常著名的端口扫描软件。该软件是一个集端口扫描、ping、主机名解析于一体的扫描器,支持使用文件列表来指定扫描主机范围。该软件还可以自定义端口范围,并附带一些常用网络工具。

(6)Nmap:寻找存在漏洞的目标主机。一旦发现有漏洞的目标,就开始对监听端口进行扫描。Nmap 通过使用 TCP 协议栈指纹准确地判断出被扫描主机的操作系统类型。通过使用它,可以让安全管理员了解黑客欲攻击的站点,安全管理员还可以发现自己网站的漏洞,并逐步加以完善。

下面对 Nmap 和 X-Scan 的使用作简单的介绍。

1. Nmap

Nmap 是在免费软件基金会 GPL(GNU general public license)上发布的。Nmap 带有图形终端,语法相当简单,Nmap 的不同选项和-s 标志组成了不同的扫描类型,例如,一个 ping-scan 命令就是“-sP”。在确定了目标主机和网络之后,即可进行扫描。如果以 root 来运行 Nmap,Nmap 的功能会大大地增强,因为超级用户可以创建便于 Nmap 利用的定制数据包。

在目标机上,Nmap 运行灵活。使用 Nmap 进行单机扫描或是整个网络的扫描都很简单,只要将带有“/mask”的目标地址指定给 Nmap 即可。地址是“victim/24”则目标是 C 类网络,地址是“victim/16”则目标是 B 类网络。另外,Nmap 允许使用各类指定的网络地址,例如,“192.168.7.*”指 192.168.7.0/24,“192.168.7.1,4,8~12”可以对所选子网下的多台主机进行扫描。

Nmap 扫描器的使用方法如下。

1) ping 扫描

扫描者(用户、网络管理员或黑客)使用 Nmap 扫描整个网络寻找目标。通过使用-sP 命

令,进行 ping 扫描。默认情况下,Nmap 给每个扫描到的主机发送一个 ICMP Echo 和一个 TCP ACK,主机对其中任何一个的响应都会被 Nmap 得到。例如,扫描 192.168.7.0 网络的命令如下:

```
# nmap -sP 192.168.7.0/24
Starting nmap V. 2.12 by Fyodor (fyodor@dhp.com,www.insecure.org/nmap/)
Host (192.168.7.11) appears to be up.
Host (192.168.7.12) appears to be up.
Host (192.168.7.76) appears to be up.
Nmap run completed -- 256 IP addresses (3 hosts up) scanned in 1 second
```

如果不发送 ICMP Echo 请求,但要检查系统的可用性,ping 扫描可能得不到一些站点的响应。在这种情况下,一个 TCP ping 就可用于扫描目标网络,TCP ping 将发送一个 ACK 到目标网络上的每个主机。网络上的主机如果在线,则会返回一个 TCP RST 响应。使用带有 ping 扫描的 TCP ping 选项,也就是 PT 选项可以对网络上指定端口进行扫描(本文例子中指的缺省端口是 80 号端口),它将可能通过目标边界路由器甚至防火墙。例如,扫描 192.168.7.0 网络指定的 80 端口的命令如下:

```
# nmap -sP -PT 80 192.168.7.0/24
TCP probe port is 80
Starting nmap V. 2.12 by Fyodor (fyodor@dhp.com,www.insecure.org/nmap/)
Host (192.168.7.11) appears to be up.
Host (192.168.7.12) appears to be up.
Host (192.168.7.76) appears to be up.
Nmap run completed -- 256 IP addresses (3 hosts up) scanned in 1 second
```

2)端口扫描

当入侵者发现了在目标网络上运行的主机后,下一步进行的是端口扫描。

Nmap 支持不同类别的端口扫描 TCP 连接,如 TCP SYN、Stealth FIN、Xmas Tree、Null 和 UDP 扫描。

一个攻击者使用 TCP 连接扫描很容易被发现,因为 Nmap 将使用 connect()系统调用打开目标机上相关端口的连接,并完成三次握手,攻击者登录到主机将显示开放的端口。一个 TCP 连接扫描使用-sT 命令如下:

```
# nmap -sT 192.168.7.16
Starting nmap V. 2.12 by Fyodor (fyodor@dhp.com,www.insecure.org/nmap/)
Interesting ports on (192.168.7.16):
Port      State      Protocol    Service
7         open       tcp         echo
9         open       tcp         discard
13        open       tcp         daytime
19        open       tcp         chargen
21        open       tcp         ftp
Nmap run completed -- 1 IP address (1 host up) scanned in 3 seconds
```

3) 隐蔽扫描

如果一个攻击者不愿在扫描时使其信息被记录在目标系统日志上,可使用 TCP SYN 扫描,它很少会在目标机上留下记录,三次握手的过程从来都不会完全实现。TCP SYN 通过发送一个 SYN 包(是 TCP 协议中的第一个包)开始一次 SYN 的扫描。任何开放的端口都将有一个 SYN|ACK 响应。然后,攻击者发送一个 RST 替代 ACK,连接将中止。三次握手得不到实现,也就很少有站点能记录这样的探测。如果是关闭的端口,对最初的 SYN 信号的响应也会是 RST,Nmap 就会知道该端口不在监听。-sS 命令将发送一个 SYN 扫描探测主机或网络:

```
# nmap -sS 192.168.7.9
Starting nmap V. 2.12 by Fyodor (fyodor@dhp.com, www.insecure.org/nmap/)
Interesting ports on saturnlink.nac.net (192.168.7.9):
Port      State      Protocol    Service
21        open       tcp         ftp
25        open       tcp         smtp
53        open       tcp         domain
80        open       tcp         http
...
Nmap run completed -- 1 IP address (1 host up) scanned in 1 second
```

4) UDP 扫描

如果要查出哪些端口在监听则进行 UDP 扫描,这样就可以知道哪些端口对 UDP 是开放的。例如,一个攻击者要寻找一个 rpcbind 或 cDc Back Orifice 漏洞,就可以使用 Nmap 发送一个 0 字节的 UDP 包到每个端口。如果主机返回端口不可达,则表示端口是关闭的。但这种方法受到时间的限制,因为大多数的 UNIX 主机限制 ICMP 错误速率。Nmap 本身会检测这种速率并自身减速,也就不会产生溢出主机的情况。在实际使用中可以用-sU 命令发送一个 UDP 扫描,使用方法如下:

```
# nmap -sU 192.168.7.7
WARNING: -sU is now UDP scan -- for TCP FIN scan use -sF
Starting nmap V. 2.12 by Fyodor (fyodor@dhp.com, www.insecure.org/nmap/)
Interesting ports on saturnlink.nac.net (192.168.7.7):
Port      State      Protocol    Service
53        open       udp         domain
111       open       udp         sunrpc
123       open       udp         ntp
137       open       udp         netbios-ns
138       open       udp         netbios-dgm
177       open       udp         xdmcp
1024      open       udp         unknown
Nmap run completed -- 1 IP address (1 host up) scanned in 2 seconds
```

2. X-Scan

X-Scan 是由 Xfocus(安全焦点)开发的一个功能强大的扫描工具。它采用多线程方式

对指定 IP 地址段(或单机)进行安全漏洞检测,支持插件功能,提供了图形界面和命令行两种操作方式。X-Scan 扫描内容包括:远程服务类型、操作系统类型及版本,各种弱口令漏洞、后门、应用服务漏洞、网络设备漏洞、拒绝服务漏洞等二十几个大类。

X-Scan v3.3 是 X-Scan 的最高版本,它修改了以前版本的一些小 bug,修改“存活主机”插件,并加入 2.3 版本中的 SNMP、NetBIOS 插件,优化了主程序及 NASL 库,更新了攻击测试脚本及中文描述。它可以运行在 Windows NT/XP 上,但在 Windows NT 4.0 系统下无法通过 TCP/IP 堆栈指纹识别远程操作系统类型。

X-Scan 是一款经典的扫描器,它和同类软件相比扫描更加全面且无时间、IP 等限制,所以更适合初学者使用。也可以用它来检查系统的漏洞,使系统的配置更加安全。下面介绍 X-Scan 的具体使用:

(1)首先,打开 X-Scan,启动界面如图 2-3 所示。



图 2-3 X-Scan 启动界面

(2)设置扫描参数。在菜单中选择“设置”,在弹出的窗口中对扫描参数进行设置,如图 2-4 所示。本例中指定 IP 范围为“192.168.1.101”;也可以设置为区段,如“192.168.1.1~192.168.1.255”。



图 2-4 扫描参数设置

(3)全局设置。设置扫描模块,如图 2-5 所示。默认的扫描报告为 HTML 格式。



图 2-5 全局设置

(4)插件设置。分别对要扫描的端口、SNMP 信息、NetBIOS 信息的显示内容以及漏洞检测脚本、CGI 和用户字典进行设置。端口设置如图 2-6 所示。

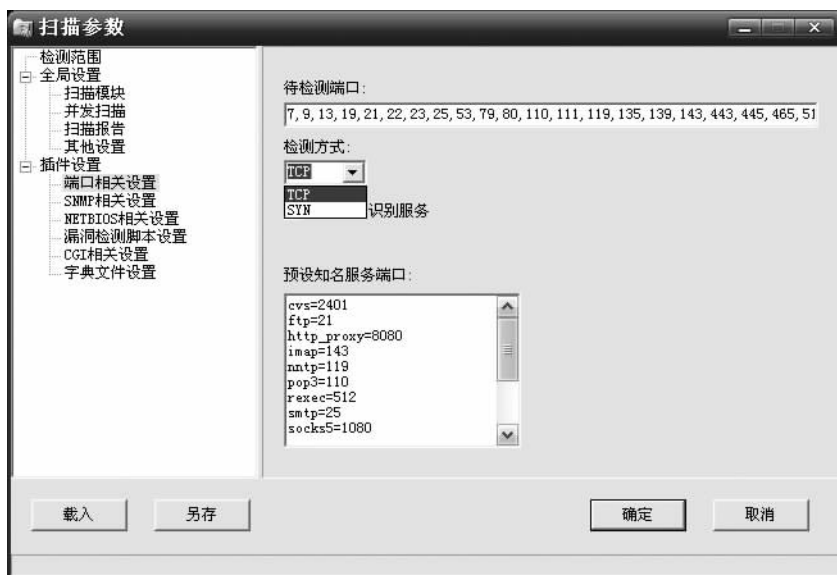


图 2-6 端口设置

(5)开始扫描。设置完成后,单击“确定”按钮,在 X-Scan 主窗口工具栏单击“开始扫描”按钮,X-Scan 会按照设置好的扫描参数进行扫描。扫描过程如图 2-7 所示。



图 2-7 扫描过程

(6)查看扫描结果。扫描完成后在浏览器中会新建页面显示扫描结果。报告中详细地介绍了各个漏洞,并可以链接上 Xfocus 的站点数据库进行查询。扫描结果如图 2-8 所示。

安全漏洞及解决方案: 192.168.1.101		
类型	端口/服务	安全漏洞及解决方案
提示	ftp (21/tcp)	开放服务 "FTP"服务运行于该端口。 BANNER信息: 220 Serv-U FTP Server v6.4 for WinSock ready... NESSUS_ID : 10330
	ftp (21/tcp)	FTP服务的版本和类型 通过登陆目标服务器并经过缓冲器接收可查出FTP服务的类型和版本。这些注册过的标识信息将给予潜在的攻击者们关于他们要攻击的系统的额外信息。版本和类型会在可能的地方被泄露。 解决方案: 将这些注册过的标识信息转变为普通类别的信息。。 风险等级: 低 Remote FTP server banner : 220 Serv-U FTP Server v6.4 for WinSock ready... NESSUS_ID : 10092

图 2-8 扫描结果

2.3 网络监听

对于进行网络攻击的攻击者来说,能攻破网关、路由器、防火墙的情况极为少见。安全管理员完全可以通过安装一些设备对网络进行监控;或者使用一些专门的设备,运行专门的

监听软件来防止任何非法访问。然而,潜入一台不引人注意的计算机中,然后悄悄地运行一个监听程序,黑客是完全可以做到的。

网络监听也称嗅探(sniffer),其目的是截获通信的信息。网络监听是提供给管理员的一种管理工具。使用这种工具,可以监视网络的状态、数据流动情况以及网络上传输的信息,但网络监听也是黑客常用的一种方法。当黑客成功地登录一台网络上的主机,并取得了这台主机的超级用户的权限之后,往往要扩大其战果,尝试夺取网络中其他主机的控制权,网络监听就是一种简单且有效的方法,它常常能轻易地获得用其他方法很难获得的信息。

2.3.1 网络监听概述

对于一台联网的计算机,最方便的是在以太网中进行监听,只需安装一个监听软件,然后就可以在机器旁浏览监听到的信息了。以太网协议将要发送的数据包发往连在一起的所有主机,包头中包含着应该接收数据包的主机的正确地址。因此,只有与数据包中目标地址一致的主机才能接收数据包。但是,如果主机工作在监听模式下,则无论数据包中的目标物理地址是什么,主机都将接收。下面就以以太网为例讲述在局域网中进行监听的原理和过程。

1. 网络监听基本原理

在 Internet 上,有许多局域网是由几台甚至十几台主机通过一条电缆或一个集线器连在一起的。在协议的高层或用户看来,当同一网络中的两台主机通信时,源主机将写有目的主机 IP 地址的数据包发向网关。但是,这种数据包并不能在协议栈的高层被直接发送出去,要发送的数据包必须从 TCP/IP 协议的 IP 层交给网络接口,即数据链路层。而网卡不能识别 IP 地址,这是由于带有 IP 地址的数据包又增加了一部分信息——以太帧的帧头。在帧头中,有两个域分别为 48 位的地址源主机和目的主机的物理地址,该地址只有网络接口才能识别。这个 48 位的地址是与 IP 地址对应的。也就是说,一个 IP 地址必然对应一个物理地址。对于作为网关的主机,由于它连接了多个网络,因此它同时具有多个 IP 地址(在每个网络中都有一个)。

在以太网中,填写了物理地址的帧从网卡中发送出去,传送到物理线路上。如果局域网由电缆连接而成,则数据包在电缆上传输,数据包的信号能够到达线路上的每一台主机。当使用集线器时,发送出去的数据包到达集线器,由集线器再发往连接在集线器上的每一条线路。在物理线路上传输的数据包的信号也能到达连接在集线器上的每一台主机。局域网的这种工作方式,可以用一个形象的例子来描述。大房间就像是一个共享的信道,目前绝大多数计算机网络使用共享的通信信道;里面的每个人好像是一台主机;人们所说的话是数据包,在大房间中到处传播。当对其中某个人说话时,所有的人都能听到,但只有名字相同的那个人,才会对这些话语作出反应和进行处理。

2. 网络监听的实现

当数据包到达一台主机的网卡时,在正常情况下,网卡读入数据帧并进行检查。如果数据帧中携带的物理地址是自己的,或者物理地址是广播地址,则将数据帧交给上层协议软件,也就是 IP 层软件,否则就将这个帧丢弃。对于每一个到达网卡的数据帧,都要进行这个过程。一般而言,网卡有几种接收数据帧的状态,如 unicast、broadcast、multicast、promiscuous 等。unicast 是指网卡在工作时接收目的地址是本机硬件地址的数据帧;broadcast 是指接收所有类型为广播报文的数据帧;multicast 是指接收特定的组播报文;promiscuous 则是通常

说的监听模式,是指对报文中的目的硬件地址不加任何检查,全部接收的工作模式。可以看到,正常的网卡应该只是接收发往自身的数据报文、广播和组播报文。当在promiscuous模式下时,在同一条物理信道上传输的所有信息都可以被接收到。当信息在网络中进行传播时,利用监听工具将网卡设置为监听模式,便可以源源不断地将网上传输的信息截获,然后进行攻击。

