



Red Hat Enterprise Linux 6 Security Guide

A Guide to Securing Red Hat Enterprise Linux

Red Hat Engineering Content Services

Red Hat Enterprise Linux 6 Security Guide

A Guide to Securing Red Hat Enterprise Linux

Red Hat Engineering Content Services

Legal Notice

Copyright 2011 Red Hat, Inc. The text of and illustrations in this document are licensed by Red Hat under a Creative Commons Attribution–Share Alike 3.0 Unported license ("CC-BY-SA"). An explanation of CC-BY-SA is available at <http://creativecommons.org/licenses/by-sa/3.0/>. In accordance with CC-BY-SA, if you distribute this document or an adaptation of it, you must provide the URL for the original version. Red Hat, as the licensor of this document, waives the right to enforce, and agrees not to assert, Section 4d of CC-BY-SA to the fullest extent permitted by applicable law. Red Hat, Red Hat Enterprise Linux, the Shadowman logo, JBoss, MetaMatrix, Fedora, the Infinity Logo, and RHCE are trademarks of Red Hat, Inc., registered in the United States and other countries. Linux is the registered trademark of Linus Torvalds in the United States and other countries. Java is a registered trademark of Oracle and/or its affiliates. XFS is a trademark of Silicon Graphics International Corp. or its subsidiaries in the United States and/or other countries. MySQL is a registered trademark of MySQL AB in the United States, the European Union and other countries. All other trademarks are the property of their respective owners. 1801 Varsity Drive Raleigh, NC 27606-2072 USA Phone: +1 919 754 3700 Phone: 888 733 4281 Fax: +1 919 754 3701

Keywords**Abstract**

This book assists users and administrators in learning the processes and practices of securing workstations and servers against local and remote intrusion, exploitation and malicious activity. Focused on Red Hat Enterprise Linux but detailing concepts and techniques valid for all Linux systems, this guide details the planning and the tools involved in creating a secured computing environment for the data center, workplace, and home. With proper administrative knowledge, vigilance, and tools, systems running Linux can be both fully functional and secured from most common intrusion and exploit methods.

Table of Contents

Preface	9
1. Document Conventions	9
1.1. Typographic Conventions	9
1.2. Pull-quote Conventions	10
1.3. Notes and Warnings	11
2. We Need Feedback!	11
Chapter 1. Security Overview	13
1.1. Introduction to Security	13
1.1.1. What is Computer Security?	13
1.1.1.1. How did Computer Security come about?	13
1.1.1.2. Security Today	14
1.1.1.3. Standardizing Security	14
1.1.2. SELinux	15
1.1.3. Security Controls	15
1.1.3.1. Physical Controls	15
1.1.3.2. Technical Controls	15
1.1.3.3. Administrative Controls	16
1.1.4. Conclusion	16
1.2. Vulnerability Assessment	16
1.2.1. Thinking Like the Enemy	16
1.2.2. Defining Assessment and Testing	17
1.2.2.1. Establishing a Methodology	18
1.2.3. Evaluating the Tools	18
1.2.3.1. Scanning Hosts with Nmap	19
1.2.3.1.1. Using Nmap	19
1.2.3.2. Nessus	19
1.2.3.3. Nikto	20
1.2.3.4. Anticipating Your Future Needs	20
1.3. Attackers and Vulnerabilities	20
1.3.1. A Quick History of Hackers	20
1.3.1.1. Shades of Gray	21
1.3.2. Threats to Network Security	21
1.3.2.1. Insecure Architectures	21
1.3.2.1.1. Broadcast Networks	21
1.3.2.1.2. Centralized Servers	21
1.3.3. Threats to Server Security	22
1.3.3.1. Unused Services and Open Ports	22
1.3.3.2. Unpatched Services	22
1.3.3.3. Inattentive Administration	22
1.3.3.4. Inherently Insecure Services	23
1.3.4. Threats to Workstation and Home PC Security	23
1.3.4.1. Bad Passwords	23
1.3.4.2. Vulnerable Client Applications	23
1.4. Common Exploits and Attacks	24
1.5. Security Updates	27
1.5.1. Updating Packages	27
1.5.2. Verifying Signed Packages	28
1.5.3. Installing Signed Packages	29
1.5.4. Applying the Changes	30
Chapter 2. Securing Your Network	33

2.1. Workstation Security	33
2.1.1. Evaluating Workstation Security	33
2.1.2. BIOS and Boot Loader Security	33
2.1.2.1. BIOS Passwords	33
2.1.2.1.1. Securing Non-x86 Platforms	34
2.1.2.2. Boot Loader Passwords	34
2.1.2.2.1. Password Protecting GRUB	34
2.1.2.2.2. Disabling Interactive Startup	35
2.1.3. Password Security	35
2.1.3.1. Creating Strong Passwords	36
2.1.3.1.1. Secure Password Creation Methodology	37
2.1.3.2. Creating User Passwords Within an Organization	38
2.1.3.2.1. Forcing Strong Passwords	38
2.1.3.2.2. Passphrases	38
2.1.3.2.3. Password Aging	39
2.1.4. Administrative Controls	41
2.1.4.1. Allowing Root Access	42
2.1.4.2. Disallowing Root Access	42
2.1.4.3. Enabling Automatic Logouts	45
2.1.4.4. Limiting Root Access	46
2.1.5. Session Locking	46
2.1.5.1. Locking GNOME Using gnome-screensaver-command	47
2.1.5.1.1. Automatic Lock on Screen Saver Activation	47
2.1.5.1.2. Remote Session Locking	48
2.1.5.2. Locking Virtual Consoles Using vlock	49
2.1.6. Available Network Services	49
2.1.6.1. Risks To Services	49
2.1.6.2. Identifying and Configuring Services	50
2.1.6.3. Insecure Services	51
2.1.7. Personal Firewalls	52
2.1.8. Security Enhanced Communication Tools	53
2.2. Server Security	53
2.2.1. Securing Services With TCP Wrappers and xinetd	53
2.2.1.1. Enhancing Security With TCP Wrappers	54
2.2.1.1.1. TCP Wrappers and Connection Banners	54
2.2.1.1.2. TCP Wrappers and Attack Warnings	54
2.2.1.1.3. TCP Wrappers and Enhanced Logging	55
2.2.1.2. Enhancing Security With xinetd	55
2.2.1.2.1. Setting a Trap	55
2.2.1.2.2. Controlling Server Resources	56
2.2.2. Securing Portmap	56
2.2.2.1. Protect portmap With TCP Wrappers	57
2.2.2.2. Protect portmap With iptables	57
2.2.3. Securing NIS	57
2.2.3.1. Carefully Plan the Network	58
2.2.3.2. Use a Password-like NIS Domain Name and Hostname	58
2.2.3.3. Edit the /var/yp/securenets File	58
2.2.3.4. Assign Static Ports and Use iptables Rules	58
2.2.3.5. Use Kerberos Authentication	59
2.2.4. Securing NFS	59
2.2.4.1. Carefully Plan the Network	59
2.2.4.2. Securing NFS Mount Options	60
2.2.4.2.1. Review the NFS Server	60
2.2.4.2.2. Review the NFS Client	60
2.2.4.3. Beware of Syntax Errors	61

2.2.4.4. Do Not Use the no_root_squash Option	61
2.2.4.5. NFS Firewall Configuration	61
2.2.5. Securing the Apache HTTP Server	62
Removing httpd Modules	63
httpd and SELinux	64
2.2.6. Securing FTP	64
2.2.6.1. FTP Greeting Banner	64
2.2.6.2. Anonymous Access	64
2.2.6.2.1. Anonymous Upload	65
2.2.6.3. User Accounts	65
2.2.6.3.1. Restricting User Accounts	65
2.2.6.4. Use TCP Wrappers To Control Access	66
2.2.7. Securing Postfix	66
2.2.7.1. Limiting a Denial of Service Attack	66
2.2.7.2. NFS and Postfix	66
2.2.7.3. Mail-only Users	67
2.2.7.4. Disable Postfix Network Listening	67
2.2.8. Securing Sendmail	67
2.2.8.1. Limiting a Denial of Service Attack	67
2.2.8.2. NFS and Sendmail	68
2.2.8.3. Mail-only Users	68
2.2.8.4. Disable Sendmail Network Listening	68
2.2.9. Verifying Which Ports Are Listening	68
2.2.10. Disable Source Routing	69
2.2.11. Reverse Path Filtering	70
2.2.11.1. Additional Resources	71
2.2.11.1.1. Installed Documentation	71
2.2.11.1.2. Useful Websites	71
2.3. Single Sign-on (SSO)	72
2.4. Pluggable Authentication Modules (PAM)	72
2.5. Kerberos	72
2.6. TCP Wrappers and xinetd	72
2.6.1. TCP Wrappers	73
2.6.1.1. Advantages of TCP Wrappers	74
2.6.2. TCP Wrappers Configuration Files	74
2.6.2.1. Formatting Access Rules	75
2.6.2.1.1. Wildcards	76
2.6.2.1.2. Patterns	76
2.6.2.1.3. Portmap and TCP Wrappers	77
2.6.2.1.4. Operators	78
2.6.2.2. Option Fields	78
2.6.2.2.1. Logging	78
2.6.2.2.2. Access Control	79
2.6.2.2.3. Shell Commands	79
2.6.2.2.4. Expansions	79
2.6.3. xinetd	80
2.6.4. xinetd Configuration Files	81
2.6.4.1. The /etc/xinetd.conf File	81
2.6.4.2. The /etc/xinetd.d/ Directory	82
2.6.4.3. Altering xinetd Configuration Files	82
2.6.4.3.1. Logging Options	82
2.6.4.3.2. Access Control Options	83
2.6.4.3.3. Binding and Redirection Options	84
2.6.4.3.4. Resource Management Options	85
2.6.5. Additional Resources	86

2.6.5.1. Installed TCP Wrappers Documentation	86
2.6.5.2. Useful TCP Wrappers Websites	86
2.6.5.3. Related Books	86
2.7. Virtual Private Networks (VPNs)	87
2.7.1. How Does a VPN Work?	87
2.7.2. Openswan	87
2.7.2.1. Overview	87
2.7.2.2. Configuration	88
2.7.2.3. Commands	89
2.7.2.4. Openswan Resources	89
2.8. Firewalls	89
2.8.1. Netfilter and IPTables	92
2.8.1.1. IPTables Overview	92
2.8.2. Basic Firewall Configuration	92
2.8.2.1. Firewall Configuration Tool	92
2.8.2.2. Enabling and Disabling the Firewall	93
2.8.2.3. Trusted Services	94
2.8.2.4. Other Ports	95
2.8.2.5. Saving the Settings	95
2.8.2.6. Activating the IPTables Service	95
2.8.3. Using IPTables	95
2.8.3.1. IPTables Command Syntax	96
2.8.3.2. Basic Firewall Policies	96
2.8.3.3. Saving and Restoring IPTables Rules	96
2.8.4. Common IPTables Filtering	97
2.8.5. FORWARD and NAT Rules	98
2.8.5.1. Postrouting and IP Masquerading	99
2.8.5.2. Prerouting	99
2.8.5.3. DMZs and IPTables	100
2.8.6. Malicious Software and Spoofed IP Addresses	100
2.8.7. IPTables and Connection Tracking	101
2.8.8. IPv6	101
2.8.9. IPTables	102
2.8.9.1. Packet Filtering	102
2.8.9.2. Command Options for IPTables	103
2.8.9.2.1. Structure of IPTables Command Options	104
2.8.9.2.2. Command Options	104
2.8.9.2.3. IPTables Parameter Options	105
2.8.9.2.4. IPTables Match Options	107
2.8.9.2.4.1. TCP Protocol	107
2.8.9.2.4.2. UDP Protocol	108
2.8.9.2.4.3. ICMP Protocol	108
2.8.9.2.4.4. Additional Match Option Modules	108
2.8.9.2.5. Target Options	109
2.8.9.2.6. Listing Options	110
2.8.9.3. Saving IPTables Rules	111
2.8.9.4. IPTables Control Scripts	111
2.8.9.4.1. IPTables Control Scripts Configuration File	113
2.8.9.5. IPTables and IPv6	113
2.8.9.6. Additional Resources	113
2.8.9.6.1. Useful Firewall Websites	114
2.8.9.6.2. Related Documentation	114
2.8.9.6.3. Installed IP Tables Documentation	114
2.8.9.6.4. Useful IP Tables Websites	114

Chapter 3. Encryption	115
------------------------------------	------------

3.1. Data at Rest	115
3.1.1. Full Disk Encryption	115
3.1.2. File Based Encryption	115
3.2. Data in Motion	115
3.2.1. Virtual Private Networks	116
3.2.2. Secure Shell	116
3.2.2.1. SSH Cryptographic Login	116
3.2.3. OpenSSL Intel AES-NI Engine	117
3.2.4. LUKS Disk Encryption	117
Overview of LUKS	117
3.2.4.1. LUKS Implementation in Red Hat Enterprise Linux	118
3.2.4.2. Manually Encrypting Directories	118
3.2.4.3. Add a new passphrase to an existing device	120
3.2.4.4. Remove a passphrase from an existing device	120
3.2.4.5. Creating Encrypted Block Devices in Anaconda	120
3.2.4.6. Links of Interest	120
3.2.5. Using GNU Privacy Guard (GnuPG)	120
3.2.5.1. Creating GPG Keys in GNOME	121
3.2.5.2. Creating GPG Keys in KDE	121
3.2.5.3. Creating GPG Keys Using the Command Line	121
3.2.5.4. About Public Key Encryption	123
Chapter 4. General Principles of Information Security	124
4.1. Tips, Guides, and Tools	124
Chapter 5. Secure Installation	125
5.1. Disk Partitions	125
5.2. Utilize LUKS Partition Encryption	125
Chapter 6. Software Maintenance	126
6.1. Install Minimal Software	126
6.2. Plan and Configure Security Updates	126
6.3. Adjusting Automatic Updates	126
6.4. Install Signed Packages from Well Known Repositories	126
Chapter 7. Federal Standards and Regulations	128
7.1. Introduction	128
7.2. Federal Information Processing Standard (FIPS)	128
7.2.1. Enabling FIPS Mode	128
7.3. National Industrial Security Program Operating Manual (NISPOM)	129
7.4. Payment Card Industry Data Security Standard (PCI DSS)	130
7.5. Security Technical Implementation Guide	130
Chapter 8. References	131
Encryption Standards	133
A.1. Synchronous Encryption	133
A.1.1. Advanced Encryption Standard - AES	133
A.1.1.1. AES History	133
A.1.2. Data Encryption Standard - DES	133
A.1.2.1. DES History	133
A.2. Public-key Encryption	133
A.2.1. Diffie-Hellman	134
A.2.1.1. Diffie-Hellman History	134
A.2.2. RSA	134
A.2.3. DSA	135
A.2.4. SSL/TLS	135

A.2.5. Cramer-Shoup Cryptosystem	135
A.2.6. ElGamal Encryption	135
Revision History	137

Preface

1. Document Conventions

This manual uses several conventions to highlight certain words and phrases and draw attention to specific pieces of information.

In PDF and paper editions, this manual uses typefaces drawn from the [Liberation Fonts](#) set. The Liberation Fonts set is also used in HTML editions if the set is installed on your system. If not, alternative but equivalent typefaces are displayed. Note: Red Hat Enterprise Linux 5 and later includes the Liberation Fonts set by default.

1.1. Typographic Conventions

Four typographic conventions are used to call attention to specific words and phrases. These conventions, and the circumstances they apply to, are as follows.