在 UNIX 系统中,设置网卡的监听状态需要超级用户的权限,这一点限制了在 UNIX 系统中的普通用户是不能进行网络监听的。只有获得超级用户权限,才能进行网络监听。但是在 Windows XP 中没有这个限制,只要运行这一类的监听软件即可。同时,在计算机上运行的这类软件具有操作方便、监听信息的综合能力强等特点。

当连接在同一条电缆或集线器上的主机被逻辑地分为几个子网时,如果一台主机处于监听模式下,通过使用不同的掩码、IP 地址和网关,它还能接收到发往与自己不在同一子网的主机的数据包。也就是说,在同一条物理信道上传输的所有信息都可以被接收到。但不能监听不在同一个网段计算机传输的信息。即一台计算机只能监听经过自己网段的数据包。如果能监听到整个 Internet,情形将会非常可怕。

从上面的讨论中可以了解到,通信信道的共享意味着计算机有可能接收到发往另一台计算机的信息。另外,要说明的是,Internet 中使用的大部分协议都是很早设计的,许多协议的实现都是基于一种非常友好的、通信双方充分信任的基础之上的。因此,直到现在网络安全还是非常脆弱的。在通常的网络环境下,用户的所有信息,包括用户名和口令信息都是以明文的方式在网上传输的。因此,对于一个网络黑客或网络攻击者来说,只要他们具有初步的网络和 TCP/IP 协议知识,便能轻易地从监听到的信息中提取出感兴趣的部分。

2.3.2 网络监听工具

网络监听工具非常多。运行在 Windows 平台上的有 Windump、Iris、Sniffer Pro、Win Sniffer、Pswmonitor 等;运行在 UNIX 平台上的有 TCPdump、Snort、Dsniff、Sniffit 等。

运行在 Windows 平台上的 Sniffer Pro 可以监听到网上传输的所有数据。网络可以是运行在各种协议之下的,包括以太网、TCP/IP、ZPX 等,也可以是集中协议的联合体系。

安装 Sniffer 程序的主机和被监听的主机必须在同一个以太网网段上,因此在外部主机上运行 Sniffer 是没有效果的。使用者必须以 Root 的身份使用 Sniffer 程序,才能够监听到以太网网段上的数据流。Sniffer 通常运行在路由器或有路由器功能的主机上。这样就能对大量的数据进行监听。可以借助 Sniffer 分析网络的流量,找出所关心的网络中潜在的问题。例如,假设网络的某一段运行得不是很好,报文的发送比较慢,而又不知道问题出在什么地方,此时就可以用 Sniffer 作出准确的判断。

下面以 Sniffer Pro 为例,介绍网络监听软件的使用:

(1)打开 Sniffer Pro,界面如图 2-9 所示。

(2)在运行程序之前,需要为 Sniffer Pro 设置用来监听的网络接口,如图 2-10 所示。

(3)设置捕获条件,如图 2-11 所示。基本的捕获条件有两种:

- 链路层捕获,按源 MAC 地址和目的 MAC 地址进行捕获,输入方式为十六进制连续输入,如 00E0FC123456。
- IP 层捕获,按源 IP 和目的 IP 进行捕获。输入方式为点间隔方式,如 10.107.1.1。如果选择 IP 层捕获,则 ARP 等报文将被过滤掉。



图 2-9 Sniffer Pro 运行界面



图 2-10 设置监听的网络接口



图 2-11 设置捕获条件

(4)开始捕获。捕获条件设置完之后,在工具栏上单击“开始捕获”按钮,Sniffer Pro 就会按照过滤器中定义好的设置进行监听,如图 2-12 所示。



图 2-12 Sniffer Pro 监听界面

2.4 木马攻击

计算机世界中特洛伊木马的名字来自著名的“木马屠城记”。传说古希腊大军围攻特洛伊城,有人献计制造一只大木马,让士兵藏匿其中,于是部队里应外合焚毁了特洛伊城。后世就称这只木马为“特洛伊木马”。如今,黑客程序借用此名,有“里应外合”之意。

2.4.1 木马攻击原理

计算机世界的特洛伊木马(Trojan Horse)是指隐藏在正常程序中的一段具有特殊功能的恶意代码,是具备破坏和删除文件、发送密码、记录键盘和拒绝服务攻击等特殊功能的后门程序。

木马程序是未经授权的,它一般包含在合法用户的程序中。木马程序是常常把有预谋的功能藏在公开的功能之中,掩盖其真实企图的程序。这个程序表面上看是完成某一功能,如登录、编辑或游戏等,而实际上完成的却是其他操作,如删除文件、窃取口令或格式化磁盘等。一般来说,特洛伊木马是在合法用户的程序运行时悄悄地进行非法操作,而且一般不易被察觉。因此,它是一种极为危险的攻击手段。

完整的木马程序同远程控制软件一样,都是一个通过端口进行通信的网络客户机/服务器程序。一般由两个部分组成:一个是服务器程序(服务器端),另一个是控制器程序(客户端)。服务器端安装在被控制的计算机中,它一般通过电子邮件或其他手段让用户在其安装的计算机中运行,以达到控制该用户计算机的目的。客户端程序是控制者所使用的,用于对受控的计算机进行控制。服务器端程序和客户端程序建立连接后就可以实现对远程计算机的控制。

木马程序运行时,其服务器端程序获得本地计算机的最高操作权限,当本地计算机连入

网络后,客户端程序可以与服务器端程序直接建立连接,并可以向服务器端程序发送各种基本的操作请求,并由服务器端程序完成这些请求,也就实现对本地计算机的控制了。木马发挥作用必须要求服务器端程序和客户端程序同时存在,所以必须要求本地计算机上有服务器端程序。服务器端程序是可执行程序,可以直接传播,也可以隐含在其他的可执行程序中传播,但木马本身不具备繁殖性和自动感染的功能。

木马程序不是计算机病毒,但越来越多的新版杀毒软件已可以查杀一些木马程序了,所以也有不少人称木马程序为木马病毒,不少杀毒软件直接提供检测和清除木马程序的功能。不过,为区别传统意义上的计算机病毒和木马程序,还是将木马程序单独列为一类比较合适。

2.4.2 木马的隐藏

世界上第一个计算机木马是出现在 1986 年的 PC-Write 木马。它伪装成共享软件 PC-Write 的 2.72 版本(编写 PC-Write 的 Quicksort 公司从未发行过 2.72 版本)。一旦用户信以为真运行该木马程序,就会导致运行该程序的计算机硬盘被格式化。此时的第一代木马还不具备传染特征。

1989 年出现了 AIDS 木马。由于当时很少有人使用电子邮件,所以 AIDS 的编写者就利用现实生活中的邮件进行散播:给其他人寄去一封封含有木马程序软盘的邮件。之所以叫这个名称,是因为软盘中包含有治疗 AIDS 和 HIV 疾病的药品和价格以及这两种疾病的预防措施等相关信息。软盘中的木马程序在运行后,虽然不会破坏数据,但是会将硬盘加密锁死,然后提示受感染用户花钱消灾。可以说第二代木马已具备了传播特征,尽管只是通过传统的邮件方式。

随着 Internet 的普及,新一代的木马出现了,它兼备伪装和传播两种特征并结合 TCP/IP 网络技术四处泛滥。木马的主要目标也不再是进行文件和系统的破坏,而是带有收集密码、远程控制等目的。这段时期比较有名的木马有国外的 BO 2000(Back Orifice)和国内的冰河木马。它们的共同特点是:基于网络的客户端/服务器应用程序,具有搜集信息、执行系统命令、重新设置机器、重新定向等功能。

在木马的发展历史中,不管是哪一种类型的木马攻击,都是将服务端即木马程序隐藏在目标主机中,控制端的控制程序通过木马端口用木马进行远程攻击。将木马程序隐藏在目标主机中是木马攻击的最为关键的步骤。木马的隐藏大致有以下几种常见形式。

1. 集成到程序中

木马程序为了不被用户轻易地删除,常常集成到程序里。一旦用户激活木马程序,那么木马文件就能和某一应用程序捆绑在一起,然后上传到服务端覆盖原文件,这样即使木马被删除了,只要运行捆绑了木马的应用程序,木马就会被安装上去。如果绑定到系统文件,那么每一次系统启动均会启动木马。

2. 隐藏在配置文件中

普通用户平时使用最多的是图形化界面的操作系统,对于配置文件大多数是不闻不问,这给木马提供了一个藏身之处。利用配置文件的特殊作用,木马很容易就能在用户的计算机中运行、发作,从而偷窥或者监视使用者。不过,这种方式不是很隐蔽,容易被发现,所以早期的在 autoexec.bat 和 config.sys 中加载木马程序的现象并不多见了,很多木马程序又另外寻找其他的配置文件进行隐藏。

3. 潜伏在 win.ini 中

木马要想达到控制或者监视计算机的目的,必须要运行,然而很少有人在自己的计算机中运行木马。木马必须找一个既安全又能在系统启动时自动运行的地方,于是木马就潜伏在 win.ini 中,因为通过 win.ini 文件可以在 Windows 系统中启动一些服务程序。一般情况下,在 win.ini 文件中的 [WINDOWS] 下面的 run 命令和 load 命令中的等号后面什么都没有。如果发现后面跟的路径和文件名不是常见的启动文件,那么这台计算机就可能中木马了。除此之外,一些木马(如 AOL Trojan 木马)常伪装成 command.exe 文件。如果不注意,可能不会发现它不是真正的系统启动文件。例如:

```
run=c:\Windows\file.exe
load=c:\Windows\file.exe           //这个 file.exe 很可能是木马
```

4. 伪装在普通文件中

这种方法出现得比较晚,但是现在很流行。对于不熟练的 Windows 操作者,非常容易受到欺骗。具体方法是把可执行文件伪装成图片或文本,在程序中把图标改成 Windows 的默认图片图标,然后把文件名改为“*.jpg.exe”,由于 Windows 默认设置是“不显示已知的文件后缀名”,文件将会显示为“*.jpg”,甚至有些在程序中嵌一张图片,一旦用户点击图标就中木马了。

5. 内置到注册表中

上面介绍的隐藏方式是木马隐藏的早期形式,容易被发现。被攻击者也渐渐地发现了木马隐藏的伎俩。木马程序的开发者认为必须要找到一个不容易被人发现的地方,如注册表中。注册表本身比较复杂,隐藏在其中很难被发现,于是就有了注册表木马。例如:

```
HKEY_LOCAL_MACHINE\Software\Microsoft\Windows\CurrentVersion
```

下所有以“run”开头的键值;

```
HKEY_CURRENT_USER\Software\Microsoft\Windows\CurrentVersion
```

下所有以“run”开头的键值;

```
HKEY_USERS\.DEFAULT\Software\Microsoft\Windows\CurrentVersion
```

下所有以“run”开头的键值。

6. 在 system.ini 中藏身

Windows 安装目录下的 system.ini 也是木马喜欢隐藏的地方。在该文件的 [boot] 字段中如果有这样的内容“shell=Explorer.exe\file.exe”,说明计算机中已经隐藏了木马,因为这里的 file.exe 就是木马服务端程序。另外,检查在 system.ini 中的 [386enh] 字段内的“driver=路径\程序名”,这里也有可能被木马所利用。还有,在 system.ini 中的 [mic]、[drivers]、[drivers32] 三个字段可以起到加载驱动程序的作用,但也是增添木马程序的好场所。

7. 隐藏于启动组中

木马更关注的是能否自动加载到系统中。因为一旦木马加载到系统中,不管用户用什么方法都无法将木马删除。因此按照这个逻辑,启动组也是木马可以藏身的好地方,因为这里的确是自动加载运行的好场所,所以要注意经常检查启动组。

启动组对应的文件夹为:

系统盘\Windows\startmenu\programs\startup

在注册表中的位置为:

```
HKEY_CURRENT_USER\Software\Microsoft\Windows\CurrentVersion\Explorer\Shell-  
FoldersStartup="C:\Windows\startmenu\programs\startup"
```

8. 隐藏在 winstart.bat 中

凡是利于木马自动加载的地方,木马都有可能隐藏其中。winstart.bat 也是一个能自动被 Windows 加载运行的文件,它多数情况下为应用程序及 Windows 自动生成,在执行了 win.com 并加载了多数驱动程序之后开始执行。由于 autoexec.bat 的功能可以由 winstart.bat 代替完成,因此,木马完全可以像在 autoexec.bat 中那样被加载运行。

9. 捆绑在启动文件中

木马还可以存在于应用程序的启动配置文件中。控制端利用这些文件能启动程序的特点,将制作好的带有木马启动命令的同名文件上传到服务端覆盖同名文件,这样就可以达到启动木马的目的。

10. 设置在超级链接中

木马的主人在网页上放置恶意代码,引诱用户点击,一旦用户点击就有可能将木马引入自己的系统中。对于网页上的链接,浏览者要谨慎点击,对于陌生和非法的网站要尽量远离。

2.4.3 木马的清除与防范

目前,由于流行的各种热门网站、客户端软件和浏览器都存在着众多漏洞和安全薄弱点,使得用户遭到攻击的渠道剧增;支撑互联网发展的多种商业模式也都遭到了盗号木马、木马点击器的侵袭,使得用户对于网络购物、网络支付、网游产业安全的信心遭到打击。

瑞星“云安全”系统提供的数据表明,2009年1月至3月,互联网上出现的木马网页累计达1.9亿多个,平均每天有889万余网民访问这些网页,累计有8亿网民遭木马攻击。大型网站、浏览器和流行软件成为黑客窥测的对象,一季度有24 202个大型网站被植入木马,这已经成为威胁国内互联网安全的最主要因素之一。

面对存在如此众多木马攻击的 Internet,网络用户应该时刻保持警惕,在使用过程中要及时清除木马和防范木马的攻击。

1. 木马的清除

从木马的隐藏可以看出,清除木马与清除病毒非常相似。首先是定位,即在木马经常隐藏的地方发现木马的存在,然后清除。如果在定位时完全确定并找到了被感染木马的文件,则只需替换被感染的文件即可;如果在定位时完全确定有木马,但没有找到感染的文件,就要重新恢复系统。清除木马时可以采用手工和工具相结合的方式,这样可以加强清除的效果。

2. 木马的防范

经常检测并及时清除并不是最有效的方法,事实上最重要的是防范。木马程序的攻击能力很强,但是其攻击的实现还是依赖于在目标计算机中运行的服务器端,如果本机的安全措施到位,没有给服务器端植入的机会,那么木马就无法完成攻击。但是,由于网络环境比较复杂,因此,为避免造成木马侵入,多作防范是保证网络安全的首要工作。防范木马的基本策略有:

(1)提高防范意识。木马程序传播时一般是通过电子邮件或文件下载的形式,因此对陌生人的邮件或带有附件的邮件要格外小心。建议在没安装杀毒软件的情况下,不要打开任

何附件,不要执行任何来历不明的程序。不要随意下载来历不明的软件,最好是在一些知名的网站下载软件,在安装软件时最好先用杀毒软件检测,确定没有病毒和木马之后再进行安装。

(2)使用杀毒软件或木马专杀工具。现在国内的杀毒软件都推出了清除某些木马的功能,如瑞星,这些杀毒软件可以不定期地在脱机的情况下进行检查和清除木马的工作。另外,有的杀毒软件还提供网络监控功能,这一功能可以在黑客从远程端执行用户机器中的文件时,提供报警或让执行失败,使黑客在向用户机器上传可执行文件后无法正确执行这些文件,从而避免了进一步的损失。但是,杀毒软件不是万能的,而木马专杀工具在这方面却可以做得更好。

(3)观察系统异常,及时断开网络。一旦感觉有被木马攻击的迹象,在没清除前,建议断开网络。例如,发现有些软件正使用不常见的端口(一般大于1024)与用户通信时,这一端口很可能就是木马的通信端口;在c:\windows或c:\windows\system下,如果存在只有文件名却没有图标的可执行程序或是注册表被修改等,也有可能是遭到了木马的攻击。当发现可疑迹象后,应立即断开网络,然后对硬盘进行认真的检查,确认其中是否有木马。

(4)及时修补漏洞并关闭可疑端口。一般木马都是通过系统漏洞在系统中打开端口、留下后门,以便上传木马程序和可执行代码。因此,在把漏洞修补上的同时,需要对端口进行检查,把可疑端口关闭。

(5)运行实时监控程序。在上网时最好运行反木马实时监控程序和个人防火墙,并定时对系统进行病毒检查。经常升级系统和更新病毒库,经常关注厂商网站的安全公告,这些网站通常都会及时地将系统漏洞、木马和更新情况公布出来,并在第一时间发布系统补丁和新的病毒库等。

目前,国内安全厂商提出了云安全的概念,通过在客户端监控网络中软件行为的异常,将发现的疑似木马的最新信息送到服务器进行分析和处理,然后再把木马的解决方案分发到客户端。云安全缩短了样本的发现时间和响应时间,为未知木马的防范开辟了新的思路。

2.5 拒绝服务攻击

2007年4月27日,爱沙尼亚政府下令拆除了位于首都塔林市中心的苏军解放塔林纪念碑和苏联战士公墓中的铜制苏联红军雕像而引发骚乱,俄罗斯对此反应强烈。从4月27日起,爱沙尼亚总统和议会的官方网站、政府各大部门网站、政党网站的访问量激增,服务器由于过于拥挤而陷于瘫痪状态。全国6大新闻机构中有3家遭到攻击,此外还有两家全国最大的银行和多家从事通信业务的公司网站纷纷中招。为了安全起见,受到攻击的网站曾一度完全关闭。爱沙尼亚的网络安全专家认为这次空前的“网络战”是一次典型的分布式拒绝服务攻击。

2009年1月6日中午开始,大量网民在访问论坛时出现故障,无法发帖回帖,或者页面出现大片空白,不止一两家论坛,而是绝大多数论坛网站都出现了这一问题。此外,网易、淘宝等一系列知名网站,包括目前在公司白领中最流行的开心网也被株连,页面频频显示出错,甚至无法打开,中国互联网出现了罕见的“集体休克”现象。以上网站受到影响的程度不同,有的是部分显示错误或功能无法实现,有的则为完全无法使用。众多网友开始质疑,是

谁导演了这次网站“集体休克”事件。

2009年7月18日上午10时,上海市7月私车额度投标拍卖会照常开始。通常,在第一阶段的最后5分钟,是人流大量上升的时段,正常情况下这一时段投标人数在8 000人左右,但18日该时段仅有361人。大约11时10分左右,竞拍界面上突然出现了一行红字:“因网络原因,本月私车额度投标拍卖会取消。何时进行另行公告。”当晚,上海市政府新闻办表示此次事件是由于10时55分以后竞拍流量比以往正常流量猛增10倍,相当于有几十万人次同时登录,且均为大量异常请求,致使系统无法正常运行造成的。警方确认在7月18日私车额度投标拍卖会期间,有明确针对系统的拒绝服务攻击行为。

以上都是拒绝服务攻击造成的网络安全事件。

2.5.1 拒绝服务攻击概述

拒绝服务简称DoS(denial of service)。DoS是一种最常见的攻击形式,这种攻击是指一个用户占据了大量的共享资源,使系统没有剩余的资源给其他用户使用,即拒绝服务,造成DoS的攻击行为被称为DoS攻击。这种攻击的结果是使系统效率降低或者失去服务能力。

DoS攻击降低了资源的可用性,这些资源可以是磁盘空间、CPU使用的时间、打印机、调制解调器等,目的是使计算机或网络无法提供正常的服务。最常见的DoS攻击有计算机网络带宽攻击和连通性攻击。带宽攻击指以极大的通信量冲击网络,使得所有可用网络资源都被消耗殆尽,最后导致合法的用户请求无法通过;连通性攻击指用大量的连接请求冲击计算机,使得所有可用的操作系统资源都被消耗殆尽,最终计算机无法再处理合法用户的请求。

DoS攻击主要是由以下两种情况引起:

(1)由程序员所编程序的错误造成。由于程序员对程序错误的编制,导致系统不停地建立进程,最终耗尽资源,只能重新启动机器。不同的系统平台都会采取某些方法防止一些特殊的用户占用过多的系统资源,建议尽量采用资源管理的方式降低这种安全威胁。

(2)由磁盘存储空间引起。假如一个用户有权利存储大量的文件,那么他可能只为系统留下很小的空间用来存储日志文件等系统信息。这种不良的操作习惯会给系统留下隐患。此时,应对系统配额作出考虑。

DoS侧重于通过对主机特定漏洞的攻击导致网络栈失效、系统崩溃、主机死机而无法提供正常的网络服务功能,从而造成拒绝服务。常见的DoS攻击手段有:

1. 死亡之 ping

由于在早期的阶段,路由器对所传输的文件包最大尺寸都有限制,许多操作系统对TCP/IP的实现在ICMP包上都是规定64 KB,并且在对包的标题头进行读取之后,要根据该标题头里包含的信息来为有效负载生成缓冲区,一旦产生畸形即声称自己的尺寸超过ICMP的上限包,也就是加载的尺寸超过上限64 KB时,就会出现内存分配错误,导致TCP/IP堆栈崩溃,致使接收方死机。这种攻击方式主要是针对Windows 9x操作系统的,而UNIX、Linux、Solaris、Mac OS都具有抵抗一般死亡之ping攻击的能力。

2. UDP Flood(UDP 洪泛)攻击

UDP洪泛由ECHO/CHARGEN服务引起。ECHO/CHARGEN服务是TCP/IP协议为TCP和UDP提供的一种服务。ECHO的作用就是把接收端将接收到的数据内容返回到

发送端,CHARGEN 则随机返回字符。这样简单的功能为网络管理员提供了进行可达性测试、协议软件测试和选路识别的重要工具,也为黑客进行“洪水”攻击提供了方便。当攻击者假冒一台主机向另一台主机的服务端口发送数据时,ECHO 服务或 CHARGEN 服务就会自动回复。两台机器之间的互相回送会形成大量数据包。多台主机之间相互产生回送数据包,最终会导致系统瘫痪。

3. SYN Flood(洪泛)攻击

一些 TCP/IP 堆栈的实现只能等待从有限数量的计算机发来的 ACK 消息,因为它们只有有限数量的内存缓冲区用于创建连接,如果这些缓冲区内充满了虚假连接的初始信息,该服务器就会对接下来的连接停止响应,直到缓冲区里的连接超时。在一些创建连接不受限制的案例中,SYN Flood 攻击具有类似的影响。

未来的 SYN Flood 攻击令人担忧,这是由于释放洪水者并不寻求响应,所以无法从一个简单高容量的传输中鉴别出来。

4. Land 攻击

Land 攻击利用 TCP/IP 三次握手的协议缺陷进行攻击。但它不是依靠伪造的地址,而是先发出一个特殊的 SYN 数据包,包中的源地址和目标地址都是目标主机。这样,就会让目标主机向自己回以 SYN/ACK 包,导致又给自己回一个 ACK,并建立自己与自己的连接。大量这样的无效连接达到一定数量,将会拒绝新的连接请求。在 Land 攻击中,每一个这样的连接都将保留直到超时。各种系统对 Land 攻击反应不同,UNIX 系统会崩溃,而 Windows NT 会变得极其缓慢。

5. Smurf 攻击

一个简单的 Smurf 攻击将回复地址设置成受害网络的广播地址,进而所有的 ICMP 应答请求 ping 数据包淹没受害主机,最终导致该网络的所有主机都对此 ICMP 应答请求作出答复,从而导致网络阻塞,所以它比死亡之 ping 洪水的流量高出一或两个数量级。更加复杂的 Smurf 将源地址改为第三方的受害者地址,最终导致第三方崩溃。

Smurf 攻击屡屡奏效,并非受害主机的操作系统或者网络协议有什么问题,而在于这种攻击的“借力”效果。黑客可以利用自己有限的带宽,同时找到一个高带宽的网络作为反弹攻击跳板,转而攻击一个有限带宽的受害系统。最终的结果就好比一个宽敞的高速公路,却有一个极狭窄的出口,可是所有车辆都必须从这个出口通行,从而导致交通堵塞。更严重的是,这种攻击不仅能使被害主机拒绝服务,还能使该主机所在网络、所有与外界网络的通信都受到影响。

6. 电子邮件炸弹

电子邮件炸弹是黑客常用的一种攻击手段。邮件炸弹指的是邮件发送者利用特殊的电子邮件软件,在很短的时间内连续不断地将邮件邮寄给同一个收信人,在数以千万计的大容量信件的袭击下收件箱肯定不堪重负,而最终使邮件系统崩溃。

邮件炸弹的实质是发送地址不详且容量庞大的恶意邮件。由于每个人的邮箱容量是有限的,当数量庞大的垃圾邮件到达邮箱的时候,就会把正常的邮件冲掉。同时,由于它占用了大量的网络资源,常常导致网络堵塞,所以使大量的用户不能正常工作。它干扰了用户的正常电子邮件系统,影响邮件系统所在的服务器系统的安全,造成整个网络系统瘫痪。同时也大量消耗了网络资源,导致网络拥塞,使大量用户不能正常上网工作。

电子邮件炸弹是最古老的匿名攻击之一,攻击者能够耗尽接收者网络的带宽。由于这

种攻击方式简单易用,也有很多发匿名邮件的工具,而且只要获悉对方的电子邮件地址就可以进行攻击,所以这是最值得大家防范的一个攻击手段。

7. 畸形消息攻击

目前,Windows、UNIX、Linux 等各类操作系统上的许多服务都存在安全隐患问题,由于这些服务在处理信息之前没有进行适当正确的错误校验,所以一旦收到畸形的信息就有可能崩溃。

8. 泪滴(teardrop)攻击

泪滴攻击利用早期某些操作系统中 TCP/IP 协议栈对 IP 分片包进行重组时的漏洞进行攻击。Windows 3.1/95/NT 以及 Linux 2.1.63 以前的版本都存在这个问题。攻击者向目标主机发送两个分片的 IP 包。第一个 IP 包的数据偏移设为 0,有效数据长度为 N 。第二个 IP 包的数据偏移设为 $K(K < N)$,有效数据长度为 $S(K + S < N)$ 。操作系统需要将分片的 IP 包组合成一个完整的 IP 包,IP 分片含有指示该分段包含的是原包的哪一段的信息。重组的时候会出现一些重叠现象,需要对此进行处理。重组时大量的重叠现象消耗了系统的可用资源导致拒绝服务。

尽管网络安全专家都在着力开发阻止 DoS 攻击的设备,但收效不大,因为 DoS 攻击利用了 TCP 协议本身的弱点。

2.5.2 分布式拒绝服务攻击

分布式拒绝服务 DDoS 攻击指借助于客户机/服务器技术,将多个计算机联合起来作为攻击平台,对一个或多个目标发动 DoS 攻击,从而成倍地提高拒绝服务攻击的威力。

1. DDoS 与 DoS

DDoS 是在传统的 DoS 攻击基础之上产生的一类攻击方式。单一的 DoS 攻击一般是采用一对一的方式,当攻击目标的 CPU 速度、内存或者网络带宽等各项性能指标不高时,它的效果是明显的。随着计算机与网络技术的发展,计算机的处理能力迅速增强,内存大大增加,同时也出现了千兆级别的网络,这使得 DoS 攻击的困难程度加大了。目标主机对恶意攻击包的“消化能力”大大加强。例如,攻击软件每秒可以发送 3 000 个攻击包,但如果主机与网络带宽每秒钟可以处理 10 000 个攻击包,这样一来攻击就不会产生什么效果。

这时,分布式拒绝服务攻击手段就应运而生了。如果计算机与网络的处理能力加大了 10 倍,用一台攻击机来攻击不再起作用,攻击者可以使用 100 台攻击机同时攻击,甚至用 1 000 台同时攻击。DDoS 攻击就是利用更多的傀儡机来发起进攻,以比从前更大的规模来进攻目标主机。

DDoS 的攻击策略侧重于通过很多傀儡主机向受害主机发送大量看似合法的网络包,从而造成网络阻塞或服务器资源耗尽而导致拒绝服务,其攻击一旦被实施,攻击网络包就会犹如洪水般涌向受害主机,从而把合法用户的网络包淹没,导致合法用户无法正常访问服务器的网络资源。因此,又被称为洪水式攻击。

2. DDoS 攻击的原理

通常,攻击者通过木马将 DDoS 主控程序安装在 Internet 上的许多计算机上,在一个设定的时间内,主控程序与大量代理程序通信,代理程序收到指令时就发动攻击。利用客户机/服务器技术,主控程序能在几秒内激活成百上千次代理程序的运行。这些计算机就如同传说中的“僵尸”一样被黑客所控制,有了大量“僵尸”或者称为傀儡机的计算机,发起 DDoS

攻击将是一件轻而易举的事。

高速广泛连接的网络给大家带来了方便,也为 DDoS 攻击创造了极为有利的条件。在低速网络时代,黑客占领攻击用的傀儡机时,总是会优先考虑离目标网络距离近的机器,因为经过路由器的跳数少,效果好。而现在电信骨干结点之间的连接带宽都是以 Gb/s 为级别的,大城市之间的连接带宽更可以达到 2.5 Gb/s,这使得攻击可以从更远的地方或者其他城市发起,攻击者的傀儡机位置可以分布在更大的范围,选择起来更灵活了。

3. DDoS 攻击的表现形式

DDoS 攻击的表现形式主要有两种:一种为流量攻击,主要是针对网络带宽的攻击,即大量攻击包导致网络带宽被阻塞,合法网络包被虚假的攻击包淹没而无法到达主机;另一种为资源耗尽攻击,主要是针对服务器主机的攻击,即通过大量攻击包导致主机的内存被耗尽或 CPU 被内核及应用程序占完而造成无法提供网络服务。

判断网站是否遭受了流量攻击,可通过 ping 命令来测试。若发现 ping 超时或丢包严重,则可能遭受了流量攻击,此时若发现和主机连接在同一交换机上的服务器也访问不了了,假如可以排除网络故障因素的话,则肯定是遭受了流量攻击;还有一个流量攻击的典型现象是,一旦遭受流量攻击,会发现用远程终端连接网站服务器失败。

相对于流量攻击而言,资源耗尽攻击要容易判断一些。假如平时 ping 网站主机和访问网站都是正常的,突然发现网站访问非常缓慢或无法访问了,而 ping 还可以连通,则很可能遭受了资源耗尽攻击。还有一种属于资源耗尽攻击的现象是,ping 自己的网站主机 ping 不通或者是丢包严重,而 ping 与自己的主机在同一交换机上的服务器则正常,造成这种现象的原因是,网站主机遭受攻击后导致系统内核或某些应用程序对 CPU 的占用率达到 100%而无法回应 ping 命令,其实带宽还是有的,否则就 ping 不通接在同一交换机上的主机了。

当受到 DDoS 攻击时还可能伴随以下现象:

- 被攻击主机上有大量等待的 TCP 连接。
- 网络中充斥着大量的无用的数据包,源地址为假。
- 制造高流量无用数据,造成网络拥塞,使受害主机无法正常和外界通信。
- 利用受害主机提供的服务或传输协议上的缺陷,反复高速地发出特定的服务请求,使受害主机无法及时处理所有正常请求。
- 严重时会造成系统死机。

防止 DDoS 攻击是一个系统工程,仅仅依靠某种系统或产品防范 DDoS 是不现实的,甚至可以说完全杜绝 DDoS 目前是不可能的。但通过适当的措施抵御 90%的 DDoS 攻击是可以做到的。例如,增强操作系统的 TCP/IP 栈、升级主机服务器硬件、安装专业抗 DDoS 攻击防火墙等,都可以收到很好的效果。攻击和防御都有成本开销,所以通过适当的办法增强抵御 DDoS 攻击的能力,也就意味着加大了攻击者的攻击成本,那么绝大多数攻击者将无法继续下去而放弃,也就相当于成功地抵御了 DDoS 攻击。