Mono-spaced Bold

Used to highlight system input, including shell commands, file names and paths. Also used to highlight keys and key combinations. For example:

To see the contents of the file **my_next_bestselling_novel** in your current working directory, enter the **cat my_next_bestselling_novel** command at the shell prompt and press **Enter** to execute the command.

The above includes a file name, a shell command and a key, all presented in mono-spaced bold and all distinguishable thanks to context.

Key combinations can be distinguished from an individual key by the plus sign that connects each part of a key combination. For example:

Press **Enter** to execute the command.

Press **Ctrl+Alt+F2** to switch to a virtual terminal.

The first example highlights a particular key to press. The second example highlights a key combination: a set of three keys pressed simultaneously.

If source code is discussed, class names, methods, functions, variable names and returned values mentioned within a paragraph will be presented as above, in **mono-spaced bold**. For example:

File-related classes include **filesystem** for file systems, **file** for files, and **dir** for directories. Each class has its own associated set of permissions.

Proportional Bold

This denotes words or phrases encountered on a system, including application names; dialog box text; labeled buttons; check-box and radio button labels; menu titles and sub-menu titles. For example:

Choose **System** → **Preferences** → **Mouse** from the main menu bar to launch **Mouse Preferences**. In the **Buttons** tab, click the **Left-handed mouse** check box and click **Close** to switch the primary mouse button from the left to the right (making the mouse suitable for use in the left hand).

To insert a special character into a **gedit** file, choose **Applications** → **Accessories** →

Character Map from the main menu bar. Next, choose **Search** → **Find...** from the **Character Map** menu bar, type the name of the character in the **Search** field and click **Next**. The character you sought will be highlighted in the **Character Table**. Double-click this highlighted character to place it in the **Text to copy** field and then click the **Copy** button. Now switch back to your document and choose **Edit** → **Paste** from the **gedit** menu bar.

The above text includes application names; system-wide menu names and items; application-specific menu names; and buttons and text found within a GUI interface, all presented in proportional bold and all distinguishable by context.

Mono-spaced Bold Italic* or *Proportional Bold Italic

Whether mono-spaced bold or proportional bold, the addition of italics indicates replaceable or variable text. Italics denotes text you do not input literally or displayed text that changes depending on circumstance. For example:

To connect to a remote machine using ssh, type **ssh *username@domain.name*** at a shell prompt. If the remote machine is **example.com** and your username on that machine is john, type **ssh john@example.com**.

The **mount -o remount *file-system*** command remounts the named file system. For example, to remount the **/home** file system, the command is **mount -o remount /home**.

To see the version of a currently installed package, use the **rpm -q *package*** command. It will return a result as follows: ***package-version-release***.

Note the words in bold italics above — username, domain.name, file-system, package, version and release. Each word is a placeholder, either for text you enter when issuing a command or for text displayed by the system.

Aside from standard usage for presenting the title of a work, italics denotes the first use of a new and important term. For example:

Publican is a *DocBook* publishing system.

1.2. Pull-quote Conventions

Terminal output and source code listings are set off visually from the surrounding text.

Output sent to a terminal is set in **mono-spaced roman** and presented thus:

```
books      Desktop  documentation  drafts  mss    photos  stuff  svn
books_tests Desktop1  downloads      images  notes  scripts svgs
```

Source-code listings are also set in **mono-spaced roman** but add syntax highlighting as follows:

```

package org.jboss.book.jca.ex1;

import javax.naming.InitialContext;

public class ExClient
{
    public static void main(String args[])
        throws Exception
    {
        InitialContext iniCtx = new InitialContext();
        Object          ref    = iniCtx.lookup("EchoBean");
        EchoHome        home   = (EchoHome) ref;
        Echo            echo    = home.create();

        System.out.println("Created Echo");

        System.out.println("Echo.echo('Hello') = " + echo.echo("Hello"));
    }
}

```

1.3. Notes and Warnings

Finally, we use three visual styles to draw attention to information that might otherwise be overlooked.



Note

Notes are tips, shortcuts or alternative approaches to the task at hand. Ignoring a note should have no negative consequences, but you might miss out on a trick that makes your life easier.



Important

Important boxes detail things that are easily missed: configuration changes that only apply to the current session, or services that need restarting before an update will apply. Ignoring a box labeled 'Important' will not cause data loss but may cause irritation and frustration.



Warning

Warnings should not be ignored. Ignoring warnings will most likely cause data loss.

2. We Need Feedback!

If you find a typographical error in this manual, or if you have thought of a way to make this manual better, we would love to hear from you! Please submit a report in Bugzilla: <http://bugzilla.redhat.com/> against the product **Red Hat Enterprise Linux**.

When submitting a bug report, be sure to mention the manual's identifier: *doc-Security_Guide* and version number: **6**.

If you have a suggestion for improving the documentation, try to be as specific as possible when describing it. If you have found an error, please include the section number and some of the surrounding

text so we can find it easily.

Chapter 1. Security Overview

Because of the increased reliance on powerful, networked computers to help run businesses and keep track of our personal information, entire industries have been formed around the practice of network and computer security. Enterprises have solicited the knowledge and skills of security experts to properly audit systems and tailor solutions to fit the operating requirements of their organization. Because most organizations are increasingly dynamic in nature, their workers are accessing critical company IT resources locally and remotely, hence the need for secure computing environments has become more pronounced.

Unfortunately, many organizations (as well as individual users) regard security as more of an afterthought, a process that is overlooked in favor of increased power, productivity, convenience, ease of use, and budgetary concerns. Proper security implementation is often enacted postmortem — *after* an unauthorized intrusion has already occurred. Taking the correct measures prior to connecting a site to an untrusted network, such as the Internet, is an effective means of thwarting many attempts at intrusion.



Note

This document makes several references to files in the `/lib` directory. When using 64-bit systems, some of the files mentioned may instead be located in `/lib64`.

1.1. Introduction to Security

1.1.1. What is Computer Security?

Computer security is a general term that covers a wide area of computing and information processing. Industries that depend on computer systems and networks to conduct daily business transactions and access critical information regard their data as an important part of their overall assets. Several terms and metrics have entered our daily business vocabulary, such as total cost of ownership (TCO), return on investment (ROI), and quality of service (QoS). Using these metrics, industries can calculate aspects such as data integrity and high-availability (HA) as part of their planning and process management costs. In some industries, such as electronic commerce, the availability and trustworthiness of data can mean the difference between success and failure.

1.1.1.1. How did Computer Security come about?

Information security has evolved over the years due to the increasing reliance on public networks not to disclose personal, financial, and other restricted information. There are numerous instances such as the Mitnick ^[1] and the Vladimir Levin ^[2] cases that prompted organizations across all industries to re-think the way they handle information, including its transmission and disclosure. The popularity of the Internet was one of the most important developments that prompted an intensified effort in data security.

An ever-growing number of people are using their personal computers to gain access to the resources that the Internet has to offer. From research and information retrieval to electronic mail and commerce transactions, the Internet has been regarded as one of the most important developments of the 20th century.

The Internet and its earlier protocols, however, were developed as a *trust-based* system. That is, the Internet Protocol (IP) was not designed to be secure in itself. There are no approved security standards built into the TCP/IP communications stack, leaving it open to potentially malicious users and processes across the network. Modern developments have made Internet communication more secure, but there are still several incidents that gain national attention and alert us to the fact that nothing is completely safe.

1.1.1.2. Security Today

In February of 2000, a Distributed Denial of Service (DDoS) attack was unleashed on several of the most heavily-trafficked sites on the Internet. The attack rendered yahoo.com, cnn.com, amazon.com, fbi.gov, and several other sites completely unreachable to normal users, as it tied up routers for several hours with large-byte ICMP packet transfers, also called a *ping flood*. The attack was brought on by unknown assailants using specially created, widely available programs that scanned vulnerable network servers, installed client applications called *Trojans* on the servers, and timed an attack with every infected server flooding the victim sites and rendering them unavailable. Many blame the attack on fundamental flaws in the way routers and the protocols used are structured to accept all incoming data, no matter where or for what purpose the packets are sent.

In 2007, a data breach exploiting the widely-known weaknesses of the Wired Equivalent Privacy (WEP) wireless encryption protocol resulted in the theft from a global financial institution of over 45 million credit card numbers. ^[3]

In a separate incident, the billing records of over 2.2 million patients stored on a backup tape were stolen from the front seat of a courier's car. ^[4]

Currently, an estimated 1.4 billion people use or have used the Internet worldwide. ^[5] At the same time:

- On any given day, there are approximately 225 major incidences of security breach reported to the CERT Coordination Center at Carnegie Mellon University. ^[6]
- The number of CERT reported incidences jumped from 52,658 in 2001, 82,094 in 2002 and to 137,529 in 2003. ^[7]
- According to the FBI, computer-related crimes cost US businesses \$67.2 Billion dollars in 2006. ^[8]

From a 2009 global survey of security and information technology professionals, "Why Security Matters Now" ^[9], undertaken by *CIO Magazine*, some notable results are:

- Just 23% of respondents have policies for using Web 2.0 technologies. These technologies, such as Twitter, Facebook and LinkedIn may provide a convenient way for companies and individuals to communicate and collaborate, however they open new vulnerabilities, primarily the leaking of confidential data.
- Even during the recent financial crisis of 2009, security budgets were found in the survey to be mostly at the same amount or increasing over previous years (nearly 2 out of 3 respondents expect spending to increase or remain the same). This is good news and reflects the importance that organizations are placing on information security today.

These results enforce the reality that computer security has become a quantifiable and justifiable expense for IT budgets. Organizations that require data integrity and high availability elicit the skills of system administrators, developers, and engineers to ensure 24x7 reliability of their systems, services, and information. Falling victim to malicious users, processes, or coordinated attacks is a direct threat to the success of the organization.

Unfortunately, system and network security can be a difficult proposition, requiring an intricate knowledge of how an organization regards, uses, manipulates, and transmits its information. Understanding the way an organization (and the people who make up the organization) conducts business is paramount to implementing a proper security plan.

1.1.1.3. Standardizing Security

Enterprises in every industry rely on regulations and rules that are set by standards-making bodies

such as the American Medical Association (AMA) or the Institute of Electrical and Electronics Engineers (IEEE). The same ideals hold true for information security. Many security consultants and vendors agree upon the standard security model known as CIA, or *Confidentiality, Integrity, and Availability*. This three-tiered model is a generally accepted component to assessing risks of sensitive information and establishing security policy. The following describes the CIA model in further detail:

- Confidentiality — Sensitive information must be available only to a set of pre-defined individuals. Unauthorized transmission and usage of information should be restricted. For example, confidentiality of information ensures that a customer's personal or financial information is not obtained by an unauthorized individual for malicious purposes such as identity theft or credit fraud.
- Integrity — Information should not be altered in ways that render it incomplete or incorrect. Unauthorized users should be restricted from the ability to modify or destroy sensitive information.
- Availability — Information should be accessible to authorized users any time that it is needed. Availability is a warranty that information can be obtained with an agreed-upon frequency and timeliness. This is often measured in terms of percentages and agreed to formally in Service Level Agreements (SLAs) used by network service providers and their enterprise clients.

1.1.2. SELinux

Red Hat Enterprise Linux includes an enhancement to the Linux kernel called SELinux, which implements a Mandatory Access Control (MAC) architecture that provides a fine-grained level of control over files, processes, users and applications in the system. Detailed discussion of SELinux is beyond the scope of this document; however, for more information on SELinux and its use in Red Hat Enterprise Linux, refer to the Red Hat Enterprise Linux SELinux User Guide. For more information on configuring and running services that are protected by SELinux, refer to the SELinux Managing Confined Services Guide. Other available resources for SELinux are listed in [Chapter 8, References](#).

1.1.3. Security Controls

Computer security is often divided into three distinct master categories, commonly referred to as *controls*:

- Physical
- Technical
- Administrative

These three broad categories define the main objectives of proper security implementation. Within these controls are sub-categories that further detail the controls and how to implement them.

1.1.3.1. Physical Controls

Physical control is the implementation of security measures in a defined structure used to deter or prevent unauthorized access to sensitive material. Examples of physical controls are:

- Closed-circuit surveillance cameras
- Motion or thermal alarm systems
- Security guards
- Picture IDs
- Locked and dead-bolted steel doors
- Biometrics (includes fingerprint, voice, face, iris, handwriting, and other automated methods used to recognize individuals)

1.1.3.2. Technical Controls

Technical controls use technology as a basis for controlling the access and usage of sensitive data

throughout a physical structure and over a network. Technical controls are far-reaching in scope and encompass such technologies as:

- Encryption
- Smart cards
- Network authentication
- Access control lists (ACLs)
- File integrity auditing software

1.1.3.3. Administrative Controls

Administrative controls define the human factors of security. They involve all levels of personnel within an organization and determine which users have access to what resources and information by such means as:

- Training and awareness
- Disaster preparedness and recovery plans
- Personnel recruitment and separation strategies
- Personnel registration and accounting

1.1.4. Conclusion

Now that you have learned about the origins, reasons, and aspects of security, you will find it easier to determine the appropriate course of action with regard to Red Hat Enterprise Linux. It is important to know what factors and conditions make up security in order to plan and implement a proper strategy. With this information in mind, the process can be formalized and the path becomes clearer as you delve deeper into the specifics of the security process.

1.2. Vulnerability Assessment

Given time, resources, and motivation, an attacker can break into nearly any system. All of the security procedures and technologies currently available cannot guarantee that any systems are completely safe from intrusion. Routers help secure gateways to the Internet. Firewalls help secure the edge of the network. Virtual Private Networks safely pass data in an encrypted stream. Intrusion detection systems warn you of malicious activity. However, the success of each of these technologies is dependent upon a number of variables, including:

- The expertise of the staff responsible for configuring, monitoring, and maintaining the technologies.
- The ability to patch and update services and kernels quickly and efficiently.
- The ability of those responsible to keep constant vigilance over the network.

Given the dynamic state of data systems and technologies, securing corporate resources can be quite complex. Due to this complexity, it is often difficult to find expert resources for all of your systems. While it is possible to have personnel knowledgeable in many areas of information security at a high level, it is difficult to retain staff who are experts in more than a few subject areas. This is mainly because each subject area of information security requires constant attention and focus. Information security does not stand still.

1.2.1. Thinking Like the Enemy

Suppose that you administer an enterprise network. Such networks commonly comprise operating systems, applications, servers, network monitors, firewalls, intrusion detection systems, and more. Now imagine trying to keep current with each of those. Given the complexity of today's software and networking environments, exploits and bugs are a certainty. Keeping current with patches and updates

for an entire network can prove to be a daunting task in a large organization with heterogeneous systems.

Combine the expertise requirements with the task of keeping current, and it is inevitable that adverse incidents occur, systems are breached, data is corrupted, and service is interrupted.

To augment security technologies and aid in protecting systems, networks, and data, you must think like a cracker and gauge the security of your systems by checking for weaknesses. Preventative vulnerability assessments against your own systems and network resources can reveal potential issues that can be addressed before a cracker exploits it.

A vulnerability assessment is an internal audit of your network and system security; the results of which indicate the confidentiality, integrity, and availability of your network (as explained in [Section 1.1.1.3, “Standardizing Security”](#)). Typically, vulnerability assessment starts with a reconnaissance phase, during which important data regarding the target systems and resources is gathered. This phase leads to the system readiness phase, whereby the target is essentially checked for all known vulnerabilities. The readiness phase culminates in the reporting phase, where the findings are classified into categories of high, medium, and low risk; and methods for improving the security (or mitigating the risk of vulnerability) of the target are discussed.