2.6 缓冲区溢出

近些年来,以缓冲区溢出作为远程攻击类型的黑客入侵事件频繁发生。这种攻击可以

使一个匿名的 Internet 用户有机会获得一台主机的部分或全部的控制权,是一种十分危险的攻击。

2.6.1 缓冲区溢出概述

如果把两升的水注入容量为一升的容器中,水会溢出。同样的道理,在计算机内部,输入数据通常被存放在一个临时空间内,这个临时存放空间被称为缓冲区。如果向一个容量有限的内存空间里存储过量数据,这时数据也会溢出存储空间。

1. 缓冲区

缓冲区是程序运行时在内存中为保存特定类型的数据而开辟的一个连续空间,可以保存相同数据类型多个实例。以 C 语言为例,在用 C 语言编写程序的过程中通常会用到数组、指针等,这些空间同 C 语言中所有的变量一样,可以被声明为静态或动态的。静态变量在程序加载时定位于数据段。动态变量在程序运行时定位于堆栈之中,操作系统所使用的缓冲区又被称为堆栈。在各个操作进程之间,指令被临时存储在堆栈中。堆栈是一块保存数据的连续内存,堆栈中的数据具有一个特性:最后一个放入堆栈中的数据总是被最先拿出来。这个特性通常称为后进先出(LIFO)。堆栈中定义了一些操作,最重要的两种操作是 PUSH 和 POP。PUSH 操作在堆栈的顶部追加一个元素,并将堆栈的大小加 1;POP 操作在堆栈顶部移去一个元素,并将堆栈的大小减 1。一个名为堆栈指针(SP)的寄存器指向堆栈的顶部。堆栈的底部在一个固定的地址。堆栈的大小在运行时由内核动态地调整。CPU 实现指令 PUSH 和 POP,向堆栈中添加元素和从中移去元素。在使用 C 语言编写程序时,最重要的技术是利用过程(procedure)和函数(function)来控制程序的运行流程,跳转(jump)也可以。这种程序流程控制的实现要依赖于上面讲到的堆栈。堆栈同时还用于给函数中使用的局部变量动态分配空间以及给函数传递参数和函数返回值。

2. 缓冲区溢出原理

缓冲区溢出是指当计算机程序向缓冲区内填充的数据位数超过了缓冲区本身的容量时,溢出的数据覆盖在合法数据上。理想情况是,程序检查数据长度并且不允许输入超过缓冲区长度的字符串。但是绝大多数程序都会假设数据长度总是与所分配的存储空间相匹配,这就为缓冲区溢出埋下了隐患;同样,堆栈也会出现缓冲区溢出。

当一个超长的数据进入到缓冲区时,超出部分就会被写入其他缓冲区,其他缓冲区存放的可能是数据、下一条指令的指针或者是其他程序的输出内容,这些内容可能会被覆盖或者破坏掉,甚至一小部分数据或者一条指令的溢出就可能导致一个程序或者操作系统崩溃。

溢出的根源在于编程,缓冲区溢出是由编程错误引起的。如果缓冲区被写满,而程序没有去检查缓冲区边界,也没有停止接收数据,这时缓冲区溢出就会发生。缓冲区边界检查被认为是不会有收益的管理支出,计算机资源不够或者内存不足是编程者不编写缓冲区边界检查语句的理由,虽然摩尔定律已经使这一理由失去了存在的基础,但是多数用户仍然在主要应用中运行十年甚至二十年前的程序代码。早在 1988 年,美国康奈尔大学 23 岁的计算机科学系研究生莫里斯(Morris)就已经利用 UNIX fingered 程序不限制输入长度的漏洞使缓冲器溢出。Morris 又写了一段程序使他的恶意程序能以 Root 身份执行,并传播到其他机器上,结果造成 6 000 台 Internet 上的服务器瘫痪,占当时总数的 10%。

通常,缓冲区溢出攻击具有一次完成攻击代码植入和程序转向攻击代码两种功能。攻击者首先将目标定为具有溢出漏洞的自动变量;然后向程序传递超长的字符串,进而引发缓冲区

溢出;最后这段精巧设计的攻击代码以一定的权限运行漏洞程序,进而获得目标主机的控制权。事实上,如果成功利用缓冲漏洞,攻击者就有可能获得对远程计算机的完全控制,并以本地系统权限执行任意指令,如安装程序,查看或更改、删除数据,格式化硬盘等,危害性不言而喻。

2000年7月,微软 Outlook 以及 Outlook Express 被发现存在漏洞,能够使攻击者仅通过发送邮件就能危及目标主机安全,只要邮件头部程序被运行,就会产生缓冲区溢出,并且触发恶意代码。2001年8月,“红色代码”利用微软 IIS 漏洞产生缓冲区溢出,成为攻击企业网络的“罪魁祸首”。2003年1月,Slammer 蠕虫利用微软 SQL 漏洞产生缓冲区溢出对全球互联网产生冲击。一种名为“冲击波”的蠕虫病毒利用微软 RPC 远程调用存在的缓冲区漏洞对 Windows 2000/XP、Windows Server 2003 进行攻击,中毒的机器会反复重启或者执行拷贝、粘贴功能但不工作等现象,危害波及全球网络系统。据 CERT 安全小组的数据显示,操作系统中超过 50%的安全漏洞都是由缓冲区溢出引起的,其中大多数与微软技术有关,这些与缓冲区溢出相关的安全漏洞正在被越来越多的蠕虫病毒所利用。

缓冲区溢出之所以泛滥,与 C 语言的广泛使用不无关系。被广泛使用的 C 语言没有建立检测机制,标准 C 语言具有许多复制和添加字符串的函数,这使得标准 C 语言很难进行边界检查。C++ 虽然作了一些限制,但是仍然存在缓冲区溢出。一般情况下,覆盖其他数据区的数据是没有意义的,最多造成应用程序错误。但是,如果输入的数据是经过黑客或者病毒精心设计的,覆盖缓冲区的数据恰恰是黑客或者病毒的入侵程序代码,一旦多余字节被编译执行,黑客或者病毒就有可能为所欲为,获取系统的控制权。

现在,使用最多的操作系统如 Windows、Linux、UNIX 等和数据库的开发大都依赖于 C 语言,所以这些操作系统、数据库等大型应用程序成为了缓冲区溢出攻击的重灾区。

2.6.2 缓冲区溢出实例分析及其防范

下面通过一个简单的例子分析缓冲区溢出的形成原因,在实例分析之前,先介绍一下堆栈的使用。堆栈可以向下(向内存低地址)增长,也可以向上增长,向下增长是很多计算机的实现方式,包括 Intel、Motorola、SPARC 和 MIPS 处理器。下面用一个简单的例子来介绍缓冲区的溢出。该例子中,堆栈是向下增长的。

```
void function(int a, int b, int c)
{
    char buffer1[5];
    char buffer2[10];
}

void main()
{
    function(1,2,3);
}
```

使用 gcc 的 -S 选项编译,产生汇编代码输出,通过查看汇编语言输出,可知对 function() 的调用被翻译成:

```
pushl $3
pushl $2
pushl $1
```

```
call function
```

以从后往前的顺序将 function 的 3 个参数压入栈中,然后调用 function()。指令 call 会把指令指针(IP)也压入栈中。被保存的 IP 称为返回地址(RET)。

在函数中所做的第一件事情是本例的开始工作:

```
pushl %ebp
movl %esp, %ebp
subl $ 20, %esp
```

执行过程将帧指针(EBP)压入栈中;然后把当前的 SP 复制到 EBP 使其成为新的指针,这个被保存的 SP 叫做 SFP;接下来将 SP 的值减小,为局部变量保留空间(内存只能以字为单位寻址,这里一个字是 4 个字节),5 个字节的缓冲区 buffer1 占用 8 个字节的内存空间,而 10 个字节的缓冲区 buffer2 占用 12 个字节的内存空间,这就是 SP 要减掉 20 的原因。这样就可以知道 function()被调用时堆栈的使用情况,如图 2-13 所示。

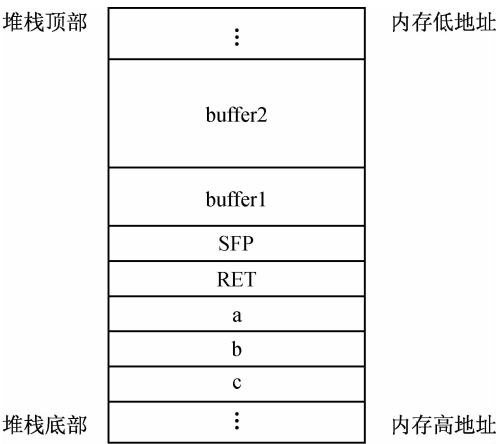


图 2-13 function()被调用时的堆栈

从图 2-13 来看,假如输入的 buffer1 超长了就直接覆盖掉后面的 SFP 和 RET,这样该函数的返回地址就被修改了。

现在试着修改第一个例子,让它可以覆盖返回地址,而且使它可以执行任意代码。堆栈中在 buffer1 之前的是 SFP,SFP 之前是返回地址 RET。RET 从 buffer1 的结尾算起是 4 个字节。buffer1 实际上是两个字即 8 个字节长。因此返回地址从 buffer1 的开头算起是 12 个字节。使用这种方法修改返回地址,跳过函数调用后面的赋值语句“x=1;”,为了做到这一点把返回地址加上 8 个字节,这样就可以跳过赋值语句而直接执行 printf。代码是这样的:

```
void function(int a, int b, int c)
{
    char buffer1[5];
    char buffer2[10];
    int * ret;
    ret=buffer1+12;
    (* ret) +=8;
}
```

```
void main()
{
    int x;
    x=0;
    function(1,2,3);
    x=1;
    printf(" %d\n",x);
}
```

缓冲区溢出的一个致命的后果是让程序执行它本来不愿意执行的函数。利用这一漏洞可以通过计算机网络攻击系统。通常,给程序输入字符串,如果这个字符串包含一些可执行代码的字节编码,就称之为漏洞入侵代码(exploit code)。所以,执行 RET 指令的效果就是跳到漏洞入侵代码段,攻击中的漏洞入侵代码实际就是一个称之为 Shell 的程序,提供给攻击者一组操作系统的函数,这样,攻击者取得系统的控制权就易如反掌了。

从“红色代码”到 Slammer,再到爆发的“冲击波”,都是利用缓冲区溢出漏洞的典型。缓冲区溢出攻击不是一种窃密和欺骗的手段,而是从计算机系统的最底层发起的攻击。因此,在它的攻击下系统的身份验证和访问权限等安全策略就形同虚设了。同时,由于攻击者传输的数据分组并无异常特征,没有任何欺骗,且用来实施缓冲区溢出攻击的字符串非常多样化,无法与正常数据进行有效区分,因而传统安全工具(如防火墙)对这种攻击方式也无能为力。防范缓冲区溢出攻击应注意以下几点:

(1)编写正确的代码。编写正确的代码是一件非常有意义却耗时的工作,特别是编写像 C 语言这种容易出错的程序。这是由只追求性能而忽视正确性的传统引起的。尽管程序员知道如何编写安全的程序,但具有安全漏洞的程序依旧出现。因此人们开发了一些工具和技术来帮助经验不足的程序员编写安全正确的程序。虽然这些工具可以帮助程序员开发更安全的程序,但是由于 C 语言的特点,这些工具不可能找出所有的缓冲区溢出漏洞。所以,这些技术只能用来减少缓冲区溢出的可能,并不能完全地避免它的发生。除非程序员能保证他的程序万无一失,否则还是要用到以下部分的内容来保证程序的可靠性。

(2)通过操作系统使得缓冲区不可执行,从而阻止攻击者植入攻击代码。这种方法有效地阻止了很多缓冲区溢出的攻击,但是攻击者并不一定要通过植入攻击代码来实现缓冲区溢出的攻击,所以这种方法还是存在很多弱点的。

(3)利用编译器的边界检查来实现缓冲区的保护。这个方法使得缓冲区溢出不可能出现,从而完全消除了缓冲区溢出的威胁,但是相对而言,代价比较大。

(4)在程序指针失效前进行完整性检查。虽然这种方法不能使得所有的缓冲区溢出失效,但它的确阻止了绝大多数的缓冲区溢出攻击。

从长远来看,要想从根本上消除缓冲区溢出攻击,需要对编程模式或 CPU 体系进行基础性修改。随着信息技术的飞速发展和人们对网络安全重视程度的不断加强,相信总会有解决缓冲区溢出攻击的最佳途径。

本章小结

随着 Internet 在各领域的广泛应用,黑客入侵和攻击已成为当前网络安全领域的重要课题。本章介绍了黑客攻击的一般步骤及常见方法,对于网络安全的初学者了解网络安全的现状以及对后续章节的学习都是很有帮助的。只有做到知己知彼,才能有的放矢针对各种攻击采取有效的防范措施,因为根据不同的攻击类型和方式而采用不同的应对手段,是提高网络安全最有效的途径。

习 题 2

1. 简述黑客攻击的一般过程。
2. 扫描器有哪些类型? 端口扫描的方法主要有哪些?
3. 为什么可以在以太网中实现网络监听?
4. 简述拒绝服务攻击的形成原因。
5. 常见的 DoS 攻击手段有哪些?
6. 简述缓冲区溢出攻击的形成过程及有效防范方法。

第 3 章 数字加密与认证

随着互联网上各种应用的开展,越来越多的网络连接在一起,安全性也变得越来越重要。企业的信息需要保护,特别是包含了企业重要信息的数据需要防止黑客、企业的员工以及竞争对手破坏。如何安全地进行用户身份认证已经成为每个开发人员和网站应用策划者急需解决的问题。

加密是为了保证信息安全而采取的一种措施。加密能够有效地保护数据文件或传输数据的内容,从而减小被非授权方窃取的可能性。加密可以检测出对数据的偶然或故意的变动,也能提供对文档作者的验证。

3.1 密码学基础

密码学由密码编码学和密码分析学组成。密码学是数字信息的实际保护、控制和识别。密码学包括信息的保密传输和身份的不可抵赖等,是一种实现安全目标自动化的方法(或者算法形式),在古代就有所应用。

3.1.1 密码学的起源及发展

作为保障数据安全的一种方式,数据加密技术起源于公元前 2000 年。虽然当时的加密技术不是现在所讲的加密技术,但作为一种加密的概念,埃及人确实是最先使用特别的象形文字作为信息编码的。随着时间推移,巴比伦、美索不达米亚和希腊文明都开始使用一些方法来保护他们的书面信息,后来这种技术被 Julius Caesar(凯撒大帝)使用,即成为凯撒密码。数据加密技术也曾用于历次战争中,包括美国独立战争、美国内战和两次世界大战。最广为人知的编码计算机是 German Enigma 机,在第二次世界大战中德国人利用它来加密信息。后来通过 Alan Turing 和 Ultra 计划以及其他人的努力,终于对德国人的密码进行了破解。当初,计算机的研究就是为了破解德国人的密码,人们并没有想到计算机所带来的信息革命。

20 世纪 70 年代中期,随着计算机的发展,信息安全技术进入大发展时期。具体标志为 1976 年 Diffie 和 Hellman 发表的文章《密码学的新动向》和 1977 年正式公布和实施的美国数据加密标准 DES。

密码学的发展大致可分为以下 3 个阶段:

第一阶段为从古代到 1949 年。这一时期可以看做是密码学科学的前夜时期,这个阶段的密码技术可以说是一种艺术,而不是一种科学。密码学专家常常是凭知觉和灵感而不是通过推理和证明来进行密码设计和分析。

第二阶段为从 1949 年到 1975 年。1949 年 Shannon 发表的《保密系统的信息理论》为近代密码学建立了理论基础,从此密码学成为了一门科学。但密码学直到今天仍具有艺术

性,是具有艺术性的一门科学。这段时期密码学理论的研究工作进展不大,公开的密码学文献很少。

第三阶段为从1976年至今。1976年Diffie和Hellman发表的《密码学的新动向》一文引发了密码学上的一场革命。他们首先证明了在发送端和接收端无密钥传输的保密通信是可能的,从而开创了公钥密码学的新纪元。此外,排列码加密解密方法使加密强度有了一个飞跃性的提高。

从以上3个发展阶段不难看出,随着计算机技术的发展,今天的加密技术在计算机数据加密中得到了广泛应用并成为保护计算机信息的重要手段。由于计算机运算能力的增强,人们不断地研究出新的数据加密方式,可以说计算机推动了数据加密技术的发展,而数据加密技术也已经成为解决当今网络安全问题的核心技术。

3.1.2 密码学概述

1. 密码系统的定义与相关基本概念

密码学这个词是从意为“秘密”、“书写”的两个希腊字演化而来的,其基础部件是密码系统。

密码系统是一个5元组 (E, D, M, K, C) , M 是明文集, K 是密钥集, C 是密文集, $E: M \times K \rightarrow C$ 是加密函数集, $D: C \times K \rightarrow M$ 是解密函数集。

下面对密码系统定义中出现的几个术语和密码学涉及的概念加以解释:

- (1)明文:被隐蔽的消息称做明文(plaintext)。
- (2)密文:隐蔽后的消息称做密文(ciphertext)或密报(cryptogram)。
- (3)加密:将明文变换成密文的过程称做加密(encryption)。
- (4)解密:由密文恢复出原明文的过程称做解密(decryption)。
- (5)加密算法:密码员对明文进行加密时采用的一组规则称做加密算法(encryption algorithm)。
- (6)发送者和接收者:传送消息的一方称做发送者(sender),简称为发方,而传送消息的预定对象称做接收者(receiver),简称为收方。
- (7)解密算法:接收者对密文进行解密时采用的一组规则称做解密算法(decryption algorithm)。
- (8)加密密钥和解密密钥:加密算法和解密算法的操作通常是在一组密钥(key)的控制下进行的,分别称为加密密钥(encryption key)和解密密钥(decryption key)。

此外,对明文进行加密操作的人员称做密码员或加密员(cryptographer)。

在所有著名密码系统中,凯撒密码是一个典型的替代密码。该密码系统可表达如下:

$M = \{\text{所有的罗马字母序列}\}$

$K = \{i \mid i \text{ 是整数, 满足 } 0 \leq i \leq 25\}$

$E = \{E_k \mid k \in K \text{ 且对所有 } m \in M, E_k(m) = (m+k) \bmod 26\}$

$D = \{D_k \mid k \in K \text{ 且对所有 } c \in C, D_k(c) = (26+c-k) \bmod 26\}$

每一个 D_k 仅仅是相应 E_k 的转换。且有:

$$C = M$$

因为 E 明显是一个满射函数。

例如,假设密钥 k 是3,那么字母A变成D,B变成E,依此类推,最后Z变成C。于是单

词 HELLO 通过凯撒密码加密后就变成 KHOOR。

以每个字母在字母表中的位置表示(A 的位置为 0),那么 HELLO 写成“7 4 11 11 14”。如果 $k=3$,那么密文为“10 7 14 14 17”,或者 KHOOR。

2. 数据加密模型

一般的数据加密模型如图 3-1 所示。

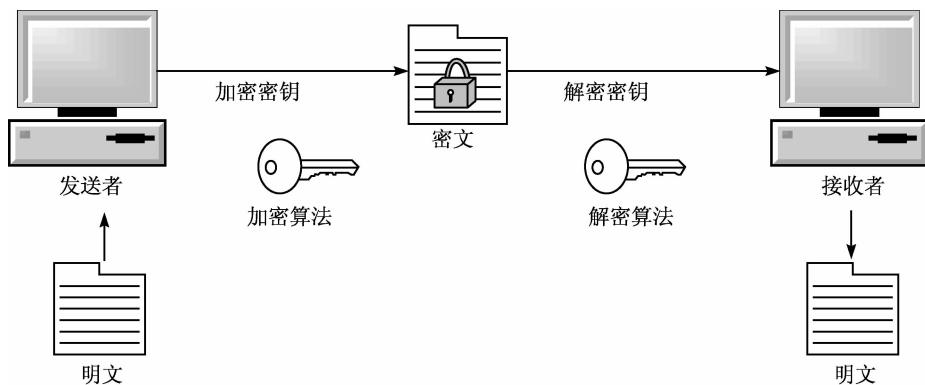


图 3-1 数据加密模型示意图

通常,数据加密的基本过程就是对原来为明文的文件或数据按某种算法进行处理变成不可读的密文,使其只能在输入相应的密钥之后才能显示出原文内容。通过这样的途径来达到保护数据不被非法窃取、阅读的目的。该过程的逆过程为解密,即将该编码信息转化为其原来数据的过程。

3. 网络数据加密的主要方式

目前,网络数据加密主要有 3 种方式:链路加密、结点加密、端到端加密。

1) 链路加密

对于在两个网络结点间的某一个通信链路,链路加密能为网上传输的数据提供安全保证。对于链路加密(又称在线加密),所有消息在被传输之前进行加密,在每一个结点对接收到的消息进行解密,然后先使用下一个链路的密钥对消息进行加密再传输。在到达目的地之前,一条消息可能要经过许多通信链路的传输。链路加密具有简单、实现起来比较容易、对用户透明的优点,但是缺点是安全性低和成本较高。

2) 结点加密

与链路加密不同,结点加密不允许消息在网络结点以明文形式存在,它先把收到的消息进行解密,然后采用另一个不同的密钥进行加密,这一过程在结点上的一个安全模块中进行。

尽管结点加密能给网络数据提供较高的安全性,但它在操作方式上与链路加密是类似的。两者均在通信链路上为传输的消息提供安全保证;都在中间结点先对消息进行解密,然后进行加密。因为要对所有传输的数据进行加密,所以加密过程对用户是透明的。结点加密的优点是成本低、安全,但缺点是结点加密要求报头和路由信息以明文形式传输,以便中间结点能得到如何处理消息的信息,因此这种方法对于防止攻击者分析通信业务是脆弱的。

3) 端到端加密

端到端加密允许数据在从源点到终点的传输过程中始终以密文形式存在。采用端到端加密,消息在到达终点之前不进行解密,在整个传输过程中均受到保护。因此,即使有结点

被损坏也不会使消息泄露。

端到端加密系统的成本比较低,并且与链路加密和结点加密相比更可靠,更容易设计、实现和维护。端到端加密还避免了其他加密系统所固有的同步问题,因为每个报文包均是独立被加密的,所以一个报文包所发生的传输错误不会影响后续的报文包。此外,从用户对安全需求的直觉上讲端到端加密更自然些。单个用户可能会选用这种加密方法,以便不影响网络上的其他用户。

端到端加密系统通常不允许对消息的目的地址进行加密,这是因为每一个消息所经过的结点都要用此地址来确定如何传输消息。由于这种加密方法不能掩盖被传输消息的源结点与目的结点,因此它对于防止攻击者分析通信业务是脆弱的。

4. 网络加密算法

根据对明文信息加密方式的不同进行分类,网络加密算法可分为分组加密算法和序列加密算法:

(1) 分组加密算法:每次只加密一个二进制位。

(2) 序列加密算法:每次对一组进行加密。

根据收发双方的密钥是否相同来进行分类,又可以分为对称式加密算法和非对称式加密算法:

(1) 对称式加密算法:加密和解密使用同一个密钥。

(2) 非对称式加密算法:加密和解密所使用的不是同一个密钥,通常有两个密钥,称为“公钥”和“私钥”,它们两个必须配对使用,否则不能打开加密文件。

“公钥”是可以对外公布的,“私钥”只能由持有人知道。采用对称式加密算法的文件在网络上传输时,如果把文件的密钥告诉对方,不管用什么方法都有可能被别人窃取。而非对称式加密算法有两个密钥,其中的“公钥”是可以公开的,收件人解密时只要用自己知道的“私钥”即可,这样就很好地避免了密钥的传输安全性问题。

3.1.3 对称密钥算法

1. 基本工作原理及特点

对称密钥算法是指加密算法的加密密钥与解密密钥是相同的,或者虽然不同但由其中一个可以很容易地推导出另一个。密钥在信息传输的双方之间需要建立安全通道进行传递和分发,如果有第三方发现该密钥则会造成失密。因此,对称密钥算法的安全性依赖于密钥,泄露密钥就意味着任何人都能对消息进行加密解密。只要通信需要保密,密钥就必须保密,密钥在传输通道中的传递与分发不适合用明文的形式。

现代计算机密码算法的典型分组长度为 64 位,这个长度可以大到足以防止分析破译,但又可以小到足以方便使用。

根据密码系统定义,对称密钥算法的加密和解密可分别表示为:

$$E_k(M)=C; D_k(C)=M$$

显然,这种算法具有特性 $D_k(E_k(M))=M$ 。

通常,采用对称密钥算法的加密方案有 5 个组成部分:明文、加密算法、密钥、密文和解密算法。其中,加密算法和解密算法互为逆运算。加密算法是指以密钥为参数对明文进行多种置换和转换的规则和步骤,最终变换结果为密文;而解密算法是加密算法的逆变换,指以密文为输入、密钥为参数的规则和步骤,最终把密文变换为明文。

对称密钥的算法简单且系统开销小,加密效率高(加/解密速度能达到数十兆每秒或更多),适合加密大量数据。此外,对称密钥算法还有一个优点,如果能够保证密钥不被第三方所获取,则信息具有等同的保密强度,只是需要维护的密钥个数会较多。一般来说,一个用户要与 N 个其他的用户进行加密通信,那么每个用户都应该有一个不同的密钥,即应该有 N 把密钥。如果网络中的 N 个用户之间都要进行两两之间的加密通信,则需要维护 $N \times (N-1)$ 把密钥来保证他们之间通信的安全。

尽管对称密钥算法有一些很好的特性,但也存在着明显的缺陷,主要有:

(1)进行安全通信前需要以安全方式进行密钥交换。这一步,在某种情况下是可行的,但在某些情况下会非常困难,甚至无法实现。

(2)密钥长度短,密码空间小,穷举方式攻击的代价小。

(3)密钥数量众多。 A 与 B 两人之间的密钥必须不同于 A 和 C 两人之间的密钥,否则, A 给 B 或 C 的消息的安全性就会受到威胁。例如,在有 1 000 个用户的团体中, A 需要保持至少 999 个密钥(更确切地说是 1 000 个,留一个密钥给自己加密数据)。对于该团体中的其他用户,此种情况同样存在。这样,这个团体一共需要将近 50 万个不同的密钥。

通过应用基于对称密码的中心服务结构,上述问题会有所缓解。在这个体系中,团体中的任何一个用户与中心服务器(通常称做密钥分配中心)共享一个密钥。因而,需要存储的密钥数量基本上和团体的人数差不多,而且中心服务器也可以为以前互相不认识的用户充当“介绍人”。但是,这个与安全密切相关的中心服务器必须随时都是在线的,因为只要服务器一掉线,用户间的通信将中断。这就意味着中心服务器是整个通信成败的关键和受攻击的焦点,它将成为一个庞大组织通信服务的瓶颈。

2. 常用的对称密钥算法

常用的对称密钥算法有 DES 算法与 IDEA 算法。

1) DES 算法

DES 算法是一种迭代的分组密码,1977 年,美国政府颁布采纳 IBM 公司设计的方案作为非机密数据的数据加密标准(data encryption standard, DES)。它的输入与输出都是 64 位,包括一个 56 位的密钥和附加的 8 位奇偶校验位。攻击该算法的主要方法是穷举密钥法,因此,DES 算法并不算非常安全,但是破译它也需要较长的时间,所以只要破译的时间超过了密文的有效期,则该加密就是有效的。现在为了提高安全性,出现了 112 位密钥,即对数据进行 3 次加密的算法,称为 3DES。

DES 算法在 POS、ATM、磁卡及智能卡(IC 卡)、加油站、高速公路收费站等领域被广泛应用,以此来实现关键数据的保密。如信用卡持卡人的 PIN 的加密传输、IC 卡与 POS 间的双向认证、金融交易数据包的 MAC 校验等均用到 DES 算法。

DES 算法的入口参数有 3 个:Key、Data、Mode。其中,Key 为 8 个字节共 64 位,是 DES 算法的工作密钥;Data 也为 8 个字节 64 位,是要被加密或被解密的数据;Mode 为 DES 的工作方式,即加密或解密。

DES 算法的工作方式为:若 Mode 为加密,则用 Key 把数据 Data 进行加密,生成 Data 的密码形式(64 位)作为 DES 的输出结果;若 Mode 为解密,则用 Key 把密码形式的数据 Data 解密,还原为 Data 的明码形式(64 位)作为 DES 的输出结果。在通信网络的两端,双方约定一致的 Key,在通信的源点用 Key 对核心数据进行 DES 加密,然后以密文形式在公共通信网(如互联网)中传输到通信网络的目的地。当数据到达目的地后用同样的 Key 对密

码数据进行解密,便再现了明码形式的数据,这样就保证了核心数据(如 PIN、MAC 等)在公共通信网中传输的安全性和可靠性。通过定期在通信网络的源端和目的端同时改用新的 Key,便能进一步提高数据的保密性。这也是现在金融交易网络的流行做法。

如图 3-2 所示是 DES 加密算法的框图,其中,明文分组长为 64 位;密钥长为 64 位,其中有效密钥长度为 56 位。图的左边是明文的处理过程,有 3 个阶段,首先是一个初始置换 IP ,用于重排明文分组的 64 位数据。然后是具有相同功能的 16 轮变换,每轮中都有置换和代换运算,第 16 轮变换的输出分为左右两部分,并被交换次序。最后再经过一个逆初始置换 IP^{-1} ,从而产生 64 位的密文。

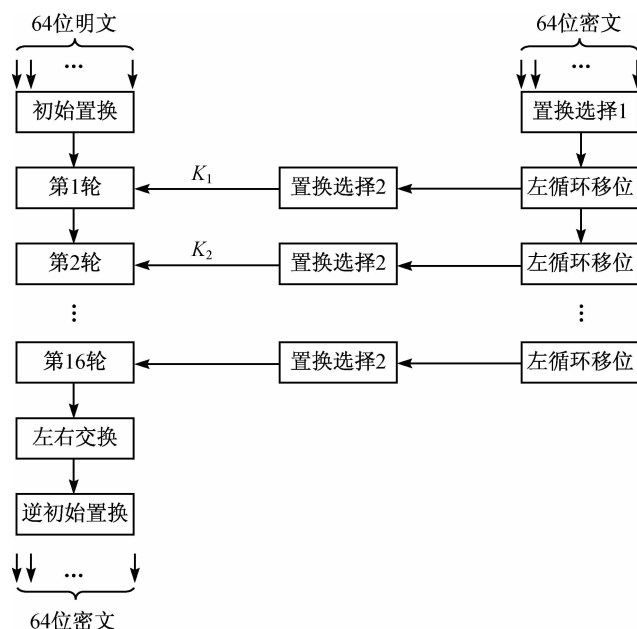


图 3-2 DES 加密算法框图

图 3-2 的右边使用 64 位密钥先通过一个置换选择,然后在加密过程的每一轮通过一个左循环移位和一个置换产生一个子密钥,其中每轮的置换都相同。由于密钥被重复迭代,所以产生的每轮子密钥不相同。

(1)初始置换与逆置换。DES 的初始置换规则如表 3-1 所示。

表 3-1 初始置换 (IP 表)