If you were to perform a vulnerability assessment of your home, you would likely check each door to your home to see if they are closed and locked. You would also check every window, making sure that they closed completely and latch correctly. This same concept applies to systems, networks, and electronic data. Malicious users are the thieves and vandals of your data. Focus on their tools, mentality, and motivations, and you can then react swiftly to their actions.

1.2.2. Defining Assessment and Testing

Vulnerability assessments may be broken down into one of two types: *outside looking in* and *inside looking around*.

When performing an outside-looking-in vulnerability assessment, you are attempting to compromise your systems from the outside. Being external to your company provides you with the cracker's viewpoint. You see what a cracker sees — publicly-routable IP addresses, systems on your *DMZ*, external interfaces of your firewall, and more. DMZ stands for “demilitarized zone”, which corresponds to a computer or small subnetwork that sits between a trusted internal network, such as a corporate private LAN, and an untrusted external network, such as the public Internet. Typically, the DMZ contains devices accessible to Internet traffic, such as Web (HTTP) servers, FTP servers, SMTP (e-mail) servers and DNS servers.

When you perform an inside-looking-around vulnerability assessment, you are at an advantage since you are internal and your status is elevated to trusted. This is the viewpoint you and your co-workers have once logged on to your systems. You see print servers, file servers, databases, and other resources.

There are striking distinctions between the two types of vulnerability assessments. Being internal to your company gives you more privileges than an outsider. In most organizations, security is configured to keep intruders out. Very little is done to secure the internals of the organization (such as departmental firewalls, user-level access controls, and authentication procedures for internal resources). Typically, there are many more resources when looking around inside as most systems are internal to a company. Once you are outside the company, your status is untrusted. The systems and resources available to you externally are usually very limited.

Consider the difference between vulnerability assessments and *penetration tests*. Think of a vulnerability assessment as the first step to a penetration test. The information gleaned from the assessment is used for testing. Whereas the assessment is undertaken to check for holes and potential vulnerabilities, the penetration testing actually attempts to exploit the findings.

Assessing network infrastructure is a dynamic process. Security, both information and physical, is dynamic. Performing an assessment shows an overview, which can turn up false positives and false negatives.

Security administrators are only as good as the tools they use and the knowledge they retain. Take any of the assessment tools currently available, run them against your system, and it is almost a guarantee that there are some false positives. Whether by program fault or user error, the result is the same. The tool may find vulnerabilities which in reality do not exist (false positive); or, even worse, the tool may not find vulnerabilities that actually do exist (false negative).

Now that the difference between a vulnerability assessment and a penetration test is defined, take the findings of the assessment and review them carefully before conducting a penetration test as part of your new best practices approach.



Warning

Attempting to exploit vulnerabilities on production resources can have adverse effects to the productivity and efficiency of your systems and network.

The following list examines some of the benefits to performing vulnerability assessments.

- Creates proactive focus on information security.
- Finds potential exploits before crackers find them.
- Results in systems being kept up to date and patched.
- Promotes growth and aids in developing staff expertise.
- Abates financial loss and negative publicity.

1.2.2.1. Establishing a Methodology

To aid in the selection of tools for a vulnerability assessment, it is helpful to establish a vulnerability assessment methodology. Unfortunately, there is no predefined or industry approved methodology at this time; however, common sense and best practices can act as a sufficient guide.

What is the target? Are we looking at one server, or are we looking at our entire network and everything within the network? Are we external or internal to the company? The answers to these questions are important as they help determine not only which tools to select but also the manner in which they are used.

To learn more about establishing methodologies, refer to the following websites:

- <http://www.isecom.org/osstmm/> *The Open Source Security Testing Methodology Manual (OSSTMM)*
- <http://www.owasp.org/> *The Open Web Application Security Project*

1.2.3. Evaluating the Tools

An assessment can start by using some form of an information gathering tool. When assessing the entire network, map the layout first to find the hosts that are running. Once located, examine each host individually. Focusing on these hosts requires another set of tools. Knowing which tools to use may be the most crucial step in finding vulnerabilities.

Just as in any aspect of everyday life, there are many different tools that perform the same job. This concept applies to performing vulnerability assessments as well. There are tools specific to operating systems, applications, and even networks (based on the protocols used). Some tools are free; others

are not. Some tools are intuitive and easy to use, while others are cryptic and poorly documented but have features that other tools do not.

Finding the right tools may be a daunting task and in the end, experience counts. If possible, set up a test lab and try out as many tools as you can, noting the strengths and weaknesses of each. Review the README file or man page for the tool. Additionally, look to the Internet for more information, such as articles, step-by-step guides, or even mailing lists specific to a tool.

The tools discussed below are just a small sampling of the available tools.

1.2.3.1. Scanning Hosts with Nmap

Nmap is a popular tool that can be used to determine the layout of a network. Nmap has been available for many years and is probably the most often used tool when gathering information. An excellent manual page is included that provides detailed descriptions of its options and usage. Administrators can use Nmap on a network to find host systems and open ports on those systems.

Nmap is a competent first step in vulnerability assessment. You can map out all the hosts within your network and even pass an option that allows Nmap to attempt to identify the operating system running on a particular host. Nmap is a good foundation for establishing a policy of using secure services and restricting unused services.

To install Nmap, run the **yum install nmap** command as the root user.

1.2.3.1.1. Using Nmap

Nmap can be run from a shell prompt by typing the **nmap** command followed by the hostname or IP address of the machine to scan:

```
nmap <hostname>
```

For example, to scan a machine with hostname **foo.example.com**, type the following at a shell prompt:

```
~]$ nmap foo.example.com
```

The results of a basic scan (which could take up to a few minutes, depending on where the host is located and other network conditions) look similar to the following:

```
Interesting ports on foo.example.com:
Not shown: 1710 filtered ports
PORT      STATE SERVICE
22/tcp    open  ssh
53/tcp    open  domain
80/tcp    open  http
113/tcp   closed auth
```

Nmap tests the most common network communication ports for listening or waiting services. This knowledge can be helpful to an administrator who wants to close down unnecessary or unused services.

For more information about using Nmap, refer to the official homepage at the following URL:

<http://www.insecure.org/>

1.2.3.2. Nessus

Nessus is a full-service security scanner. The plug-in architecture of Nessus allows users to customize it for their systems and networks. As with any scanner, Nessus is only as good as the signature database it relies upon. Fortunately, Nessus is frequently updated and features full reporting, host scanning, and real-time vulnerability searches. Remember that there could be false positives and false negatives, even in a tool as powerful and as frequently updated as Nessus.



Note

The Nessus client and server software requires a subscription to use. It has been included in this document as a reference to users who may be interested in using this popular application.

For more information about Nessus, refer to the official website at the following URL:

<http://www.nessus.org/>

1.2.3.3. Nikto

Nikto is an excellent common gateway interface (CGI) script scanner. Nikto not only checks for CGI vulnerabilities but does so in an evasive manner, so as to elude intrusion detection systems. It comes with thorough documentation which should be carefully reviewed prior to running the program. If you have Web servers serving up CGI scripts, Nikto can be an excellent resource for checking the security of these servers.

More information about Nikto can be found at the following URL:

<http://cirt.net/nikto2>

1.2.3.4. Anticipating Your Future Needs

Depending upon your target and resources, there are many tools available. There are tools for wireless networks, Novell networks, Windows systems, Linux systems, and more. Another essential part of performing assessments may include reviewing physical security, personnel screening, or voice/PBX network assessment. New concepts, such as *war walking* and *wardriving*, which involves scanning the perimeter of your enterprise's physical structures for wireless network vulnerabilities, are some concepts that you should investigate and, if needed, incorporate into your assessments. Imagination and exposure are the only limits of planning and conducting vulnerability assessments.

1.3. Attackers and Vulnerabilities

To plan and implement a good security strategy, first be aware of some of the issues which determined, motivated attackers exploit to compromise systems. However, before detailing these issues, the terminology used when identifying an attacker must be defined.

1.3.1. A Quick History of Hackers

The modern meaning of the term *hacker* has origins dating back to the 1960s and the Massachusetts Institute of Technology (MIT) Tech Model Railroad Club, which designed train sets of large scale and intricate detail. Hacker was a name used for club members who discovered a clever trick or workaround for a problem.

The term hacker has since come to describe everything from computer buffs to gifted programmers. A common trait among most hackers is a willingness to explore in detail how computer systems and networks function with little or no outside motivation. Open source software developers often consider themselves and their colleagues to be hackers, and use the word as a term of respect.

Typically, hackers follow a form of the *hacker ethic* which dictates that the quest for information and expertise is essential, and that sharing this knowledge is the hackers duty to the community. During this quest for knowledge, some hackers enjoy the academic challenges of circumventing security controls on computer systems. For this reason, the press often uses the term hacker to describe those who illicitly access systems and networks with unscrupulous, malicious, or criminal intent. The more accurate term for this type of computer hacker is *cracker* — a term created by hackers in the mid-1980s to differentiate the two communities.

1.3.1.1. Shades of Gray

Within the community of individuals who find and exploit vulnerabilities in systems and networks are several distinct groups. These groups are often described by the shade of hat that they "wear" when performing their security investigations and this shade is indicative of their intent.

The *white hat hacker* is one who tests networks and systems to examine their performance and determine how vulnerable they are to intrusion. Usually, white hat hackers crack their own systems or the systems of a client who has specifically employed them for the purposes of security auditing. Academic researchers and professional security consultants are two examples of white hat hackers.

A *black hat hacker* is synonymous with a cracker. In general, crackers are less focused on programming and the academic side of breaking into systems. They often rely on available cracking programs and exploit well known vulnerabilities in systems to uncover sensitive information for personal gain or to inflict damage on the target system or network.

The *gray hat hacker*, on the other hand, has the skills and intent of a white hat hacker in most situations but uses his knowledge for less than noble purposes on occasion. A gray hat hacker can be thought of as a white hat hacker who wears a black hat at times to accomplish his own agenda.

Gray hat hackers typically subscribe to another form of the hacker ethic, which says it is acceptable to break into systems as long as the hacker does not commit theft or breach confidentiality. Some would argue, however, that the act of breaking into a system is in itself unethical.

Regardless of the intent of the intruder, it is important to know the weaknesses a cracker may likely attempt to exploit. The remainder of the chapter focuses on these issues.

1.3.2. Threats to Network Security

Bad practices when configuring the following aspects of a network can increase the risk of attack.

1.3.2.1. Insecure Architectures

A misconfigured network is a primary entry point for unauthorized users. Leaving a trust-based, open local network vulnerable to the highly-insecure Internet is much like leaving a door ajar in a crime-ridden neighborhood — nothing may happen for an arbitrary amount of time, but *eventually* someone exploits the opportunity.

1.3.2.1.1. Broadcast Networks

System administrators often fail to realize the importance of networking hardware in their security schemes. Simple hardware such as hubs and routers rely on the broadcast or non-switched principle; that is, whenever a node transmits data across the network to a recipient node, the hub or router sends a broadcast of the data packets until the recipient node receives and processes the data. This method is the most vulnerable to address resolution protocol (*ARP*) or media access control (*MAC*) address spoofing by both outside intruders and unauthorized users on local hosts.

1.3.2.1.2. Centralized Servers

Another potential networking pitfall is the use of centralized computing. A common cost-cutting measure

for many businesses is to consolidate all services to a single powerful machine. This can be convenient as it is easier to manage and costs considerably less than multiple-server configurations. However, a centralized server introduces a single point of failure on the network. If the central server is compromised, it may render the network completely useless or worse, prone to data manipulation or theft. In these situations, a central server becomes an open door which allows access to the entire network.

1.3.3. Threats to Server Security

Server security is as important as network security because servers often hold a great deal of an organization's vital information. If a server is compromised, all of its contents may become available for the cracker to steal or manipulate at will. The following sections detail some of the main issues.

1.3.3.1. Unused Services and Open Ports

A full installation of Red Hat Enterprise Linux 6 contains 1000+ application and library packages. However, most server administrators do not opt to install every single package in the distribution, preferring instead to install a base installation of packages, including several server applications.

A common occurrence among system administrators is to install the operating system without paying attention to what programs are actually being installed. This can be problematic because unneeded services may be installed, configured with the default settings, and possibly turned on. This can cause unwanted services, such as Telnet, DHCP, or DNS, to run on a server or workstation without the administrator realizing it, which in turn can cause unwanted traffic to the server, or even, a potential pathway into the system for crackers. Refer To [Section 2.2, “Server Security”](#) for information on closing ports and disabling unused services.

1.3.3.2. Unpatched Services

Most server applications that are included in a default installation are solid, thoroughly tested pieces of software. Having been in use in production environments for many years, their code has been thoroughly refined and many of the bugs have been found and fixed.

However, there is no such thing as perfect software and there is always room for further refinement. Moreover, newer software is often not as rigorously tested as one might expect, because of its recent arrival to production environments or because it may not be as popular as other server software.

Developers and system administrators often find exploitable bugs in server applications and publish the information on bug tracking and security-related websites such as the Bugtraq mailing list (<http://www.securityfocus.com>) or the Computer Emergency Response Team (CERT) website (<http://www.cert.org>). Although these mechanisms are an effective way of alerting the community to security vulnerabilities, it is up to system administrators to patch their systems promptly. This is particularly true because crackers have access to these same vulnerability tracking services and will use the information to crack unpatched systems whenever they can. Good system administration requires vigilance, constant bug tracking, and proper system maintenance to ensure a more secure computing environment.

Refer to [Section 1.5 “Security Updates”](#) for more information about keeping a system up-to-date.

1.3.3.3. Inattentive Administration

Administrators who fail to patch their systems are one of the greatest threats to server security. According to the *SysAdmin, Audit, Network, Security Institute* (SANS), the primary cause of computer security vulnerability is to “assign untrained people to maintain security and provide neither the training nor the time to make it possible to do the job.” ^[10] This applies as much to inexperienced administrators as it does to overconfident or amotivated administrators.

Some administrators fail to patch their servers and workstations, while others fail to watch log messages from the system kernel or network traffic. Another common error is when default passwords or keys to services are left unchanged. For example, some databases have default administration passwords because the database developers assume that the system administrator changes these passwords immediately after installation. If a database administrator fails to change this password, even an inexperienced cracker can use a widely-known default password to gain administrative privileges to the database. These are only a few examples of how inattentive administration can lead to compromised servers.

1.3.3.4. Inherently Insecure Services

Even the most vigilant organization can fall victim to vulnerabilities if the network services they choose are inherently insecure. For instance, there are many services developed under the assumption that they are used over trusted networks; however, this assumption fails as soon as the service becomes available over the Internet — which is itself inherently untrusted.

One category of insecure network services are those that require unencrypted usernames and passwords for authentication. Telnet and FTP are two such services. If packet sniffing software is monitoring traffic between the remote user and such a service usernames and passwords can be easily intercepted.

Inherently, such services can also more easily fall prey to what the security industry terms the *man-in-the-middle* attack. In this type of attack, a cracker redirects network traffic by tricking a cracked name server on the network to point to his machine instead of the intended server. Once someone opens a remote session to the server, the attacker's machine acts as an invisible conduit, sitting quietly between the remote service and the unsuspecting user capturing information. In this way a cracker can gather administrative passwords and raw data without the server or the user realizing it.

Another category of insecure services include network file systems and information services such as NFS or NIS, which are developed explicitly for LAN usage but are, unfortunately, extended to include WANs (for remote users). NFS does not, by default, have any authentication or security mechanisms configured to prevent a cracker from mounting the NFS share and accessing anything contained therein. NIS, as well, has vital information that must be known by every computer on a network, including passwords and file permissions, within a plain text ASCII or DBM (ASCII-derived) database. A cracker who gains access to this database can then access every user account on a network, including the administrator's account.

By default, Red Hat Enterprise Linux is released with all such services turned off. However, since administrators often find themselves forced to use these services, careful configuration is critical. Refer to [Section 2.2, “Server Security”](#) for more information about setting up services in a safe manner.

1.3.4. Threats to Workstation and Home PC Security

Workstations and home PCs may not be as prone to attack as networks or servers, but since they often contain sensitive data, such as credit card information, they are targeted by system crackers. Workstations can also be co-opted without the user's knowledge and used by attackers as “slave” machines in coordinated attacks. For these reasons, knowing the vulnerabilities of a workstation can save users the headache of reinstalling the operating system, or worse, recovering from data theft.

1.3.4.1. Bad Passwords

Bad passwords are one of the easiest ways for an attacker to gain access to a system. For more on how to avoid common pitfalls when creating a password, refer to [Section 2.1.3, “Password Security”](#).

1.3.4.2. Vulnerable Client Applications

Although an administrator may have a fully secure and patched server, that does not mean remote users

are secure when accessing it. For instance, if the server offers Telnet or FTP services over a public network, an attacker can capture the plain text usernames and passwords as they pass over the network, and then use the account information to access the remote user's workstation.

Even when using secure protocols, such as SSH, a remote user may be vulnerable to certain attacks if they do not keep their client applications updated. For instance, v.1 SSH clients are vulnerable to an X-forwarding attack from malicious SSH servers. Once connected to the server, the attacker can quietly capture any keystrokes and mouse clicks made by the client over the network. This problem was fixed in the v.2 SSH protocol, but it is up to the user to keep track of what applications have such vulnerabilities and update them as necessary.

[Section 2.1, “Workstation Security”](#) discusses in more detail what steps administrators and home users should take to limit the vulnerability of computer workstations.

1.4. Common Exploits and Attacks

[Table 1.1, “Common Exploits”](#) details some of the most common exploits and entry points used by intruders to access organizational network resources. Key to these common exploits are the explanations of how they are performed and how administrators can properly safeguard their network against such attacks.

Table 1.1. Common Exploits

Exploit	Description	Notes
Null or Default Passwords	Leaving administrative passwords blank or using a default password set by the product vendor. This is most common in hardware such as routers and firewalls, but some services that run on Linux can contain default administrator passwords as well (though Red Hat Enterprise Linux does not ship with them).	Commonly associated with networking hardware such as routers, firewalls, VPNs, and network attached storage (NAS) appliances. Common in many legacy operating systems, especially those that bundle services (such as UNIX and Windows.) Administrators sometimes create privileged user accounts in a rush and leave the password null, creating a perfect entry point for malicious users who discover the account.
Default Shared Keys	Secure services sometimes package default security keys for development or evaluation testing purposes. If these keys are left unchanged and are placed in a production environment on the Internet, <i>all</i> users with the same default keys have access to that shared-key resource, and any sensitive information that it contains.	Most common in wireless access points and preconfigured secure server appliances.
IP Spoofing	A remote machine acts as a node on your local network, finds vulnerabilities with your servers, and installs a backdoor program or trojan horse to gain control over your network resources.	Spoofing is quite difficult as it involves the attacker predicting TCP/IP sequence numbers to coordinate a connection to target systems, but several tools are available to assist crackers in performing such a vulnerability. Depends on target system running services (such as rsh, telnet, FTP and others) that use <i>source-based</i> authentication techniques, which are not recommended when compared to PKI or other forms of encrypted authentication used in ssh or SSL/TLS.
Eavesdropping	Collecting data that passes between two active nodes on a network by eavesdropping on the connection between the two nodes.	This type of attack works mostly with plain text transmission protocols such as Telnet, FTP, and HTTP transfers. Remote attacker must have access to a compromised system on a LAN in order to perform such an attack; usually the cracker has used an active attack (such as IP spoofing or man-in-

		the-middle) to compromise a system on the LAN. Preventative measures include services with cryptographic key exchange, one-time passwords, or encrypted authentication to prevent password snooping; strong encryption during transmission is also advised.
Service Vulnerabilities	An attacker finds a flaw or loophole in a service run over the Internet; through this vulnerability, the attacker compromises the entire system and any data that it may hold, and could possibly compromise other systems on the network.	HTTP-based services such as CGI are vulnerable to remote command execution and even interactive shell access. Even if the HTTP service runs as a non-privileged user such as "nobody", information such as configuration files and network maps can be read, or the attacker can start a denial of service attack which drains system resources or renders it unavailable to other users. Services sometimes can have vulnerabilities that go unnoticed during development and testing; these vulnerabilities (such as <i>buffer overflows</i>, where attackers crash a service using arbitrary values that fill the memory buffer of an application, giving the attacker an interactive command prompt from which they may execute arbitrary commands) can give complete administrative control to an attacker. Administrators should make sure that services do not run as the root user, and should stay vigilant of patches and errata updates for applications from vendors or security organizations such as CERT and CVE.
Application Vulnerabilities	Attackers find faults in desktop and workstation applications (such as e-mail clients) and execute arbitrary code, implant trojan horses for future compromise, or crash systems. Further exploitation can occur if the compromised workstation has administrative privileges on the rest of the network.	Workstations and desktops are more prone to exploitation as workers do not have the expertise or experience to prevent or detect a compromise; it is imperative to inform individuals of the risks they are taking when they install unauthorized software or open unsolicited email attachments. Safeguards can be implemented such that email client software does not automatically open or execute

		attachments. Additionally, the automatic update of workstation software via Red Hat Network or other system management services can alleviate the burdens of multi-seat security deployments.
Denial of Service (DoS) Attacks	Attacker or group of attackers coordinate against an organization's network or server resources by sending unauthorized packets to the target host (either server, router, or workstation). This forces the resource to become unavailable to legitimate users.	The most reported DoS case in the US occurred in 2000. Several highly-trafficked commercial and government sites were rendered unavailable by a coordinated ping flood attack using several compromised systems with high bandwidth connections acting as <i>zombies</i>, or redirected broadcast nodes. Source packets are usually forged (as well as rebroadcasted), making investigation as to the true source of the attack difficult. Advances in ingress filtering (IETF rfc2267) using iptables and Network Intrusion Detection Systems such as snort assist administrators in tracking down and preventing distributed DoS attacks.

1.5. Security Updates

As security vulnerabilities are discovered, the affected software must be updated in order to limit any potential security risks. If the software is part of a package within a Red Hat Enterprise Linux distribution that is currently supported, Red Hat is committed to releasing updated packages that fix the vulnerability as soon as is possible. Often, announcements about a given security exploit are accompanied with a patch (or source code that fixes the problem). This patch is then applied to the Red Hat Enterprise Linux package and tested and released as an errata update. However, if an announcement does not include a patch, a developer first works with the maintainer of the software to fix the problem. Once the problem is fixed, the package is tested and released as an errata update.

If an errata update is released for software used on your system, it is highly recommended that you update the affected packages as soon as possible to minimize the amount of time the system is potentially vulnerable.

1.5.1. Updating Packages

When updating software on a system, it is important to download the update from a trusted source. An attacker can easily rebuild a package with the same version number as the one that is supposed to fix the problem but with a different security exploit and release it on the Internet. If this happens, using security measures such as verifying files against the original RPM does not detect the exploit. Thus, it is very important to only download RPMs from trusted sources, such as from Red Hat and to check the signature of the package to verify its integrity.

**Note**

Red Hat Enterprise Linux includes a convenient panel icon that displays visible alerts when there is an update available.

1.5.2. Verifying Signed Packages

All Red Hat Enterprise Linux packages are signed with the Red Hat GPG key. GPG stands for GNU Privacy Guard, or GnuPG, a free software package used for ensuring the authenticity of distributed files. For example, a private key (secret key) locks the package while the public key unlocks and verifies the package. If the public key distributed by Red Hat Enterprise Linux does not match the private key during RPM verification, the package may have been altered and therefore cannot be trusted.

The RPM utility within Red Hat Enterprise Linux 6 automatically tries to verify the GPG signature of an RPM package before installing it. If the Red Hat GPG key is not installed, install it from a secure, static location, such as a Red Hat installation CD-ROM or DVD.

Assuming the disc is mounted in `/mnt/cdrom`, use the following command as the root user to import it into the *keyring* (a database of trusted keys on the system):

```
~]# rpm --import /mnt/cdrom/RPM-GPG-KEY
```

Now, the Red Hat GPG key is located in the `/etc/pki/rpm-gpg/` directory.

To display a list of all keys installed for RPM verification, execute the following command:

```
~]# rpm -qa gpg-pubkey*
gpg-pubkey-db42a60e-37ea5438
```

To display details about a specific key, use the `rpm -qi` command followed by the output from the previous command, as in this example:

```
~]# rpm -qi gpg-pubkey-db42a60e-37ea5438
Name           : gpg-pubkey                Relocations: (not relocatable)
Version        : 2fa658e0                 Vendor: (none)
Release        : 45700c69                 Build Date: Fri 07 Oct 2011 02:04:51
PM CEST
Install Date:  Fri 07 Oct 2011 02:04:51 PM CEST    Build Host: localhost
Group          : Public Keys               Source RPM: (none)
[output truncated]
```

It is extremely important to verify the signature of the RPM files before installing them to ensure that they have not been altered from the original source of the packages. To verify all the downloaded packages at once, issue the following command:

```
~]# rpm -K /root/updates/*.rpm
alsa-lib-1.0.22-3.el6.x86_64.rpm: rsa sha1 (md5) pgp md5 OK
alsa-utils-1.0.21-3.el6.x86_64.rpm: rsa sha1 (md5) pgp md5 OK
aspell-0.60.6-12.el6.x86_64.rpm: rsa sha1 (md5) pgp md5 OK
```

For each package, if the GPG key verifies successfully, the command returns **pgp OK**. If it doesn't, make sure you are using the correct Red Hat public key, as well as verifying the source of the content. Packages that do not pass GPG verification should not be installed, as they may have been altered by a third party.

After verifying the GPG key and downloading all the packages associated with the errata report, install the packages as root at a shell prompt.

Alternatively, you may use the Yum utility to verify signed packages. Yum provides secure package management by enabling GPG signature verification on GPG-signed packages to be turned on for all package repositories (that is, package sources), or for individual repositories. When signature verification is enabled, Yum will refuse to install any packages not GPG-signed with the correct key for that repository. This means that you can trust that the RPM packages you download and install on your system are from a trusted source, such as Red Hat, and were not modified during transfer.

In order to have automatic GPG signature verification enabled when installing or updating packages via Yum, ensure you have the following option defined under the **[main]** section of your **/etc/yum.conf** file:

```
gpgcheck=1
```

1.5.3. Installing Signed Packages

Installation for most packages can be done safely (except kernel packages) by issuing the following command as root:

```
rpm -Uvh <package>...
```

For example, to install all packages in the **/tmp/updates/** directory, run:

```
~]# rpm -Uvh /root/updates/*.rpm
Preparing...                               ##### [100%]
 1:alsa-lib                               ##### [ 33%]
 2:alsa-utils                             ##### [ 67%]
 3:aspell                                 ##### [100%]
```

For kernel packages, as root use the command in the following form:

```
rpm -ivh <kernel-package>
```

For example, to install *kernel-2.6.32-220.el6.x86_64.rpm*, type the following at a shell prompt:

```
~]# rpm -ivh /tmp/updates/kernel-2.6.32-220.el6.x86_64.rpm
Preparing...                               ##### [100%]
 1:kernel                                 ##### [100%]
```

Once the machine has been safely rebooted using the new kernel, the old kernel may be removed using the following command:

```
rpm -e <old-kernel-package>
```

For instance, to remove *kernel-2.6.32-206.el6.x86_64*, type:

```
~]# rpm -e kernel-2.6.32-206.el6.x86_64
```

Alternatively, to install packages with Yum, run, as root, the following command:

```
~]# yum install kernel-2.6.32-220.el6.x86_64.rpm
```

To install local packages with Yum, run, as root, the following command:

```
~]# yum localinstall /root/updates/emacs-23.1-21.el6_2.3.x86_64.rpm
```



Note

It is not a requirement that the old kernel be removed. The default boot loader, GRUB, allows for multiple kernels to be installed, then chosen from a menu at boot time.



Important

Before installing any security errata, be sure to read any special instructions contained in the errata report and execute them accordingly. Refer to [Section 1.5.4, “Applying the Changes”](#) for general instructions about applying the changes made by an errata update.

1.5.4. Applying the Changes

After downloading and installing security errata and updates, it is important to halt usage of the older software and begin using the new software. How this is done depends on the type of software that has been updated. The following list itemizes the general categories of software and provides instructions for using the updated versions after a package upgrade.



Note

In general, rebooting the system is the surest way to ensure that the latest version of a software package is used; however, this option is not always required, or available to the system administrator.

Applications

User-space applications are any programs that can be initiated by a system user. Typically, such applications are used only when a user, script, or automated task utility launches them and they do not persist for long periods of time.

Once such a user-space application is updated, halt any instances of the application on the system and launch the program again to use the updated version.

Kernel

The kernel is the core software component for the Red Hat Enterprise Linux operating system. It manages access to memory, the processor, and peripherals as well as schedules all tasks.

Because of its central role, the kernel cannot be restarted without also stopping the computer. Therefore, an updated version of the kernel cannot be used until the system is rebooted.

Shared Libraries

Shared libraries are units of code, such as **glibc**, which are used by a number of applications and services. Applications utilizing a shared library typically load the shared code when the application is initialized, so any applications using the updated library must be halted and relaunched.

To determine which running applications link against a particular library, use the **ldsof** command:

```
ldsof <path>
```

For example, to determine which running applications link against the **libwrap.so** library, type:

```
~]# ldsof /lib64/libwrap.so*
COMMAND      PID      USER    FD    TYPE  DEVICE  SIZE/OFF  NODE  NAME
sshd          13600   root    mem    REG   253,0    43256  400501
/lib64/libwrap.so.0.7.6
sshd          13603   juan    mem    REG   253,0    43256  400501
/lib64/libwrap.so.0.7.6
gnome-set     14898   juan    mem    REG   253,0    43256  400501
/lib64/libwrap.so.0.7.6
metacity      14925   juan    mem    REG   253,0    43256  400501
/lib64/libwrap.so.0.7.6
[output truncated]
```

This command returns a list of all the running programs which use TCP wrappers for host access control. Therefore, any program listed must be halted and relaunched if the **tcp_wrappers** package is updated.

SysV Services

SysV services are persistent server programs launched during the boot process. Examples of SysV services include **sshd**, **vsftpd**, and **xinetd**.

Because these programs usually persist in memory as long as the machine is booted, each updated SysV service must be halted and relaunched after the package is upgraded. This can be done using the **Services Configuration Tool** or by logging into a root shell prompt and issuing the **/sbin/service** command:

```
/sbin/service <service-name> restart
```

Replace **<service-name>** with the name of the service, such as **sshd**.

xinetd Services

Services controlled by the **xinetd** super service only run when there is an active connection. Examples of services controlled by **xinetd** include Telnet, IMAP, and POP3.

Because new instances of these services are launched by **xinetd** each time a new request is received, connections that occur after an upgrade are handled by the updated software. However, if there are active connections at the time the **xinetd** controlled service is upgraded, they are serviced by the older version of the software.

To kill off older instances of a particular **xinetd** controlled service, upgrade the package for the service then halt all processes currently running. To determine if the process is running, use the **ps** or **pgrep** command and then use the **kill** or **killall** command to halt current instances of the service.