58	50	42	34	26	18	10	2	60	52	44	36	28	20	12	4
62	54	46	38	30	22	14	6	64	56	48	40	32	24	16	8
57	49	41	33	25	17	9	1	59	51	43	35	27	19	11	3
61	53	45	37	29	21	13	5	63	55	47	39	31	23	15	7

将输入的第 58 位换到第 1 位,第 50 位换到第 2 位,依此类推,最后 1 位是原来的第 7 位。 L_0 、 R_0 则是换位输出后的两部分, L_0 是输出的左 32 位, R_0 是右 32 位。例如,设置换前的输入值为 $D_1D_2D_3\cdots D_{64}$,则经过初始置换后的结果为: $L_0=D_{58}D_{50}\cdots D_8$, $R_0=D_{57}D_{49}\cdots D_7$ 。

逆初始置换正好是初始置换的逆运算。例如,第 1 位经过初始置换后处于第 40 位,而通过逆初始置换又将第 40 位换回到第 1 位,其逆初始置换规则如表 3-2 所示。

表 3-2 逆初始置换(IP^{-1} 表)

40	8	48	16	56	24	64	32	39	7	47	15	55	23	63	31
38	6	46	14	54	22	62	30	37	5	45	13	53	21	61	29
36	4	44	12	52	20	60	28	35	3	43	11	51	19	59	27
34	2	42	10	50	18	58	26	33	1	41	9	49	17	57	25

表 3-1 和表 3-2 分别给出了初始置换规则和逆初始置换规则,为了显示这两个置换的确是彼此互逆的,考虑下面 64 位的输入 M :

M_1 M_2 M_3 M_4 M_5 M_6 M_7 M_8
 M_9 M_{10} M_{11} M_{12} M_{13} M_{14} M_{15} M_{16}
 M_{17} M_{18} M_{19} M_{20} M_{21} M_{22} M_{23} M_{24}
 M_{25} M_{26} M_{27} M_{28} M_{29} M_{30} M_{31} M_{32}
 M_{33} M_{34} M_{35} M_{36} M_{37} M_{38} M_{39} M_{40}
 M_{41} M_{42} M_{43} M_{44} M_{45} M_{46} M_{47} M_{48}
 M_{49} M_{50} M_{51} M_{52} M_{53} M_{54} M_{55} M_{56}
 M_{57} M_{58} M_{59} M_{60} M_{61} M_{62} M_{63} M_{64}

其中, M_i 是二元数字。由表 3-1 得 $X=IP(M)$ 为:

M_{58} M_{50} M_{42} M_{34} M_{26} M_{18} M_{10} M_2
 M_{60} M_{52} M_{44} M_{36} M_{28} M_{20} M_{12} M_4
 M_{62} M_{54} M_{46} M_{38} M_{30} M_{22} M_{14} M_6
 M_{64} M_{56} M_{48} M_{40} M_{32} M_{24} M_{16} M_8
 M_{57} M_{49} M_{41} M_{33} M_{25} M_{17} M_9 M_1
 M_{59} M_{51} M_{43} M_{35} M_{27} M_{19} M_{11} M_3
 M_{61} M_{53} M_{45} M_{37} M_{29} M_{21} M_{13} M_5
 M_{63} M_{55} M_{47} M_{39} M_{31} M_{23} M_{15} M_7

如果再取逆初始置换 $Y=IP^{-1}(X)=IP^{-1}(IP(M))$,可以看出, M 各位的初始顺序将被恢复。

(2) 轮结构。图 3-3 是 DES 加密算法的轮结构。

在图 3-3 所示轮结构中,首先对密钥进行移位,然后从密钥的 56 位中选出 48 位。随后通过一个扩展置换将待加密数据的右 32 位扩展为 48 位,并与移位置换后的 48 位密钥进行异或操作和代换,生成新的 32 位数据,再将其置换一次。这 4 步运算构成了函数 F (图 3-3 中虚线框部分)。然后,通过另一个异或运算,函数 F 的输出与明文左半部分 32 位结合,其结果成为新的右半部分,原来的右半部分成为新的左半部分。

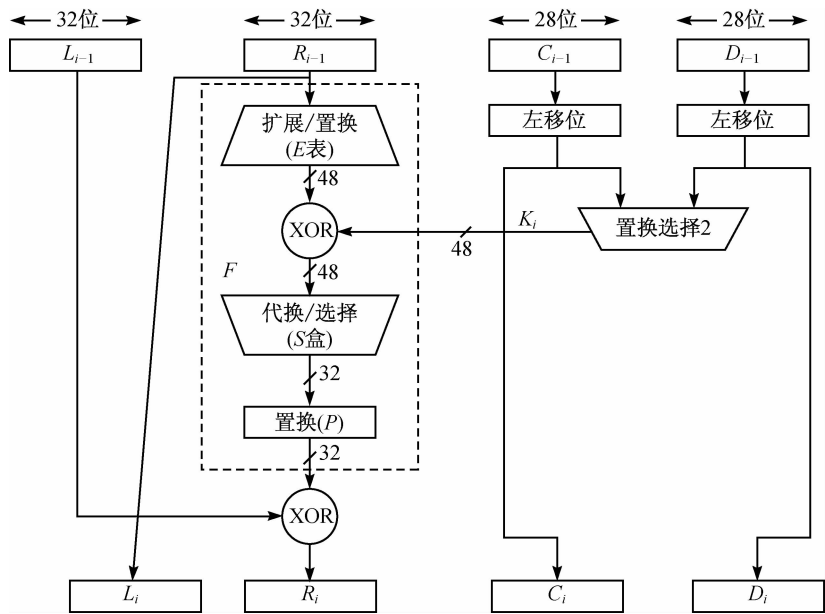


图 3-3 DES 加密算法的轮结构

其中,扩展/置换部分按照 E 表(如表 3-3 所示)将 $R_{i-1}=r_1r_2\cdots r_{32}$ 从 32 位扩展到 48 位;随后与轮密钥 K_i 进行异或运算并将结果分割为 8 个 6 位宽度的字符串: $B_1B_2\cdots B_8$,然后按照 S 盒,如表 3-4 所示的规律分别将 $B_1B_2\cdots B_8$ 映射为 8 个 4 位二进制(共 32 位)的输出。

表 3-3 扩展表(E 表)

32	1	2	3	4	5
4	5	6	7	8	9
8	9	10	11	12	13
12	13	14	15	16	17
16	17	18	19	20	21
20	21	22	23	24	25
24	25	26	27	28	29
28	29	30	31	32	1

表 3-4 S 盒变换表

列		0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
S_1	0	14	4	13	1	2	15	11	8	3	10	6	12	5	9	0	7
	1	0	15	7	4	14	2	13	1	10	6	12	11	9	5	3	8
	2	4	1	14	8	13	6	2	11	15	12	9	7	3	10	5	0
	3	15	12	8	2	4	9	1	7	5	11	3	14	10	0	6	13
S_2	0	15	1	8	14	6	11	3	4	9	7	2	13	12	0	5	10
	1	3	13	4	7	15	2	8	14	12	0	1	10	6	9	11	5
	2	0	14	7	11	10	4	13	1	5	8	12	6	9	3	2	15
	3	13	8	10	1	3	15	4	2	11	6	7	12	0	5	14	9

续表																	
列		0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
S_3	0	10	0	9	14	6	3	15	5	1	13	12	7	11	4	2	8
	1	13	7	0	9	3	4	6	10	2	8	5	14	12	11	15	1
	2	13	6	4	9	8	15	3	0	11	1	2	12	5	10	14	7
	3	1	10	13	0	6	9	8	7	4	15	14	3	11	5	2	12
S_4	0	7	13	14	3	0	6	9	10	1	2	8	5	11	12	4	15
	1	13	8	11	5	6	15	0	3	4	7	2	12	1	10	14	9
	2	10	6	9	0	12	11	7	13	15	1	3	14	5	2	8	4
	3	3	15	0	6	10	1	13	8	9	4	5	11	12	7	2	14
S_5	0	2	12	4	1	7	10	11	6	8	5	3	15	13	0	14	9
	1	14	11	2	12	4	7	13	1	5	0	15	10	3	9	8	6
	2	4	2	1	11	10	13	7	8	15	9	12	5	6	3	0	14
	3	11	8	12	7	1	14	2	13	6	15	0	9	10	4	5	3
S_6	0	12	1	10	12	9	2	6	8	0	13	3	4	14	7	5	11
	1	10	15	4	2	7	12	9	5	6	1	13	14	0	11	3	8
	2	9	14	15	5	2	8	12	3	7	0	4	10	1	13	11	6
	3	4	3	2	12	9	5	15	10	11	14	1	7	6	0	8	13
S_7	0	4	11	2	14	15	0	8	13	3	12	9	7	5	10	6	1
	1	13	0	11	7	4	9	1	10	14	3	5	12	2	15	8	6
	2	1	4	11	13	12	3	7	14	10	15	6	8	0	5	9	2
	3	6	11	13	8	1	4	10	7	9	5	0	15	14	2	3	12
S_8	0	13	2	8	4	6	15	11	1	10	9	3	14	5	0	12	7
	1	1	15	13	8	10	3	7	4	12	5	6	11	0	14	9	2
	2	7	11	4	1	9	12	14	2	0	6	10	13	15	3	5	8
	3	2	1	14	7	4	10	8	13	15	12	9	0	3	5	6	11

(3)密钥的产生。从图 3-2 和图 3-3 中可以看出,输入算法的 56 为有效密钥首先经过一个置换运算,然后将置换后的 56 位各分为 28 位的左、右两半,分别记为 C_0 和 D_0 。如表 3-5 所示。在第 i 轮分别对 C_{i-1} 和 D_{i-1} 进行左循环移位,移位前后的对应关系如表 3-6 所示。

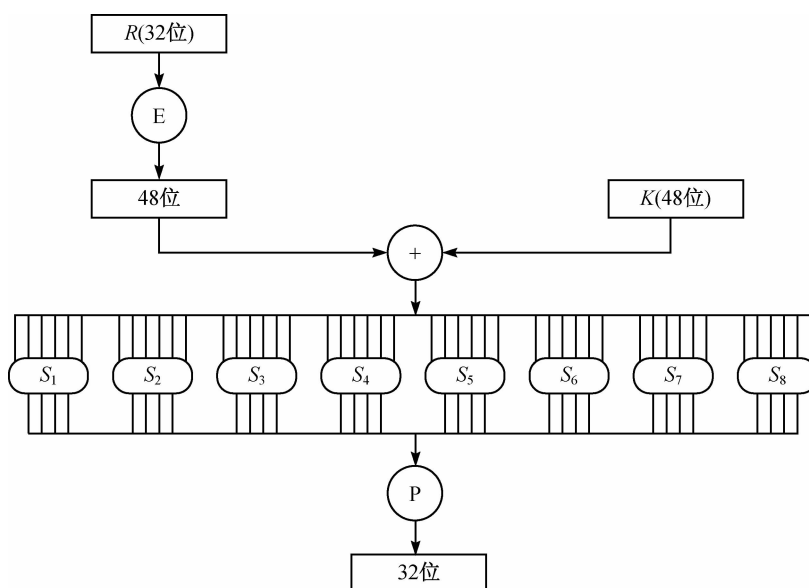
表 3-5 密钥置换安排 PC1

C_0						
57	49	41	33	25	17	9
1	58	50	42	34	26	18
10	2	59	51	43	35	27
19	11	3	60	52	44	36
D_0						
63	55	47	39	31	23	15
7	62	54	46	38	30	22
14	6	61	53	45	37	29
21	13	5	28	20	12	4

表 3-6 密钥置换安排 PC2

14	17	11	24	1	5
3	28	15	6	21	10
23	19	12	4	26	8
16	7	27	20	13	2
41	52	31	37	47	55
30	40	51	45	33	48
44	49	39	56	34	53
46	42	50	36	29	32

移位后的结果作为求下一轮子密钥的输入,同时也作为“置换选择 2”的输入。通过“置换选择 2”产生的 48 位的 K_i ,即为本轮的子密钥,作为函数 $F(R_{i-1}, K_i)$ 的输入。如图 3-4 所示。

图 3-4 函数 $F(R, K)$ 的计算过程

(4)解密。DES 的解密和加密使用同一算法,但子密钥使用的顺序相反。

2) IDEA 算法

国际数据加密算法 IDEA 是由瑞士联邦技术学院的中国学者来学嘉博士和著名的密码专家 James L. Massey 共同提出的。它在 1990 年被正式公布并在以后得到增强。这种算法是在 DES 算法的基础上发展出来的,类似于 3DES。发展 IDEA 也是因为 DES 密钥太短。IDEA 的密钥为 128 位,在今后若干年内此密钥应该是安全的,比 DES 算法更有效。

类似于 DES, IDEA 算法也是一种数据块加密算法,它设计了一系列加密轮次,每轮加密都使用从完整的加密密钥中生成的一个子密钥。与 DES 的不同之处在于,它采用软件实现和采用硬件实现同样迅速。

由于 IDEA 是在美国之外提出并发展起来的,避开了美国法律上对加密技术的诸多限

制。因此,有关 IDEA 算法和实现技术的书籍都可以自由出版和交流,极大地促进 IDEA 的发展和完善。但由于该算法出现的时间不长,针对它的攻击还不多,还未经过较长时间实际应用的考验。

3.1.4 公开密钥算法

对称密钥算法的加密、解密使用同样的密钥,在加密和解密时由发送者和接收者分别保存。采用这种方法的主要问题是密钥的生成、注入、存储、管理、分发等很复杂,特别是随着用户的增加,密钥的需求量成倍增加。在网络通信中,大量密钥的分配是一个难以解决的问题。

20 世纪 70 年代,美国斯坦福大学的两名学者 Diffie 和 Hellman 提出了一种新的加密方法 PKE,即公开密钥加密方法。与传统的加密方法不同,该技术采用两个不同的密钥来对信息加密和解密,它也称为非对称式加密方法。每个用户有一个对外公开的加密算法 E 和对外保密的解密算法 D ,它们须满足如下条件:

- (1) D 是 E 的逆,即 $D[E(X)] = X$ 。
- (2) E 和 D 都容易计算。
- (3) 由 E 出发去求解 D 十分困难。

从上述条件可看出,公开密钥密码体制下加密密钥不等于解密密钥。加密密钥可对外公开,使任何用户都可将传送给此用户的信息用公开密钥加密发送,而该用户唯一保存的私人密钥是保密的,只有它能将密文解密。虽然解密密钥理论上可由加密密钥推算出来,但这种算法设计在实际上是是不可能的,或者虽然能够推算出,但要花费很长的时间,因而是不可行的。所以,将加密密钥公开也不会危害密钥的安全。

1977 年出现了著名的 RSA 算法,该算法是基于单向陷门函数理论而设计的,数学上的单向陷门函数的特点是:单个方向求值很容易,其逆向计算却很困难。许多形式为 $y=f(x)$ 的函数。对于给定的自变量 x 值,很容易计算出函数 y 的值;而由给定的 y 值,在很多情况下依照函数关系 $f(x)$ 计算 x 值却十分困难。例如,两个大素数 p 和 q 相乘得到乘积 n 比较容易计算,但从它们的乘积 n 分解为两个大素数 p 和 q 则十分困难。如果 n 为足够大,当前的算法更不可能在有效的时间内实现。

单向陷门函数理论使得 RSA 算法的保密特性加强,主要为公用网络上信息的加密和鉴别提供了一种基本的方法。它通常是先生成一对 RSA 密钥,其中一个为解密密钥,由用户保存;另一个为公开密钥,可对外公开,甚至可在网络服务器中注册。为提高保密强度,RSA 密钥至少为 500 位长,一般推荐使用 1 024 位。这就使加密的计算量很大。为减少计算量,在传送信息时常采用传统加密方法与公开密钥加密方法相结合的方式,即信息采用改进的 DES 或 IDEA 对话密钥加密,然后使用 RSA 密钥加密对话密钥和信息摘要。对方收到信息后用不同的密钥解密并核对信息摘要。