For example, if security errata **imap** packages are released, upgrade the packages, then type the following command as root into a shell prompt:

```
~]# pgrep -l imap
1439 imapd
1788 imapd
1793 imapd
```

This command returns all active IMAP sessions. Individual sessions can then be terminated by issuing the following command as root:

```
kill <PID>
```

If this fails to terminate the session, use the following command instead:

```
kill -9 <PID>
```

In the previous examples, replace **<PID>** with the process identification number (found in the second column of the **pgrep -l** command) for an IMAP session.

To kill all active IMAP sessions, issue the following command:

```
~]# killall imapd
```

[1] <http://law.jrank.org/pages/3791/Kevin-Mitnick-Case-1999.html>

[2] http://www.livinginternet.com/i/ia_hackers_levin.htm

[3] http://www.theregister.co.uk/2007/05/04/tjx_nonfeasance/

[4] <http://www.fudzilla.com/home/item/3178-university-of-utah-loses-22-million-patient-records>

[5] <http://www.internetworldstats.com/stats.htm>

[6] <http://www.cert.org>

[7] http://www.cert.org/stats/cert_stats.html

[8] http://news.cnet.com/Computer-crime-costs-67-billion,-FBI-says/2100-7349_3-6028946.html

[9] http://www.cio.com/article/504837/Why_Security_Matters_Now

[10] <http://www.sans.org/resources/errors.php>

Chapter 2. Securing Your Network

2.1. Workstation Security

Securing a Linux environment begins with the workstation. Whether locking down a personal machine or securing an enterprise system, sound security policy begins with the individual computer. A computer network is only as secure as its weakest node.

2.1.1. Evaluating Workstation Security

When evaluating the security of a Red Hat Enterprise Linux workstation, consider the following:

- *BIOS and Boot Loader Security* — Can an unauthorized user physically access the machine and boot into single user or rescue mode without a password?
- *Password Security* — How secure are the user account passwords on the machine?
- *Administrative Controls* — Who has an account on the system and how much administrative control do they have?
- *Available Network Services* — What services are listening for requests from the network and should they be running at all?
- *Personal Firewalls* — What type of firewall, if any, is necessary?
- *Security Enhanced Communication Tools* — Which tools should be used to communicate between workstations and which should be avoided?

2.1.2. BIOS and Boot Loader Security

Password protection for the BIOS (or BIOS equivalent) and the boot loader can prevent unauthorized users who have physical access to systems from booting using removable media or obtaining root privileges through single user mode. The security measures you should take to protect against such attacks depends both on the sensitivity of the information on the workstation and the location of the machine.

For example, if a machine is used in a trade show and contains no sensitive information, then it may not be critical to prevent such attacks. However, if an employee's laptop with private, unencrypted SSH keys for the corporate network is left unattended at that same trade show, it could lead to a major security breach with ramifications for the entire company.

If the workstation is located in a place where only authorized or trusted people have access, however, then securing the BIOS or the boot loader may not be necessary.

2.1.2.1. BIOS Passwords

The two primary reasons for password protecting the BIOS of a computer are ^[11]:

1. *Preventing Changes to BIOS Settings* — If an intruder has access to the BIOS, they can set it to boot from a diskette or CD-ROM. This makes it possible for them to enter rescue mode or single user mode, which in turn allows them to start arbitrary processes on the system or copy sensitive data.
2. *Preventing System Booting* — Some BIOSes allow password protection of the boot process. When activated, an attacker is forced to enter a password before the BIOS launches the boot loader.

Because the methods for setting a BIOS password vary between computer manufacturers, consult the computer's manual for specific instructions.

If you forget the BIOS password, it can either be reset with jumpers on the motherboard or by

disconnecting the CMOS battery. For this reason, it is good practice to lock the computer case if possible. However, consult the manual for the computer or motherboard before attempting to disconnect the CMOS battery.

2.1.2.1.1. Securing Non-x86 Platforms

Other architectures use different programs to perform low-level tasks roughly equivalent to those of the BIOS on x86 systems. For instance, Intel® Itanium™ computers use the *Extensible Firmware Interface (EFI)* shell.

For instructions on password protecting BIOS-like programs on other architectures, refer to the manufacturer's instructions.

2.1.2.2. Boot Loader Passwords

The primary reasons for password protecting a Linux boot loader are as follows:

1. *Preventing Access to Single User Mode* — If attackers can boot the system into single user mode, they are logged in automatically as root without being prompted for the root password.



Warning

Protecting access to single user mode with a password by editing the **SINGLE** parameter in the `/etc/sysconfig/init` file is not recommended. An attacker can bypass the password by specifying a custom initial command (using the **init=** parameter) on the kernel command line in GRUB. It is recommended to password-protect the GRUB boot loader as specified in [Section 2.1.2.2.1, “Password Protecting GRUB”](#).

2. *Preventing Access to the GRUB Console* — If the machine uses GRUB as its boot loader, an attacker can use the GRUB editor interface to change its configuration or to gather information using the **cat** command.
3. *Preventing Access to Insecure Operating Systems* — If it is a dual-boot system, an attacker can select an operating system at boot time (for example, DOS), which ignores access controls and file permissions.

Red Hat Enterprise Linux 6 ships with the GRUB boot loader on the x86 platform. For a detailed look at GRUB, refer to the Red Hat Installation Guide.

2.1.2.2.1. Password Protecting GRUB

You can configure GRUB to address the first two issues listed in [Section 2.1.2.2, “Boot Loader Passwords”](#) by adding a password directive to its configuration file. To do this, first choose a strong password, open a shell, log in as root, and then type the following command:

```
/sbin/grub-md5-crypt
```

When prompted, type the GRUB password and press **Enter**. This returns an MD5 hash of the password.

Next, edit the GRUB configuration file `/boot/grub/grub.conf`. Open the file and below the **timeout** line in the main section of the document, add the following line:

```
password --md5 <password-hash>
```

Replace **<password-hash>** with the value returned by `/sbin/grub-md5-crypt` ^[12].

The next time the system boots, the GRUB menu prevents access to the editor or command interface without first pressing **p** followed by the GRUB password.

Unfortunately, this solution does not prevent an attacker from booting into an insecure operating system in a dual-boot environment. For this, a different part of the `/boot/grub/grub.conf` file must be edited.

Look for the **title** line of the operating system that you want to secure, and add a line with the **lock** directive immediately beneath it.

For a DOS system, the stanza should begin similar to the following:

```
title DOS
lock
```



Warning

A **password** line must be present in the main section of the `/boot/grub/grub.conf` file for this method to work properly. Otherwise, an attacker can access the GRUB editor interface and remove the lock line.

To create a different password for a particular kernel or operating system, add a **lock** line to the stanza, followed by a password line.

Each stanza protected with a unique password should begin with lines similar to the following example:

```
title DOS
lock
password --md5 <password-hash>
```

2.1.2.2.2. Disabling Interactive Startup

Pressing the **I** key at the beginning of the boot sequence allows you to start up your system interactively. During an interactive startup, the system prompts you to start up each service one by one. However, this may allow an attacker who gains physical access to your system to disable the security-related services and gain access to the system.

To prevent users from starting up the system interactively, as root, disable the **PROMPT** parameter in the `/etc/sysconfig/init` file:

```
PROMPT=no
```

2.1.3. Password Security

Passwords are the primary method that Red Hat Enterprise Linux uses to verify a user's identity. This is why password security is so important for protection of the user, the workstation, and the network.

For security purposes, the installation program configures the system to use *Secure Hash Algorithm 512* (SHA512) and shadow passwords. It is highly recommended that you do not alter these settings.

If shadow passwords are deselected during installation, all passwords are stored as a one-way hash in the world-readable `/etc/passwd` file, which makes the system vulnerable to offline password cracking attacks. If an intruder can gain access to the machine as a regular user, he can copy the `/etc/passwd` file to his own machine and run any number of password cracking programs against it. If there is an

insecure password in the file, it is only a matter of time before the password cracker discovers it.

Shadow passwords eliminate this type of attack by storing the password hashes in the file `/etc/shadow`, which is readable only by the root user.

This forces a potential attacker to attempt password cracking remotely by logging into a network service on the machine, such as SSH or FTP. This sort of brute-force attack is much slower and leaves an obvious trail as hundreds of failed login attempts are written to system files. Of course, if the cracker starts an attack in the middle of the night on a system with weak passwords, the cracker may have gained access before dawn and edited the log files to cover his tracks.

In addition to format and storage considerations is the issue of content. The single most important thing a user can do to protect his account against a password cracking attack is create a strong password.

2.1.3.1. Creating Strong Passwords

When creating a secure password, it is a good idea to follow these guidelines:

- *Do Not Use Only Words or Numbers* — Never use only numbers or words in a password.

Some insecure examples include the following:

- 8675309
- juan
- hackme

- *Do Not Use Recognizable Words* — Words such as proper names, dictionary words, or even terms from television shows or novels should be avoided, even if they are bookended with numbers.

Some insecure examples include the following:

- john1
- DS-9
- mentat123

- *Do Not Use Words in Foreign Languages* — Password cracking programs often check against word lists that encompass dictionaries of many languages. Relying on foreign languages for secure passwords is not secure.

Some insecure examples include the following:

- cheguevara
- bienvenido1
- 1dumbKopf

- *Do Not Use Hacker Terminology* — If you think you are elite because you use hacker terminology — also called l337 (LEET) speak — in your password, think again. Many word lists include LEET speak.

Some insecure examples include the following:

- H4X0R
- 1337

- *Do Not Use Personal Information* — Avoid using any personal information in your passwords. If the attacker knows your identity, the task of deducing your password becomes easier. The following is a list of the types of information to avoid when creating a password:

Some insecure examples include the following:

- Your name
- The names of pets
- The names of family members
- Any birth dates

- Your phone number or zip code
- *Do Not Invert Recognizable Words* — Good password checkers always reverse common words, so inverting a bad password does not make it any more secure. Some insecure examples include the following:
 - R0X4H
 - nauj
 - 9-DS
- *Do Not Write Down Your Password* — Never store a password on paper. It is much safer to memorize it.
- *Do Not Use the Same Password For All Machines* — It is important to make separate passwords for each machine. This way if one system is compromised, all of your machines are not immediately at risk.

The following guidelines will help you to create a strong password:

- *Make the Password at Least Eight Characters Long* — The longer the password, the better. If using MD5 passwords, it should be 15 characters or longer. With DES passwords, use the maximum length (eight characters).
- *Mix Upper and Lower Case Letters* — Red Hat Enterprise Linux is case sensitive, so mix cases to enhance the strength of the password.
- *Mix Letters and Numbers* — Adding numbers to passwords, especially when added to the middle (not just at the beginning or the end), can enhance password strength.
- *Include Non-Alphanumeric Characters* — Special characters such as &, \$, and > can greatly improve the strength of a password (this is not possible if using DES passwords).
- *Pick a Password You Can Remember* — The best password in the world does little good if you cannot remember it; use acronyms or other mnemonic devices to aid in memorizing passwords.

With all these rules, it may seem difficult to create a password that meets all of the criteria for good passwords while avoiding the traits of a bad one. Fortunately, there are some steps you can take to generate an easily-remembered, secure password.

2.1.3.1.1. Secure Password Creation Methodology

There are many methods that people use to create secure passwords. One of the more popular methods involves acronyms. For example:

- Think of an easily-remembered phrase, such as:
"over the river and through the woods, to grandmother's house we go."
- Next, turn it into an acronym (including the punctuation).
otrattw, tghwg.
- Add complexity by substituting numbers and symbols for letters in the acronym. For example, substitute **7** for **t** and the at symbol (**@**) for **a**:
o7r@77w, 7ghwg.
- Add more complexity by capitalizing at least one letter, such as **H**.
o7r@77w, 7gHwg.
- Finally, do not use the example password above for any systems, ever.

While creating secure passwords is imperative, managing them properly is also important, especially for system administrators within larger organizations. The following section details good practices for creating and managing user passwords within an organization.

2.1.3.2. Creating User Passwords Within an Organization

If an organization has a large number of users, the system administrators have two basic options available to force the use of good passwords. They can create passwords for the user, or they can let users create their own passwords, while verifying the passwords are of acceptable quality.

Creating the passwords for the users ensures that the passwords are good, but it becomes a daunting task as the organization grows. It also increases the risk of users writing their passwords down.

For these reasons, most system administrators prefer to have the users create their own passwords, but actively verify that the passwords are good and, in some cases, force users to change their passwords periodically through password aging.

2.1.3.2.1. Forcing Strong Passwords

To protect the network from intrusion it is a good idea for system administrators to verify that the passwords used within an organization are strong ones. When users are asked to create or change passwords, they can use the command line application **passwd**, which is *Pluggable Authentication Modules (PAM)* aware and therefore checks to see if the password is too short or otherwise easy to crack. This check is performed using the **pam_cracklib.so** PAM module. For more information about **pam_cracklib.so** and its options for password checking, refer to the **pam_cracklib(8)** man page. Since PAM is customizable, it is possible to add more password integrity checkers, such as **pam_passwdqc** (available from <http://www.openwall.com/passwdqc/>) or to write a new module. For a list of available PAM modules, refer to http://uw714doc.sco.com/en/SEC_pam/pam-6.html. For more information about PAM, refer to the *Managing Single Sign-On and Smart Cards* guide.

The password check that is performed at the time of their creation does not discover bad passwords as effectively as running a password cracking program against the passwords.

Many password cracking programs are available that run under Red Hat Enterprise Linux, although none ship with the operating system. Below is a brief list of some of the more popular password cracking programs:

- **John The Ripper** — A fast and flexible password cracking program. It allows the use of multiple word lists and is capable of brute-force password cracking. It is available online at <http://www.openwall.com/john/>.
- **Crack** — Perhaps the most well known password cracking software, **Crack** is also very fast, though not as easy to use as **John The Ripper**.
- **Slurpie** — **Slurpie** is similar to **John The Ripper** and **Crack**, but it is designed to run on multiple computers simultaneously, creating a distributed password cracking attack. It can be found along with a number of other distributed attack security evaluation tools online at <http://www.ussrback.com/distributed.htm>.



Warning

Always get authorization in writing before attempting to crack passwords within an organization.

2.1.3.2.2. Passphrases

Passphrases and passwords are the cornerstone to security in most of today's systems. Unfortunately, techniques such as biometrics and two-factor authentication have not yet become mainstream in many systems. If passwords are going to be used to secure a system, then the use of passphrases should be considered. Passphrases are longer than passwords and provide better protection than a password even when implemented with non-standard characters such as numbers and symbols.

2.1.3.2.3. Password Aging

Password aging is another technique used by system administrators to defend against bad passwords within an organization. Password aging means that after a specified period (usually 90 days), the user is prompted to create a new password. The theory behind this is that if a user is forced to change his password periodically, a cracked password is only useful to an intruder for a limited amount of time. The downside to password aging, however, is that users are more likely to write their passwords down.

There are two primary programs used to specify password aging under Red Hat Enterprise Linux: the **chage** command or the graphical **User Manager (system-config-users)** application.



Important

Shadow passwords must be enabled to use the **chage** command. For more information, see the *Red Hat Enterprise Linux 6 Deployment Guide*.

The **-M** option of the **chage** command specifies the maximum number of days the password is valid. For example, to set a user's password to expire in 90 days, use the following command:

```
chage -M 90 <username>
```

In the above command, replace **<username>** with the name of the user. To disable password expiration, it is traditional to use a value of **99999** after the **-M** option (this equates to a little over 273 years).

For more information on the options available with the **chage** command, refer to the table below.

Table 2.1. chage command line options

Option	Description
-d days	Specifies the number of days since January 1, 1970 the password was changed.
-E date	Specifies the date on which the account is locked, in the format YYYY-MM-DD. Instead of the date, the number of days since January 1, 1970 can also be used.
-I days	Specifies the number of inactive days after the password expiration before locking the account. If the value is 0 , the account is not locked after the password expires.
-l	Lists current account aging settings.
-m days	Specify the minimum number of days after which the user must change passwords. If the value is 0 , the password does not expire.
-M days	Specify the maximum number of days for which the password is valid. When the number of days specified by this option plus the number of days specified with the -d option is less than the current day, the user must change passwords before using the account.
-W days	Specifies the number of days before the password expiration date to warn the user.

You can also use the **chage** command in interactive mode to modify multiple password aging and account details. Use the following command to enter interactive mode:

```
chage <username>
```

The following is a sample interactive session using this command:

```
~]# chage juan
Changing the aging information for juan
Enter the new value, or press ENTER for the default
Minimum Password Age [0]: 10
Maximum Password Age [99999]: 90
Last Password Change (YYYY-MM-DD) [2006-08-18]:
Password Expiration Warning [7]:
Password Inactive [-1]:
Account Expiration Date (YYYY-MM-DD) [1969-12-31]:
```

You can configure a password to expire the first time a user logs in. This forces users to change passwords immediately.

1. Set up an initial password. There are two common approaches to this step: you can either assign a default password, or you can use a null password.

To assign a default password, type the following at a shell prompt as **root**:

```
passwd username
```

To assign a null password instead, use the following command:

```
passwd -d username
```



Avoid using null passwords whenever possible

Using a null password, while convenient, is a highly insecure practice, as any third party can log in first and access the system using the insecure username. Always make sure that the user is ready to log in before unlocking an account with a null password.

2. Force immediate password expiration by running the following command as **root**:

```
chage -d 0 username
```

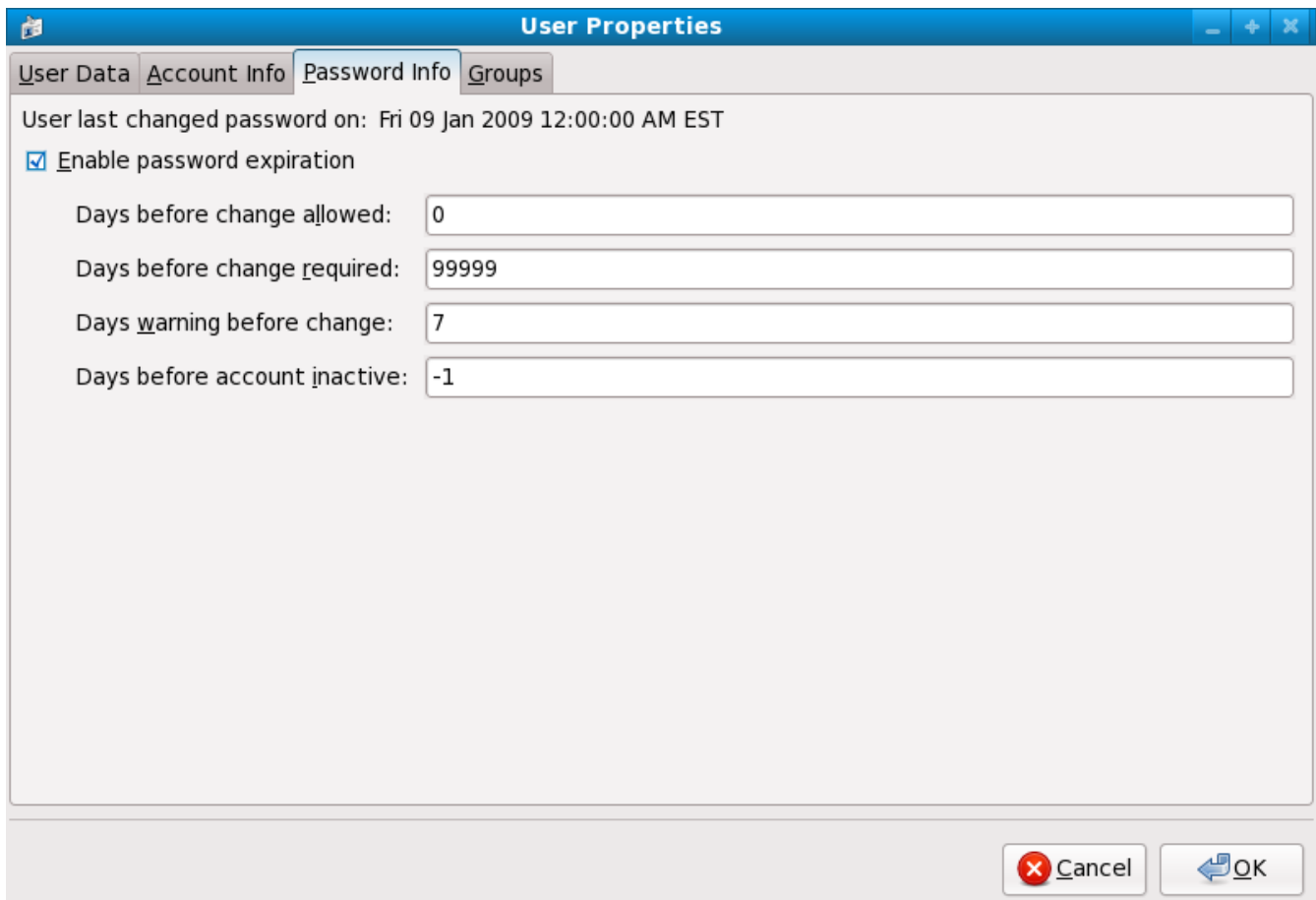
This command sets the value for the date the password was last changed to the epoch (January 1, 1970). This value forces immediate password expiration no matter what password aging policy, if any, is in place.

Upon the initial log in, the user is now prompted for a new password.

You can also use the graphical **User Manager** application to create password aging policies, as follows. Note: you need Administrator privileges to perform this procedure.

1. Click the **System** menu on the Panel, point to **Administration** and then click **Users and Groups** to display the User Manager. Alternatively, type the command **system-config-users** at a shell prompt.
2. Click the **Users** tab, and select the required user in the list of users.
3. Click **Properties** on the toolbar to display the User Properties dialog box (or choose **Properties** on the **File** menu).

4. Click the **Password Info** tab, and select the check box for **Enable password expiration**.
5. Enter the required value in the **Days before change required** field, and click **OK**.



User Properties

User Data Account Info **Password Info** Groups

User last changed password on: Fri 09 Jan 2009 12:00:00 AM EST

☒ Enable password expiration

Days before change allowed: 0

Days before change required: 99999

Days warning before change: 7

Days before account inactive: -1

Cancel OK

Figure 2.1. Specifying password aging options

2.1.4. Administrative Controls

When administering a home machine, the user must perform some tasks as the root user or by acquiring effective root privileges via a *setuid* program, such as **sudo** or **su**. A *setuid* program is one that operates with the user ID (*UID*) of the program's owner rather than the user operating the program. Such programs are denoted by an **s** in the owner section of a long format listing, as in the following example:

```
~]$ ls -l /bin/su
-rwsr-xr-x. 1 root root 34904 Mar 10 2011 /bin/su
```



Note

The **s** may be upper case or lower case. If it appears as upper case, it means that the underlying permission bit has not been set.

For the system administrators of an organization, however, choices must be made as to how much administrative access users within the organization should have to their machine. Through a PAM module called **pam_console.so**, some activities normally reserved only for the root user, such as rebooting and mounting removable media are allowed for the first user that logs in at the physical

console (refer to *Managing Single Sign-On and Smart Cards* for more information about the **pam_console.so** module.) However, other important system administration tasks, such as altering network settings, configuring a new mouse, or mounting network devices, are not possible without administrative privileges. As a result, system administrators must decide how much access the users on their network should receive.

2.1.4.1. Allowing Root Access

If the users within an organization are trusted and computer-literate, then allowing them root access may not be an issue. Allowing root access by users means that minor activities, like adding devices or configuring network interfaces, can be handled by the individual users, leaving system administrators free to deal with network security and other important issues.

On the other hand, giving root access to individual users can lead to the following issues:

- ▶ *Machine Misconfiguration* — Users with root access can misconfigure their machines and require assistance to resolve issues. Even worse, they might open up security holes without knowing it.
- ▶ *Running Insecure Services* — Users with root access might run insecure servers on their machine, such as FTP or Telnet, potentially putting usernames and passwords at risk. These services transmit this information over the network in plain text.
- ▶ *Running Email Attachments As Root* — Although rare, email viruses that affect Linux do exist. The only time they are a threat, however, is when they are run by the root user.
- ▶ *Keeping the audit trail intact* — Because the root account is often shared by multiple users, so that multiple system administrators can maintain the system, it is impossible to figure out which of those users was root at a given time. When using separate logins, the account a user logs in with, as well as a unique number for session tracking purposes, is put into the task structure, which is inherited by every process that the user starts. When using concurrent logins, the unique number can be used to trace actions to specific logins. When an action generates an audit event, it is recorded with the login account and the session associated with that unique number. Use the **aulast** command to view these logins and sessions. The **--proof** option of the **aulast** command can be used suggest a specific **ausearch** query to isolate auditable events generated by a particular session.

2.1.4.2. Disallowing Root Access

If an administrator is uncomfortable allowing users to log in as root for these or other reasons, the root password should be kept secret, and access to runlevel one or single user mode should be disallowed through boot loader password protection (refer to [Section 2.1.2.2, “Boot Loader Passwords”](#) for more information on this topic.)

The following are four different ways that an administrator can further ensure that root logins are disallowed:

Changing the root shell

To prevent users from logging in directly as root, the system administrator can set the root account's shell to **/sbin/nologin** in the **/etc/passwd** file.

Table 2.2. Disabling the Root Shell

Effects	Does Not Affect
Prevents access to the root shell and logs any such attempts. The following programs are prevented from accessing the root account:	Programs that do not require a shell, such as FTP clients, mail clients, and many setuid programs. The following programs are <i>not</i> prevented from accessing the root account:
» login » gdm » kdm » xdm » su » ssh » scp » sftp	» sudo » FTP clients » Email clients

Disabling root access via any console device (tty)

To further limit access to the root account, administrators can disable root logins at the console by editing the `/etc/securetty` file. This file lists all devices the root user is allowed to log into. If the file does not exist at all, the root user can log in through any communication device on the system, whether via the console or a raw network interface. This is dangerous, because a user can log in to their machine as root via Telnet, which transmits the password in plain text over the network.

By default, Red Hat Enterprise Linux's `/etc/securetty` file only allows the root user to log in at the console physically attached to the machine. To prevent the root user from logging in, remove the contents of this file by typing the following command at a shell prompt as root:

```
echo > /etc/securetty
```

To enable **securetty** support in the KDM, GDM, and XDM login managers, add the following line:

```
auth [user_unknown=ignore success=ok ignore=ignore default=bad]
pam_securetty.so
```

to the files listed below:

- » `/etc/pam.d/gdm`
- » `/etc/pam.d/gdm-autologin`
- » `/etc/pam.d/gdm-fingerprint`
- » `/etc/pam.d/gdm-password`
- » `/etc/pam.d/gdm-smartcard`
- » `/etc/pam.d/kdm`
- » `/etc/pam.d/kdm-np`
- » `/etc/pam.d/xdm`



Warning

A blank `/etc/securetty` file does *not* prevent the root user from logging in remotely using the OpenSSH suite of tools because the console is not opened until after authentication.

Table 2.3. Disabling Root Logins

Effects	Does Not Affect
Prevents access to the root account via the console or the network. The following programs are prevented from accessing the root account: ▶ login ▶ gdm ▶ kdm ▶ xdm ▶ Other network services that open a tty	Programs that do not log in as root, but perform administrative tasks through <code>setuid</code> or other mechanisms. The following programs are <i>not</i> prevented from accessing the root account: ▶ su ▶ sudo ▶ ssh ▶ scp ▶ sftp

Disabling root SSH logins

To prevent root logins via the SSH protocol, edit the SSH daemon's configuration file, `/etc/ssh/sshd_config`, and change the line that reads:

```
#PermitRootLogin yes
```

to read as follows:

```
PermitRootLogin no
```

Table 2.4. Disabling Root SSH Logins

Effects	Does Not Affect
Prevents root access via the OpenSSH suite of tools. The following programs are prevented from accessing the root account: ▶ ssh ▶ scp ▶ sftp	Programs that are not part of the OpenSSH suite of tools.

Using PAM to limit root access to services

PAM, through the `/lib/security/pam_listfile.so` module, allows great flexibility in denying specific accounts. The administrator can use this module to reference a list of users who are not allowed to log in. To limit root access to a system service, edit the file for the target service in the `/etc/pam.d/` directory and make sure the `pam_listfile.so` module is

required for authentication.

The following is an example of how the module is used for the **vsftpd** FTP server in the **/etc/pam.d/vsftpd** PAM configuration file (the `\` character at the end of the first line is *not* necessary if the directive is on a single line):

```
auth    required    /lib/security/pam_listfile.so    item=user \
sense=deny file=/etc/vsftpd.ftpusers onerr=succeed
```

This instructs PAM to consult the **/etc/vsftpd.ftpusers** file and deny access to the service for any listed user. The administrator can change the name of this file, and can keep separate lists for each service or use one central list to deny access to multiple services.

If the administrator wants to deny access to multiple services, a similar line can be added to the PAM configuration files, such as **/etc/pam.d/pop** and **/etc/pam.d/imap** for mail clients, or **/etc/pam.d/ssh** for SSH clients.

For more information about PAM, refer to the chapter titled *Using Pluggable Authentication Modules (PAM)* in the *Managing Single Sign-On and Smart Cards* guide.

Table 2.5. Disabling Root Using PAM

Effects	Does Not Affect
Prevents root access to network services that are PAM aware. The following services are prevented from accessing the root account: ▶ login ▶ gdm ▶ kdm ▶ xdm ▶ ssh ▶ scp ▶ sftp ▶ FTP clients ▶ Email clients ▶ Any PAM aware services	Programs and services that are not PAM aware.

2.1.4.3. Enabling Automatic Logouts

When the user is logged in as **root**, an unattended login session may pose a significant security risk. To reduce this risk, you can configure the system to automatically log out idle users after a fixed period of time:

1. Make sure the *screen* package is installed. You can do so by running the following command as **root**:

```
yum install screen
```

For more information on how to install packages in Red Hat Enterprise Linux, refer to *Red Hat Enterprise Linux 6 Deployment Guide*.

2. As **root**, add the following line at the beginning of the **/etc/profile** file to make sure the processing of this file cannot be interrupted:

```
trap "" 1 2 3 15
```

3. Add the following lines at the end of the **/etc/profile** file to start a **screen** session each time a user logs in to a virtual console or remotely:

```
SCREENEXEC="screen"
if [ -w $(tty) ]; then
    trap "exec $SCREENEXEC" 1 2 3 15
    echo -n 'Starting session in 10 seconds'
    sleep 10
    exec $SCREENEXEC
fi
```

Note that each time a new session starts, a message will be displayed and the user will have to wait ten seconds. To adjust the time to wait before starting a session, change the value after the **sleep** command.

4. Add the following lines to the **/etc/screenrc** configuration file to close the **screen** session after a given period of inactivity:

```
idle 120 quit autodetach off
```

This will set the time limit to 120 seconds. To adjust this limit, change the value after the **idle** directive.