RSA 算法的加密密钥和加密算法分开,使得密钥分配更为方便。它特别适合计算机网络环境。对于网络上的大量用户,可以将加密密钥用电话簿的方式打印出来。如果某用户想与另一用户进行保密通信,只需从公钥簿上查出对方的加密密钥,用它对所传送的信息加密发出即可。对方收到信息后用仅为自己所知的解密密钥将信息解密后查看报文的内容。由此可看出,RSA 算法解决了大量网络用户密钥管理的难题。

RSA 并不能替代 DES,它们是互补的。RSA 的密钥很长且加密速度慢,而采用 DES 正

好弥补了 RSA 的缺点,即 DES 适用于明文加密,RSA 可用于 DES 密钥的加密。由于 DES 加密速度快,所以适合加密较长的报文;而 RSA 可解决 DES 密钥分配的问题。美国的保密增强邮件(PEM)就是采用了 RSA 和 DES 结合的方法。目前,此方法已成为 E-mail 保密通信标准。具体的 RSA 算法,将在 3.4.1 节中再作详述。

3.1.5 密钥管理

密钥要求保密,就涉及密钥的管理问题。如果管理不好,密钥同样可能被无意识地泄露。任何保密只是相对的,并且是有时效的。密钥管理包括从密钥的产生到密钥的销毁过程中的各个方面,主要表现在管理体制、管理协议和密钥的产生、分配、更换和注入等。对于军用计算机网络系统,由于用户机动性强、隶属关系和协同作战指挥等方式复杂,因此对密钥管理提出了更高的要求。一个好的密钥管理系统应该做到以下几点:

- (1) 密钥难以被窃取。
- (2) 在一定条件下窃取了密钥也没有用,密钥的使用有范围和时间的限制。
- (3) 密钥的分配和更换过程对用户透明,用户不一定要亲自掌管密钥。

如果用户可以一次又一次地使用同样的密钥与别人交换信息,那么密钥也同其他任何密码一样存在着一定的安全性风险。用户的私钥虽然是不对外公开的,但是也很难保证私钥长期的保密性。如果某人偶然知道了用户的密钥,那么用户曾经和另一个人交换的每一条消息都不再是保密的了。另外,使用一个特定密钥加密的信息越多,提供给窃听者的材料也就越多,也就越不安全。因此,一般仅将一个对话密钥用于一条信息或一次对话中,或者建立一种按时更换密钥的机制以减小密钥暴露的可能性。通常,人们使用以下几种密钥管理技术:

(1) 多密钥的管理。如果某机构中有 100 个人,任意两人之间要求可以进行秘密对话,那么总共需要 4 950 个密钥,而且每个人应记住 99 个密钥。如果机构的人数非常多的话,那么,用这种办法管理密钥将是一件可怕的事情。

为了能在 Internet 上提供一个实用的解决方案,Kerberos 建立了一个安全的、可信任的密钥分发中心(key distribution center, KDC),每个用户只要记住一个和 KDC 进行会话的密钥就可以了,而不需要记住成百上千个不同的密钥。

(2) 对称密钥管理。对称加密是基于共同保守秘密来实现的。采用对称加密技术的双方必须要保证采用的是相同的密钥,要保证彼此密钥的交换是安全可靠的,同时还要设定防止密钥泄露和更改密钥的程序。这样,对称密钥的管理和分发工作将变成一个存在潜在危险的和繁琐的过程。通过公开密钥加密技术实现对称密钥的管理,使相应的管理变得简单和更加安全,同时还解决了纯对称密钥模式中存在的可靠性问题和鉴别问题。

(3) 公开密钥管理/数字证书。贸易伙伴间可以使用数字证书(公开密钥证书)来交换公开密钥。国际电信联盟(ITU)制定的标准 X.509 对数字证书进行了定义。数字证书通常包含唯一标识证书所有者(即贸易方)的名称、发布者的名称、证书所有者的公开密钥、证书发布者的数字签名、证书的有效期限及证书的序列号等。证书发布者一般称为证书管理机构(CA),它是贸易各方都信赖的机构。数字证书能够起到标识贸易方的作用,是目前电子商务广泛采用的技术之一。

(4) 数字签名。数字签名是公开密钥加密技术的另一种应用。它的主要方式是:报文的发送方从报文文本中生成一个 128 位的散列值(或报文摘要)。发送方用自己的专用密钥对

这个散列值进行加密来形成发送方的数字签名。然后,这个数字签名将作为报文的附件和报文一起发送给报文的接收方。报文的接收方首先从接收到的原始报文中计算出 128 位的散列值(或报文摘要),接着再用发送方的公开密钥来对报文附加的数字签名进行解密。如果两个散列值相同,那么接收方就能确认该数字签名是发送方的。通过数字签名能够实现对原始报文的鉴别和不可抵赖性。

目前,国际有关的标准化机构都正在着手制订关于密钥管理的技术标准规范。ISO 与 IEC 下属的信息技术委员会(JTC1)已起草了关于密钥管理的国际标准规范。该规范主要由 3 部分组成:一是密钥管理框架,二是采用对称技术的机制,三是采用非对称技术的机制。该规范现已进入到国际标准草案表决阶段,并将很快成为正式的国际标准。

3.1.6 密码分析

密码分析(cryptanalysis)着眼于找到密码系统的弱点,它通过研究密码、密文或密码系统(即秘密代码系统),以期在不知道密钥和解密算法的情况下从密文中得到原文。这也被称做破译密码、密文或密码系统。

破解通常指的是找到密码设计或执行上的弱点并减少密钥的数量,以便进行暴力破解。例如,假设一个均衡的密码执行使用 2 128 位的密钥长度,暴力破解需要尝试 2 128 个可能的组合以确信找到正确的密钥并将密文转换成原文,这对于现在和未来的计算能力看起来都是不可能实现的。然而,密码分析学揭示了密码的加密技术,破解出原文只需要 240 次尝试。这对于现有的计算机资源来说是很容易的。

在密码分析上有许多技术依赖于密码破译者想要获取的原文、密文或密码系统的其他方面。下面是一些典型的攻击:

(1)知道原文的分析:在这种情况下,密码破译者知道密文中的一部分原文,利用这些信息,密码破译者尝试找到产生密文的密钥。

(2)选择性原文分析:破译者拥有密文和原文,不过密钥并没有被分析处理,密码破译者尝试比较整个密文和原文来推断密钥。RSA 加密技术易受这种类型分析的攻击。

(3)只有密文的分析:密码破译者没有任何原文信息,就只能分析密文,这需要能够正确猜测有什么信息被加密。

(4)中间人攻击:攻击者截取两个主机之间的合法通信信息,并在双方不知道的情况下,删除或更改由一方发送给另一方的信息内容。如此,通过冒充原发送方或接收方的身份,攻击者达到非法访问通信双方保密信息的目的。

(5)调速/微分分析:这是一种诞生于 1998 年 6 月的新技术,在对抗智能卡方面非常有用,它测量一段时间内具有安全信息功能的芯片中电量的不同。这种技术用来获得在加密算法和其他功能的安全设备上的密钥信息。

3.2 数字签名与数字证书

数字签名通过某种密码运算生成一系列符号及代码组成电子密码进行签名,替代了手写签名或印章。对于这种电子式的签名还可进行技术验证,其验证的准确度是一般手写签名和印章无法比拟的。

数字签名是目前电子商务、电子政务中应用最普遍、技术最成熟、可操作性最强的一种电子签名方法。它能验证出文件的原文在传输过程中有无变动,确保传输电子文件的完整性、真实性和不可抵赖性。而基于互联网的电子交易要求买方和卖方都必须拥有合法的身份,并且在网上能够有效无误地被验证。数字证书能够验证身份,其作用类似于司机的驾驶执照或日常生活中的身份证,它由 CA 证书授权中心发行。人们可以在互联网上的交易中用它来识别对方的身份。在数字证书认证的过程中,证书认证中心作为权威的、公正的、可信赖的第三方,其作用是至关重要的。

3.2.1 电子签名

电子签名的概念没有统一的表述。美国《统一电子交易法》规定“电子签名”泛指“与电子记录相连的或在逻辑上相连的电子声音、符号或程序,而该电子声音、符号或程序是某人为签署电子记录的目的而签订或采用的”;联合国《电子商务示范法》中规定,电子签名是包含、附加在某一数据电文内,或逻辑上与某一数据电文相联系的电子形式的数据,它能被用来证实与此数据电文有关的签名人的身份,并表明该签名人认可该数据电文所载信息;欧盟的《电子签名指令》规定“电子签名”泛指“与其他电子记录相连的或在逻辑上相连并以此作为认证方法的电子形式数据”。《中华人民共和国电子签名法》是国内互联网领域首部具有法律效力的法案,该法案赋予电子签名与文本签名同等的法律效力,并明确电子认证服务市场准入制度,保障电子交易的安全,其意义非同一般。

从上述规定来看,凡是能在电子通信中起到证明当事人的身份、证明当事人对文件内容的认可的电子技术手段,都可被称为电子签名。总之,所谓电子签名,是指数据电文中以电子形式所含、所附用于识别签名人身份并表明签名人认可其中内容的数据,是现代认证技术的一般性概念,它是电子商务安全的重要保障手段。通俗地说,电子签名就是通过密码技术对电子文档的电子形式的签名,并非是书面签名的数字图像化,它类似于手写签名或印章,也可以说它就是电子印章。

目前,可以通过多种技术手段实现电子签名,在确认了签署者的确切身份后,电子签名承认人们可以用多种不同的方法签署一份电子记录。例如,基于 PKI 的公钥密码技术的数字签名,以生物特征统计学为基础的识别标识,手印、声音印记或视网膜扫描的识别,一个让收件人能识别发件人身份的密码代号、密码或个人识别码 PIN,基于量子力学的计算机等。在世界先进国家和我国普遍使用的电子签名技术是基于 PKI(public key infrastructure)的数字签名技术。

2004 年 8 月 28 日,第十届全国人大常委会第十一次会议审议通过了《中华人民共和国电子签名法》(以下简称《电子签名法》),确立了电子签名的法律效力,明确规定了“可靠的电子签名与手写签名或者盖章具有同等的法律效力”,为我国信息化建设提供了重要的法律制度保障。

根据《电子签名法》的规定,电子签名是指数据电文中以电子形式所含、所附用于识别签名人身份并表明签名人认可其中内容的数据。与手写签名或者盖章一样,电子签名有两个基本功能:一是用于识别签名人的身份,二是表明签名人对文件内容的认可。电子签名制作数据和电子签名验证数据是与电子签名有着密切联系的两个概念。电子签名制作数据,就是用于生成电子签名并将电子签名与电子签名人可靠地联系起来的字符、编码等数据。电子签名验证数据,就是用于验证电子签名的数据,包括代码、口令、算法或者公钥等。

电子签名认证证书是指可证实电子签名人与电子签名制作数据有唯一关联的数据信息或者其他电子记录,在网络环境中相当于现实社会中人们持有一个可以确认自己身份的法律证件“居民身份证”。为了简便起见,电子签名认证证书也可称为电子证书或数字证书。

3.2.2 认证机构

现实社会中用于确认身份的居民身份证是由公安机关签发的,公安机关是大家都信得过的一个权威机构。网络环境中,也需要有一个大家都信任的权威机构来签发电子签名认证证书,这个权威机构就是《电子签名法》中所说的电子签名认证服务提供者,称之为 CA (certificate authority),其主要作用就是对申请者发放、管理、取消电子签名认证证书。习惯上,也可以把 CA 称做 CA 中心、CA 认证中心等。

1. 认证中心简介

为保证网上数字信息的传输安全,除了在通信传输中采用更强的加密算法等措施之外,必须建立一种信任及信任验证机制,即参加电子商务的各方必须有一个可以被验证的标识,这就是数字证书。数字证书是各实体(持卡人/个人、商户/企业、网关/银行等)在网上信息交流及商务交易活动中的身份证明。该数字证书具有唯一性。它将实体的公开密钥同实体本身联系在一起,为实现这一目的,必须使数字证书符合 X. 509 国际标准,同时数字证书的来源必须是可靠的。这就意味着应有一个网上各方都信任的机构,专门负责数字证书的发放和管理,确保网上信息的安全,这个机构就是 CA 认证机构。各级 CA 认证机构的存在组成了整个电子商务的信任链。如果 CA 认证机构不安全或发放的数字证书不具有权威性、公正性和可信赖性,电子商务就无从谈起。

数字证书认证中心是整个网上电子交易安全的关键环节。它主要负责产生、分配并管理所有参与网上交易的实体所需的身份认证数字证书。每一份数字证书都与上一级的数字签名证书相关联,最终通过安全链追溯到一个已知的并被广泛认为是安全、权威、足以信赖的机构,即根认证中心(根 CA)。电子交易的各方都必须拥有合法的身份,即有数字证书认证中心(CA)签发的数字证书,在交易的各个环节、交易的各方都需检验对方数字证书的有效性,从而解决了用户信任问题。CA 涉及电子交易中各交易方的身份信息、严格的加密技术和认证程序。基于其牢固的安全机制,CA 应用可扩大到一切有安全要求的网上数据传输服务。

数字证书认证解决了网上交易和结算中的安全问题,其中包括:建立电子商务各主体之间的信任关系,即建立安全认证体系(CA);选择安全标准(如 SET、SSL);采用高强度的加/解密技术等。其中,安全认证体系的建立是关键,它决定了网上交易和结算能否安全进行,因此,数字证书认证中心的建立对电子商务的开展具有非常重要的意义。

认证中心主要由以下 3 部分组成:

(1)注册服务器:通过 Web Server 建立的站点,可为客户提供每天 24 小时的服务。因此客户可在自己方便的时候在网上提出证书申请和填写相应的证书申请表,免去了排队等候等烦恼。

(2)证书申请受理和审核机构:负责证书的申请和审核。它的主要功能是接受客户证书申请并进行审核。

(3)认证中心服务器:是数字证书生成、发放的运行实体,同时提供发放证书的管理、证书废止列表(CRL)的生成和处理等服务。

2. RA 简介

在数字证书认证的过程中,证书认证中心(CA)作为权威的、公正的、可信赖的第三方,其作用是至关重要的。而另一个重要机构是 RA, RA(registration authority)是数字证书注册审批机构。RA 系统是 CA 的证书发放、管理的延伸。它负责证书申请者的信息录入、审核以及证书发放等工作;同时,对发放的证书完成相应的管理功能。发放的数字证书可以存放于 IC 卡、硬盘或软盘等介质中。RA 系统是整个 CA 中心得以正常运营不可缺少的一部分。

3. 认证中心的功能

概括地说,认证中心(CA)的功能有证书发放、证书更新、证书撤销和证书验证。CA 的核心功能就是发放和管理数字证书,具体描述如下:

- (1)接收最终用户验证数字证书的申请。
- (2)确定是否接收最终用户数字证书的申请及证书的审批。
- (3)向申请者颁发或拒绝颁发数字证书及证书的发放。
- (4)接收、处理最终用户的数字证书更新请求及证书的更新。
- (5)接收最终用户数字证书的查询、撤销。
- (6)产生和发布证书废止列表(CRL)。
- (7)数字证书的归档。
- (8)密钥归档。
- (9)历史数据归档。

3.2.3 数字签名

数字签名(digital signature)是指用户用自己的私钥对原始数据的哈希摘要进行加密所得的数据。信息接收者使用信息发送者的公钥对附在原始信息后的数字签名进行解密后获得哈希摘要,并通过与自己收到的原始数据产生的哈希摘要对照,便可确任原始信息是否被篡改。这样就保证了消息来源的真实性和数据传输的完整性。

从根本上说,数字签名是一种确保电子文档(电子邮件、电子表格、文本文件等)真实可靠的方法。真实可靠的含义是:知道文档是谁创建的,并且知道在作者创建该文档之后,没有人对其进行过任何形式的修改。