Alternatively, you can configure the system to only lock the session by using the following lines instead:

```
idle 120 lockscreen autodetach off
```

This way, a password will be required to unlock the session.

The changes take effect the next time a user logs in to the system.

2.1.4.4. Limiting Root Access

Rather than completely denying access to the root user, the administrator may want to allow access only via **setuid** programs, such as **su** or **sudo**. For more information on **su** and **sudo**, refer to the *Red Hat Enterprise Linux 6 Deployment Guide* and the **su(1)** and **sudo(8)** man pages.

2.1.5. Session Locking

Users may need to leave their workstation unattended for a number of reasons during everyday operation. This could present an opportunity for an attacker to physically access the machine, especially in environments with insufficient physical security measures (see [Section 1.1.3.1, “Physical Controls”](#)). Laptops are especially exposed since their mobility interferes with physical security. You can alleviate these risks by using session locking features which prevent access to the system until a correct password is entered.

**Note**

The main advantage of locking the screen instead of logging out is that a lock allows the user's processes (such as file transfers) to continue running. Logging out would stop these processes.

2.1.5.1. Locking GNOME Using `gnome-screensaver-command`

The default desktop environment for Red Hat Enterprise Linux 6, GNOME, includes a feature which allows users to lock their screen at any time. There are several ways to activate the lock:

- Press the key combination specified in **System** → **Preferences** → **Keyboard Shortcuts** → **Desktop** → **Lock screen**. The default combination is **Ctrl+Alt+L**.
- Select **System** → **Lock screen** on the panel.
- Execute the following command from a command line interface:

```
gnome-screensaver-command -l
```

All of the techniques described have the same result: the screen saver is activated and the screen is locked. Users can then press any key to deactivate the screen saver, enter their password and continue working.

Keep in mind that this function requires the **gnome-screensaver** process to be running. You can check whether this is the case by using any command which provides information about processes. For example, execute the following command from the terminal:

```
pidof gnome-screensaver
```

If the **gnome-screensaver** process is currently running, a number denoting its identification number (PID) will be displayed on the screen after executing the command. If the process is not currently running, the command will provide no output at all.

Refer to the **gnome-screensaver-command(1)** man page for additional information.

**Important**

The means of locking the screen described above rely on manual activation. Administrators should therefore advise their users to lock their computers every time they leave them unattended, even if only for a short period of time.

2.1.5.1.1. Automatic Lock on Screen Saver Activation

As the name **gnome-screensaver-command** suggests, the locking functionality is tied to GNOME's screen saver. It is possible to tie the lock to the screen saver's activation, locking the workstation every time it is left unattended for a set period of time. This function is activated by default with a five minute timeout.

To change the automatic locking settings, select **System** → **Preferences** → **Screensaver** on the main panel. This opens a window which allows setting the timeout period (the **Regard the computer as idle after** slider) and activating or deactivating the automatic lock (the **Lock screen when screensaver is active** check box).

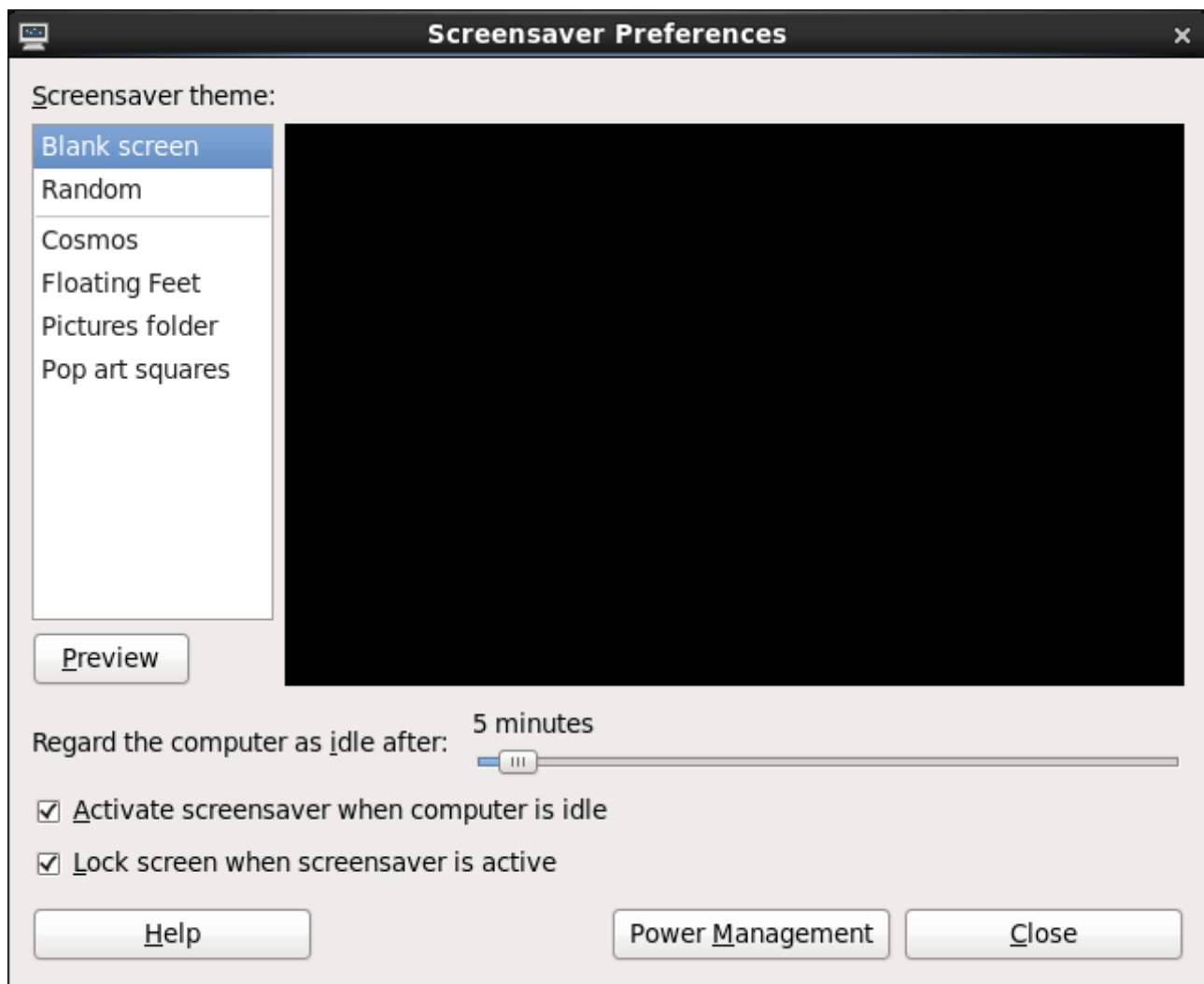


Figure 2.2. Changing the screen saver preferences



Note

Disabling the **Activate screensaver when computer is idle** option in the **Screensaver Preferences** dialog prevents the screen saver from starting automatically. Automatic locking is therefore disabled as well, but it is still possible to lock the workstation manually using the techniques described in [Section 2.1.5.1, “Locking GNOME Using gnome-screensaver-command”](#).

2.1.5.1.2. Remote Session Locking

You can also lock a GNOME session remotely using **ssh** as long as the target workstation accepts connections over this protocol. To remotely lock the screen on a machine you have access to, execute the following command:

```
ssh -X <username>@<server> "export DISPLAY=:0; gnome-screensaver-command -l"
```

Replace **<username>** with your user name and **<server>** with the IP address of the workstation you wish to lock.

Refer to [Section 3.2.2, “Secure Shell”](#) for more information regarding **ssh**.

2.1.5.2. Locking Virtual Consoles Using **vlock**

Users may also need to lock a virtual console. This can be done using a utility called **vlock**. To install this utility, execute the following command as root:

```
~]# yum install vlock
```

After installation, any console session can be locked using the **vlock** command without any additional parameters. This locks the currently active virtual console session while still allowing access to the others. To prevent access to all virtual consoles on the workstation, execute the following:

```
vlock -a
```

In this case, **vlock** locks the currently active console and the **-a** option prevents switching to other virtual consoles.

Refer to the **vlock(1)** man page for additional information.



Important

There are several known issues relevant to the version of **vlock** currently available for Red Hat Enterprise Linux 6:

- ▶ The program does not currently allow unlocking consoles using the root password. Additional information can be found in [BZ#895066](#).
- ▶ Locking a console does not clear the screen and scrollbar buffer, allowing anyone with physical access to the workstation to view previously issued commands and any output displayed in the console. Refer to [BZ#807369](#) for more information.

2.1.6. Available Network Services

While user access to administrative controls is an important issue for system administrators within an organization, monitoring which network services are active is of paramount importance to anyone who administers and operates a Linux system.

Many services under Red Hat Enterprise Linux 6 behave as network servers. If a network service is running on a machine, then a server application (called a *daemon*), is listening for connections on one or more network ports. Each of these servers should be treated as a potential avenue of attack.

2.1.6.1. Risks To Services

Network services can pose many risks for Linux systems. Below is a list of some of the primary issues:

- ▶ *Denial of Service Attacks (DoS)* — By flooding a service with requests, a denial of service attack can render a system unusable as it tries to log and answer each request.
- ▶ *Distributed Denial of Service Attack (DDoS)* — A type of DoS attack which uses multiple compromised machines (often numbering in the thousands or more) to direct a coordinated attack on a service, flooding it with requests and making it unusable.
- ▶ *Script Vulnerability Attacks* — If a server is using scripts to execute server-side actions, as Web servers commonly do, a cracker can attack improperly written scripts. These script vulnerability attacks can lead to a buffer overflow condition or allow the attacker to alter files on the system.

- **Buffer Overflow Attacks** — Services that connect to ports numbered 0 through 1023 must run as an administrative user. If the application has an exploitable buffer overflow, an attacker could gain access to the system as the user running the daemon. Because exploitable buffer overflows exist, crackers use automated tools to identify systems with vulnerabilities, and once they have gained access, they use automated rootkits to maintain their access to the system.



Note

The threat of buffer overflow vulnerabilities is mitigated in Red Hat Enterprise Linux by *ExecShield*, an executable memory segmentation and protection technology supported by x86-compatible uni- and multi-processor kernels. ExecShield reduces the risk of buffer overflow by separating virtual memory into executable and non-executable segments. Any program code that tries to execute outside of the executable segment (such as malicious code injected from a buffer overflow exploit) triggers a segmentation fault and terminates.

Execshield also includes support for *No eXecute* (NX) technology on AMD64 platforms and *eXecute Disable* (XD) technology on Itanium and Intel® 64 systems. These technologies work in conjunction with ExecShield to prevent malicious code from running in the executable portion of virtual memory with a granularity of 4KB of executable code, lowering the risk of attack from stealthy buffer overflow exploits.



Important

To limit exposure to attacks over the network, all services that are unused should be turned off.

2.1.6.2. Identifying and Configuring Services

To enhance security, most network services installed with Red Hat Enterprise Linux are turned off by default. There are, however, some notable exceptions:

- **cupsd** — The default print server for Red Hat Enterprise Linux.
- **lpd** — An alternative print server.
- **xinetd** — A super server that controls connections to a range of subordinate servers, such as **gssftp** and **telnet**.
- **sendmail** — The Sendmail *Mail Transport Agent* (MTA) is enabled by default, but only listens for connections from the localhost.
- **sshd** — The OpenSSH server, which is a secure replacement for Telnet.

When determining whether to leave these services running, it is best to use common sense and err on the side of caution. For example, if a printer is not available, do not leave **cupsd** running. The same is true for **portmap**. If you do not mount NFSv3 volumes or use NIS (the **ypbind** service), then **portmap** should be disabled.

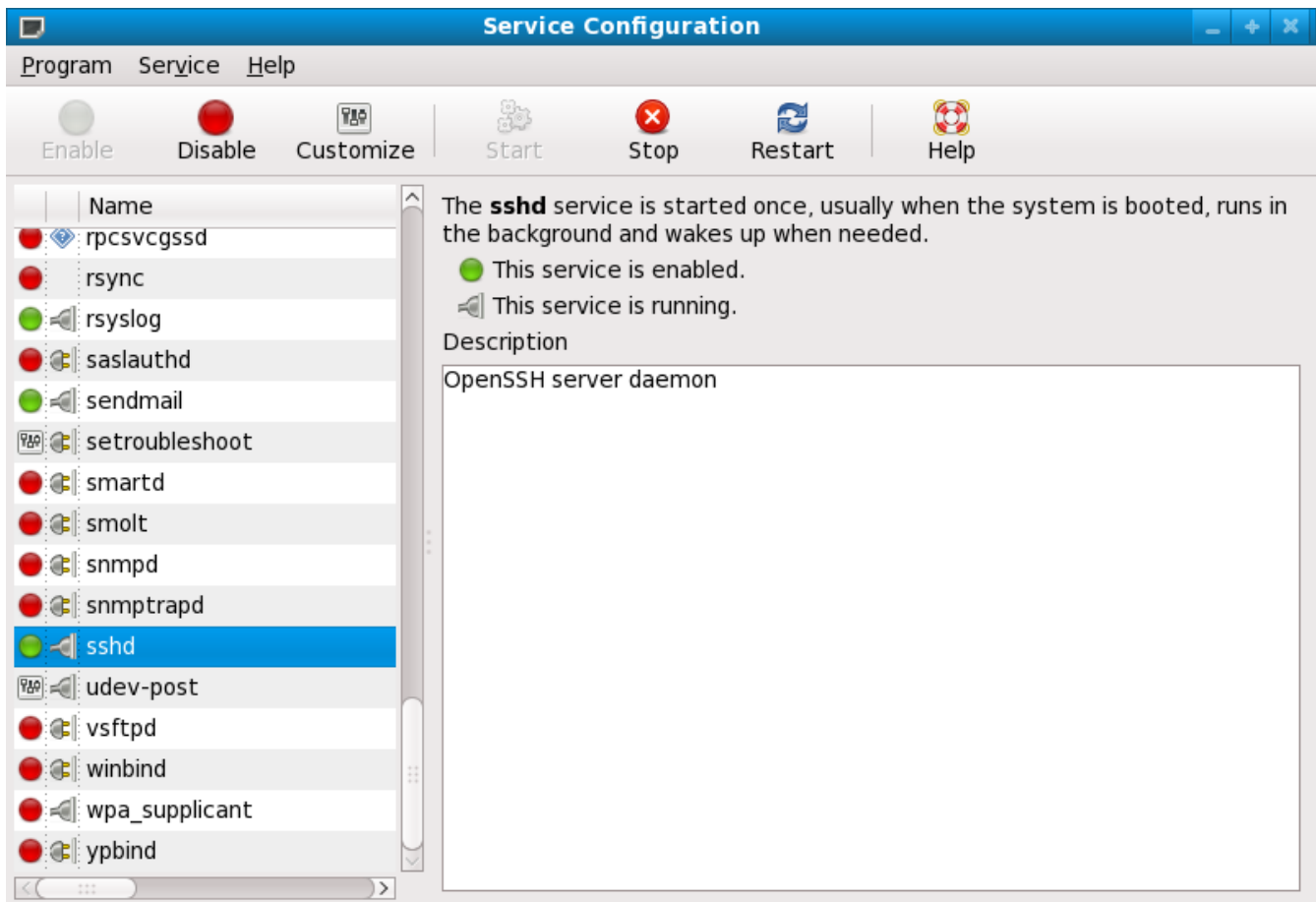


Figure 2.3. Services Configuration Tool

If unsure of the purpose for a particular service, the **Services Configuration Tool** has a description field, illustrated in [Figure 2.3, “Services Configuration Tool”](#), that provides additional information.