数字签名依靠某些类型的加密技术来验证身份。由加密和身份验证这两个过程共同实现数字签名的功能。

对人员或计算机上的信息进行身份验证的方法如下:

(1)密码。用户名和密码的使用是最常见的身份验证方式。在计算机提示下输入用户名和密码。计算机根据安全文件对两者进行核对并确认。如果用户名或密码中有一个不匹配,计算机就不允许进行进一步访问。

(2)校验和。校验和也许是确保数据正确的最古老的方法之一,它提供了一种身份验证方式,因为无效的校验和表明数据受到了某种形式的损坏。有两种方法可以用来确定校验和。假设数据包的校验和为 1 个字节,意味着校验和可以包含的最大值为 255。如果该数据包中其他字节的和是 255 或更小,则校验和将包含那个具体的值。但如果其他字节的和大于 255,则校验和为总值除以 256 后得到的余数。

(3)私钥加密。私钥是指每台计算机都有一个密钥,在向另一台计算机发送信息包之前,可以使用该密钥对信息包进行加密。私钥要求知道哪些计算机将互相通信,并在每台计

计算机上安装相应的密钥。私钥加密的原理与密码相同,即两台计算机必须互相认识,才能对信息进行解码。密码提供了对消息进行解码的密钥。可以这样理解:如果创建了一条要发送给朋友的编码消息,其中每个字母都用它之后的第二个字母来代替。这样 A 变成了 C, B 变成了 D。然后告诉自己信任的朋友,密码是“后移两位”。这样,对方在收到消息时就可以进行解码,从而得知消息的内容。任何其他获得该消息的人看到的只是无意义的内容。

(4)公钥加密。公钥加密使用公钥和私钥相结合的方法。私钥只有自己的计算机知道,而公钥由自己的计算机告诉将要进行安全通信的所有计算机。要对加密消息进行解码时,计算机必须使用发送消息的计算机所提供的公钥和它自己的私钥。

密钥基于散列值。这个值是使用散列算法,根据一个基本输入数字计算出来的。关于散列值的重要一点是,如果不知道用于创建散列值的数据,则几乎不可能推导出原始输入数字。

要大规模实施公钥加密,就需要另外的方法,而这正是数字证书的用途所在。从根本上说,数字证书是一小段信息,它声明 Web 服务器受到证书颁发机构的独立来源的信任。证书颁发机构扮演两台计算机都信任的中间人的角色。它确认每台计算机都确实具有所表明的身分,然后为每台计算机提供另一台计算机的公钥。

1. 数字签名协议原理

在文件上手写签名长期以来被作为签名者身份的证明,或表明签名者同意文件的内容。实际上,签名体现了以下几个方面的保证:

- (1)签名是可信的。签名使文件的接收者相信签名者是慎重地在文件上签名的。
- (2)签名是不可伪造的。签名证明是签名者而不是其他人在文件上签的字。
- (3)签名不可重用。签名是文件的一部分,不可能将签名移动到不同的文件上。
- (4)签名后的文件是不可变的。文件在签名以后,就不能改变。

(5)签名是不可抵赖的。签名和文件是不可分离的,签名者事后不能声称他没有签过这个文件。

在计算机上进行数字签名并使这些保证能够继续有效还存在一些问题。首先,计算机文件易于复制,即使某人的签名难以伪造,但是将有效的签名从一个文件复制和粘贴到另一个文件是很容易的,这就使签名失去了意义;其次,文件在签名后也易于修改,并且不会留下任何修改的痕迹。

有几种公开密钥算法都能用做数字签名,这些公开密钥算法的特点是不仅用公开密钥加密的消息可以用私钥解密,而且反过来用私钥加密的消息也可以用公开密钥解密。其基本协议很简单,下面以图 3-5 为例讲解数字签名协议原理。

- (1)Mary 用她的私钥对文件加密,从而对文件签名。
- (2)Mary 将签名后的文件传给 Michael。
- (3)Michael 用 Mary 的公钥解密文件,从而验证签名。

在实际过程中,这种做法的准备效率太低了。为了节省时间,数字签名协议常常与单向散列函数一起使用。Mary 并不对整个文件签名,而是只对文件的散列值签名。

在下面的协议中,单向散列函数和数字签名算法是事先协商好的:

- (1)Mary 产生文件的单向散列值。
- (2)Mary 用她的私人密钥对散列值加密,以此表示对文件的签名。
- (3)Mary 将文件和散列签名送给 Michael。
- (4)Michael 用 Mary 发送的文件产生文件的单向散列值,同时用 Mary 的公钥对签名的

散列值解密。如果签名的散列值与自己产生的散列值匹配,则签名是有效的。

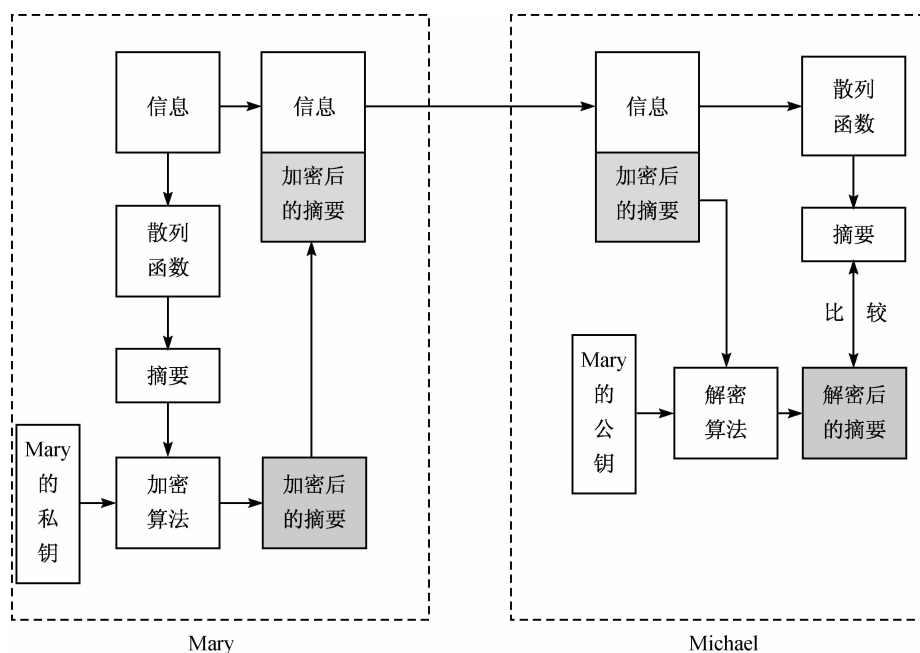


图 3-5 数字签名协议原理

由于两个不同的文件具有相同的 160 位散列值的概率为 $1/2^{160}$, 所以在这个协议中使用散列函数的签名与使用文件的签名是一样安全的。

2. 数字签名的应用实例

现在 Mary 向 Michael 传送数字信息, 为了保证信息传送的保密性、真实性、完整性和不可否认性, 需要对要传送的信息进行数字加密和数字签名, 其传送过程如下:

- (1) Mary 准备好要传送的数字信息(明文)。
- (2) Mary 对数字信息进行哈希(hash)运算, 得到一个信息摘要。
- (3) Mary 用自己的私钥(SK)对信息摘要进行加密得到 Mary 的数字签名, 并将其附在数字信息上。
- (4) Mary 随机产生一个加密密钥(DES 密钥), 并用此密钥对要发送的信息进行加密, 形成密文。
- (5) Mary 用 Michael 的公钥(PK)对刚才随机产生的加密密钥进行加密, 将加密后的 DES 密钥连同密文一起传送给 Michael。
- (6) Michael 收到 Mary 传送过来的密文和加过密的 DES 密钥, 先用自己的私钥(SK)对加密的 DES 密钥进行解密, 得到 DES 密钥。
- (7) Michael 用 DES 密钥对收到的密文进行解密, 得到明文的数字信息, 然后将 DES 密钥抛弃(即 DES 密钥作废)。
- (8) Michael 用 Mary 的公钥(PK)对 Mary 的数字签名进行解密, 得到信息摘要。
- (9) Michael 用相同的 hash 算法对收到的明文再进行一次 hash 运算, 得到一个新的信息摘要。
- (10) Michael 将收到的信息摘要和新产生的信息摘要进行比较, 如果一致, 说明收到的

信息没有被修改过。

3.2.4 公钥基础设施

公钥基础设施(PKI)采用证书管理公钥,通过可信任的第三方机构即认证中心把用户的公钥和用户的其他标识信息捆绑在一起,在 Internet 上验证用户的身份。PKI 把公钥密码和对称密码结合起来,在 Internet 上实现密钥的自动管理,保证网上数据的安全传输。

1. PKI 的概念

从广义上讲,所有提供公钥加密和数字签名服务的系统都可叫做 PKI 系统。PKI 的主要目的是通过自动管理密钥和证书,为用户建立起一个安全的网络运行环境,使用户可以在多种应用环境下方便地使用加密和数字签名技术,从而保证网上数据的机密性、完整性、有效性。数据的机密性是指数据在传输过程中,不能被非授权者偷看;数据的完整性是指数据在传输过程中不能被非法篡改;数据的有效性是指数据的传输不能被否认。

2. PKI 的组成

完整的 PKI 包括认证政策的制定、遵循的技术标准、各 CA 之间关系、安全策略、安全程度、服务对象、管理原则和框架、认证规则、运作制度的制定、所涉及各方的法律关系内容以及技术的实现。一个有效的 PKI 系统必须是安全和透明的,用户在获得加密和数字签名服务时,不需要详细地了解 PKI 是怎样管理证书和密钥的。一个典型、完整、有效的 PKI 应用系统至少应具有以下几个部分。

1) 认证中心 CA

CA 是整个 PKI 系统的核心,是 PKI 应用中权威的、可信任的、公正的第三方机构。CA 负责生成和管理 PKI 结构下的所有用户(包括各种应用程序)的证书,把用户的公钥和用户的其他信息捆绑在一起,在网上验证用户的身份。(在 3.2.2 节中已对 CA 进行了详细介绍。)

2) 注册中心 RA

注册中心 RA 是认证中心 CA 的延伸部分,它与 CA 在逻辑上是一个整体,执行不同的功能。RA 按照特定的政策和管理规范对用户的资格进行审查,以决定是否为该用户发放证书。如果审查通过,即可实时或批量地向 CA 提出申请,要求为用户签发证书。

3) 证书发布库

证书发布库是证书的集中存放地,提供公众查询服务。证书库可以是关系数据库,也可以是目录。

4) 密钥备份及恢复系统

对用户的解密密钥进行备份,当丢失时进行恢复,而签名密钥不能备份和恢复。

5) 证书作废处理系统

证书由于某些原因需要作废,终止使用时,需通过证书撤销列表(CRL)来完成。

6) PKI 应用接口系统

PKI 应用接口系统为各种应用提供安全、一致、可信任的方式与 PKI 交互,确保所建立起来的网络环境安全可信,并降低管理成本。

3. PKI 的功能操作

PKI 具体的功能操作有:产生、验证和分发密钥,签名和验证,证书的获取,验证证书,保存证书,本地保存的证书的获取,证书撤销的申请,密钥的恢复,CRL 的获取,密钥的更新,

审计,存档(证书及 CRL)等。

4. PKI 的服务

1) PKI 的核心服务

一般认为 PKI 所提供的核心服务有 3 个:

(1)认证。认证即为身份识别与鉴别,向一个实体确认另一个实体的身份,通常有两种鉴别方法:

- 实体鉴别。实体鉴别一般会产生一个明确的结果,由此允许实体进行某些操作或通信。
- 数据来源鉴别。就是鉴定某个特定的数据是否来源于某个特定的实体。

(2)完整性。数据完整性服务就是确定数据有没有被修改,即对数据无论是在传输还是在存储过程中检查有没有被修改。这一般通过数字签名技术和消息认证码两种方式来实现。

(3)保密性。保密性服务就是确保数据的秘密性,除了指定实体外,无人能解读数据的关键部分。

2) PKI 的附加服务

PKI 的附加服务也称 PKI 的支撑服务。这些服务不是任何 PKI 都具备的功能,但这些服务都建立在 PKI 的核心服务之上。PKI 的附加服务有:

(1)不可否认服务。不可否认服务是指从技术上用于保证实体对用户的行为的确认。最主要的是:对数据来源的不可否认,即用户不能否认敏感消息或文件是来源于他;接收后的不可否认性,即用户不能否认已接收到了敏感文件。此外,还包括其他类型的不可否认,如传输的不可否认,创建的不可否认以及同意的不可否认。

(2)安全时间戳。安全时间戳就是一个可信的时间权威,它用一段可认证的完整的数据表示时间戳。

(3)公证。PKI 中的公证服务指的是数据认证。PKI 公证人是一个被其他 PKI 实体所信任的实体,能够正确公正地提供公证服务。一般通过数据签名机制和时间戳服务来证明数据的正确性。

3.2.5 数字证书

数字证书就是网络通信中标志通信各方身份信息的一系列数据,其作用类似于现实生活中的身份证。它是由 CA 发行的,可以在网络中用来识别对方的身份。

使用数字证书,通过运用对称和非对称密码体制等密码技术建立起一套严密的身份认证系统,从而保证信息不被他人窃取,信息在传输过程中不被篡改。发送方能够通过数字证书来确认接收方的身份,发送方对于自己发送的信息不能抵赖。

1. 证书的类型

1) X. 509 证书

X. 509 证书是应用最广泛的一种证书。X. 509 证书符合国际电信联盟远程通信标准化组织(ITU-T)部分标准和国际标准化组织(ISO)的证书格式标准。X. 509 证书是随 PKI 的形成而新发展起来的安全机制,当前版本是 X. 509v3。相对于以前的版本,它加入了扩展字段,增强了证书的灵活性。X. 509 证书实现了身份的鉴别与识别、完整性、保密性及不可否认性等安全服务。X. 509 证书的格式如图 3-6 所示。基本证书定义如表 3-7 所示。

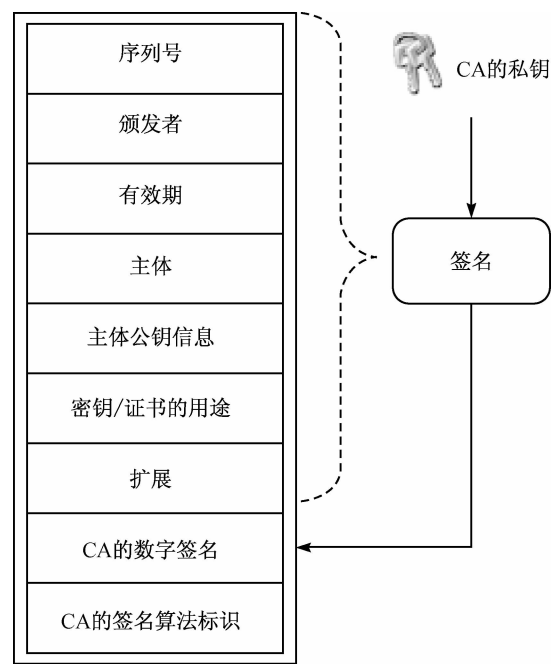


图 3-6 X.509 证书的格式

表 3-7 基本证书定义

域	定 义
标准域	证书版本号(certificate format version)
	证书序列号(certificate serial number)
	签名算法标识(signature algorithm identifier for CA)
	证书颁发者 CA 名称(issuer X.509 name)
	证书有效期(validity period)
	用户名称(subject X.509 name)
	用户公钥信息(subject public key information)
	颁发者唯一标识(issuer unique ID)
	主体证书拥有唯一标识(subject unique ID)
扩展域	CA 的公钥标识(authority key identifier)
	用户的公钥标识(subject key identifier)
	CRL 分布(CRL distribution point)
	证书中的公钥用途(key usage)
	CA 承认的证书政策列表(certificate policies)
	...
CA 签名	...
CA 签名算法	...