Checking which network services are available to start at boot time is only part of the story. You should also check which ports are open and listening. Refer to [Section 2.2.9, “Verifying Which Ports Are Listening”](#) for more information.

2.1.6.3. Insecure Services

Potentially, any network service is insecure. This is why turning off unused services is so important. Exploits for services are routinely revealed and patched, making it very important to regularly update packages associated with any network service. Refer to [Section 1.5, “Security Updates”](#) for more information.

Some network protocols are inherently more insecure than others. These include any services that:

- *Transmit Usernames and Passwords Over a Network Unencrypted* — Many older protocols, such as Telnet and FTP, do not encrypt the authentication session and should be avoided whenever possible.
- *Transmit Sensitive Data Over a Network Unencrypted* — Many protocols transmit data over the network unencrypted. These protocols include Telnet, FTP, HTTP, and SMTP. Many network file systems, such as NFS and SMB, also transmit information over the network unencrypted. It is the user's responsibility when using these protocols to limit what type of data is transmitted.

Remote memory dump services, like **netdump**, transmit the contents of memory over the network unencrypted. Memory dumps can contain passwords or, even worse, database entries and other

sensitive information.

Other services like **finger** and **rwhod** reveal information about users of the system.

Examples of inherently insecure services include **rlogin**, **rsh**, **telnet**, and **vsftpd**.

All remote login and shell programs (**rlogin**, **rsh**, and **telnet**) should be avoided in favor of SSH. Refer to [Section 2.1.8, “Security Enhanced Communication Tools”](#) for more information about **sshd**.

FTP is not as inherently dangerous to the security of the system as remote shells, but FTP servers must be carefully configured and monitored to avoid problems. Refer to [Section 2.2.6, “Securing FTP”](#) for more information about securing FTP servers.

Services that should be carefully implemented and behind a firewall include:

- ▶ **finger**
- ▶ **authd** (this was called **identd** in previous Red Hat Enterprise Linux releases.)
- ▶ **netdump**
- ▶ **netdump-server**
- ▶ **nfs**
- ▶ **rwhod**
- ▶ **sendmail**
- ▶ **smb** (Samba)
- ▶ **yppasswdd**
- ▶ **ypserv**
- ▶ **ypxfrd**

More information on securing network services is available in [Section 2.2, “Server Security”](#).

The next section discusses tools available to set up a simple firewall.

2.1.7. Personal Firewalls

After the *necessary* network services are configured, it is important to implement a firewall.



Important

You should configure the necessary services and implement a firewall *before* connecting to the Internet or any other network that you do not trust.

Firewalls prevent network packets from accessing the system's network interface. If a request is made to a port that is blocked by a firewall, the request is ignored. If a service is listening on one of these blocked ports, it does not receive the packets and is effectively disabled. For this reason, care should be taken when configuring a firewall to block access to ports not in use, while not blocking access to ports used by configured services.

For most users, the best tool for configuring a simple firewall is the graphical firewall configuration tool which ships with Red Hat Enterprise Linux: the **Firewall Configuration Tool (system-config-firewall)**. This tool creates broad **iptables** rules for a general-purpose firewall using a control panel interface.

Refer to [Section 2.8.2, “Basic Firewall Configuration”](#) for more information about using this application

and its available options.

For advanced users and server administrators, manually configuring a firewall with **iptables** is probably a better option. Refer to [Section 2.8, “Firewalls”](#) for more information. Refer to [Section 2.8.9, “IPTables”](#) for a comprehensive guide to the **iptables** command.

2.1.8. Security Enhanced Communication Tools

As the size and popularity of the Internet has grown, so has the threat of communication interception. Over the years, tools have been developed to encrypt communications as they are transferred over the network.

Red Hat Enterprise Linux 6 ships with two basic tools that use high-level, public-key-cryptography-based encryption algorithms to protect information as it travels over the network.

- *OpenSSH* — A free implementation of the SSH protocol for encrypting network communication.
- *Gnu Privacy Guard (GPG)* — A free implementation of the PGP (Pretty Good Privacy) encryption application for encrypting data.

OpenSSH is a safer way to access a remote machine and replaces older, unencrypted services like **telnet** and **rsh**. OpenSSH includes a network service called **sshd** and three command line client applications:

- **ssh** — A secure remote console access client.
- **scp** — A secure remote copy command.
- **sftp** — A secure pseudo-ftp client that allows interactive file transfer sessions.

Refer to [Section 3.2.2, “Secure Shell”](#) for more information regarding OpenSSH.



Important

Although the **sshd** service is inherently secure, the service *must* be kept up-to-date to prevent security threats. Refer to [Section 1.5, “Security Updates”](#) for more information.

GPG is one way to ensure private email communication. It can be used both to email sensitive data over public networks and to protect sensitive data on hard drives.

2.2. Server Security

When a system is used as a server on a public network, it becomes a target for attacks. Hardening the system and locking down services is therefore of paramount importance for the system administrator.

Before delving into specific issues, review the following general tips for enhancing server security:

- Keep all services current, to protect against the latest threats.
- Use secure protocols whenever possible.
- Serve only one type of network service per machine whenever possible.
- Monitor all servers carefully for suspicious activity.

2.2.1. Securing Services With TCP Wrappers and xinetd

TCP Wrappers provide access control to a variety of services. Most modern network services, such as SSH, Telnet, and FTP, make use of TCP Wrappers, which stand guard between an incoming request

and the requested service.

The benefits offered by TCP Wrappers are enhanced when used in conjunction with **xinetd**, a super server that provides additional access, logging, binding, redirection, and resource utilization control.



Note

It is a good idea to use iptables firewall rules in conjunction with TCP Wrappers and **xinetd** to create redundancy within service access controls. Refer to [Section 2.8, “Firewalls”](#) for more information about implementing firewalls with iptables commands.

The following subsections assume a basic knowledge of each topic and focus on specific security options.

2.2.1.1. Enhancing Security With TCP Wrappers

TCP Wrappers are capable of much more than denying access to services. This section illustrates how they can be used to send connection banners, warn of attacks from particular hosts, and enhance logging functionality. Refer to the **hosts_options** man page for information about the TCP Wrapper functionality and control language. Refer to the **xinetd.conf** man page available online at <http://linux.die.net/man/5/xinetd.conf> for available flags, which act as options you can apply to a service.

2.2.1.1.1. TCP Wrappers and Connection Banners

Displaying a suitable banner when users connect to a service is a good way to let potential attackers know that the system administrator is being vigilant. You can also control what information about the system is presented to users. To implement a TCP Wrappers banner for a service, use the **banner** option.

This example implements a banner for **vsftpd**. To begin, create a banner file. It can be anywhere on the system, but it must have same name as the daemon. For this example, the file is called **/etc/banners/vsftpd** and contains the following lines:

```
220-Hello, %c
220-All activity on ftp.example.com is logged.
220-Inappropriate use will result in your access privileges being removed.
```

The **%c** token supplies a variety of client information, such as the username and hostname, or the username and IP address to make the connection even more intimidating.

For this banner to be displayed to incoming connections, add the following line to the **/etc/hosts.allow** file:

```
vsftpd : ALL : banners /etc/banners/
```

2.2.1.1.2. TCP Wrappers and Attack Warnings

If a particular host or network has been detected attacking the server, TCP Wrappers can be used to warn the administrator of subsequent attacks from that host or network using the **spawn** directive.

In this example, assume that a cracker from the 206.182.68.0/24 network has been detected attempting to attack the server. Place the following line in the **/etc/hosts.deny** file to deny any connection attempts from that network, and to log the attempts to a special file:

```
ALL : 206.182.68.0 : spawn /bin/echo `date` %c %d >> /var/log/intruder_alert
```

The **%d** token supplies the name of the service that the attacker was trying to access.

To allow the connection and log it, place the **spawn** directive in the **/etc/hosts.allow** file.



Note

Because the **spawn** directive executes any shell command, it is a good idea to create a special script to notify the administrator or execute a chain of commands in the event that a particular client attempts to connect to the server.

2.2.1.1.3. TCP Wrappers and Enhanced Logging

If certain types of connections are of more concern than others, the log level can be elevated for that service using the **severity** option.

For this example, assume that anyone attempting to connect to port 23 (the Telnet port) on an FTP server is a cracker. To denote this, place an **emerg** flag in the log files instead of the default flag, **info**, and deny the connection.

To do this, place the following line in **/etc/hosts.deny**:

```
in.telnetd : ALL : severity emerg
```

This uses the default **authpriv** logging facility, but elevates the priority from the default value of **info** to **emerg**, which posts log messages directly to the console.

2.2.1.2. Enhancing Security With xinetd

This section focuses on using **xinetd** to set a trap service and using it to control resource levels available to any given **xinetd** service. Setting resource limits for services can help thwart *Denial of Service* (DoS) attacks. Refer to the man pages for **xinetd** and **xinetd.conf** for a list of available options.

2.2.1.2.1. Setting a Trap

One important feature of **xinetd** is its ability to add hosts to a global **no_access** list. Hosts on this list are denied subsequent connections to services managed by **xinetd** for a specified period or until **xinetd** is restarted. You can do this using the **SENSOR** attribute. This is an easy way to block hosts attempting to scan the ports on the server.

The first step in setting up a **SENSOR** is to choose a service you do not plan on using. For this example, Telnet is used.

Edit the file **/etc/xinetd.d/telnet** and change the **flags** line to read:

```
flags = SENSOR
```

Add the following line:

```
deny_time = 30
```

This denies any further connection attempts to that port by that host for 30 minutes. Other acceptable

values for the **deny_time** attribute are **FOREVER**, which keeps the ban in effect until **xinetd** is restarted, and **NEVER**, which allows the connection and logs it.

Finally, the last line should read:

```
disable          = no
```

This enables the trap itself.

While using **SENSOR** is a good way to detect and stop connections from undesirable hosts, it has two drawbacks:

- It does not work against stealth scans.
- An attacker who knows that a **SENSOR** is running can mount a Denial of Service attack against particular hosts by forging their IP addresses and connecting to the forbidden port.

2.2.1.2.2. Controlling Server Resources

Another important feature of **xinetd** is its ability to set resource limits for services under its control.

It does this using the following directives:

- **cps = <number_of_connections> <wait_period>** — Limits the rate of incoming connections. This directive takes two arguments:
 - **<number_of_connections>** — The number of connections per second to handle. If the rate of incoming connections is higher than this, the service is temporarily disabled. The default value is fifty (50).
 - **<wait_period>** — The number of seconds to wait before re-enabling the service after it has been disabled. The default interval is ten (10) seconds.
- **instances = <number_of_connections>** — Specifies the total number of connections allowed to a service. This directive accepts either an integer value or **UNLIMITED**.
- **per_source = <number_of_connections>** — Specifies the number of connections allowed to a service by each host. This directive accepts either an integer value or **UNLIMITED**.
- **rlimit_as = <number[K|M]>** — Specifies the amount of memory address space the service can occupy in kilobytes or megabytes. This directive accepts either an integer value or **UNLIMITED**.
- **rlimit_cpu = <number_of_seconds>** — Specifies the amount of time in seconds that a service may occupy the CPU. This directive accepts either an integer value or **UNLIMITED**.

Using these directives can help prevent any single **xinetd** service from overwhelming the system, resulting in a denial of service.

2.2.2. Securing Portmap

The **portmap** service is a dynamic port assignment daemon for RPC services such as NIS and NFS. It has weak authentication mechanisms and has the ability to assign a wide range of ports for the services it controls. For these reasons, it is difficult to secure.



Note

Securing **portmap** only affects NFSv2 and NFSv3 implementations, since NFSv4 no longer requires it. If you plan to implement an NFSv2 or NFSv3 server, then **portmap** is required, and the following section applies.

If running RPC services, follow these basic rules.

2.2.2.1. Protect portmap With TCP Wrappers

It is important to use TCP Wrappers to limit which networks or hosts have access to the **portmap** service since it has no built-in form of authentication.

Further, use *only* IP addresses when limiting access to the service. Avoid using hostnames, as they can be forged by DNS poisoning and other methods.

2.2.2.2. Protect portmap With iptables

To further restrict access to the **portmap** service, it is a good idea to add iptables rules to the server and restrict access to specific networks.

Below are two example iptables commands. The first allows TCP connections to the port 111 (used by the **portmap** service) from the 192.168.0.0/24 network. The second allows TCP connections to the same port from the localhost. This is necessary for the **sgi_fam** service used by **Nautilus**. All other packets are dropped.

```
~]# iptables -A INPUT -p tcp -s ! 192.168.0.0/24 --dport 111 -j DROP
~]# iptables -A INPUT -p tcp -s 127.0.0.1 --dport 111 -j ACCEPT
```

To similarly limit UDP traffic, use the following command:

```
~]# iptables -A INPUT -p udp -s ! 192.168.0.0/24 --dport 111 -j DROP
```



Note

Refer to [Section 2.8, “Firewalls”](#) for more information about implementing firewalls with iptables commands.

2.2.3. Securing NIS

The *Network Information Service* (NIS) is an RPC service, called **ypserv**, which is used in conjunction with **portmap** and other related services to distribute maps of usernames, passwords, and other sensitive information to any computer claiming to be within its domain.

A NIS server is comprised of several applications. They include the following:

- ▶ **/usr/sbin/rpc.yppasswdd** — Also called the **yppasswdd** service, this daemon allows users to change their NIS passwords.
- ▶ **/usr/sbin/rpc.ypxfrd** — Also called the **ypxfrd** service, this daemon is responsible for NIS map transfers over the network.
- ▶ **/usr/sbin/yppush** — This application propagates changed NIS databases to multiple NIS servers.
- ▶ **/usr/sbin/ypserv** — This is the NIS server daemon.

NIS is somewhat insecure by today's standards. It has no host authentication mechanisms and transmits all of its information over the network unencrypted, including password hashes. As a result, extreme care must be taken when setting up a network that uses NIS. This is further complicated by the fact that the default configuration of NIS is inherently insecure.

It is recommended that anyone planning to implement a NIS server first secure the **portmap** service as

outlined in [Section 2.2.2, “Securing Portmap”](#), then address the following issues, such as network planning.

2.2.3.1. Carefully Plan the Network

Because NIS transmits sensitive information unencrypted over the network, it is important the service be run behind a firewall and on a segmented and secure network. Whenever NIS information is transmitted over an insecure network, it risks being intercepted. Careful network design can help prevent severe security breaches.

2.2.3.2. Use a Password-like NIS Domain Name and Hostname

Any machine within a NIS domain can use commands to extract information from the server without authentication, as long as the user knows the NIS server's DNS hostname and NIS domain name.

For instance, if someone either connects a laptop computer into the network or breaks into the network from outside (and manages to spoof an internal IP address), the following command reveals the `/etc/passwd` map:

```
ypcat -d <NIS_domain> -h <DNS_hostname> passwd
```

If this attacker is a root user, they can obtain the `/etc/shadow` file by typing the following command:

```
ypcat -d <NIS_domain> -h <DNS_hostname> shadow
```



Note

If Kerberos is used, the `/etc/shadow` file is not stored within a NIS map.

To make access to NIS maps harder for an attacker, create a random string for the DNS hostname, such as `o7hfawtgmhgw.domain.com`. Similarly, create a *different* randomized NIS domain name. This makes it much more difficult for an attacker to access the NIS server.

2.2.3.3. Edit the `/var/yp/securenets` File

If the `/var/yp/securenets` file is blank or does not exist (as is the case after a default installation), NIS listens to all networks. One of the first things to do is to put netmask/network pairs in the file so that `ypserv` only responds to requests from the appropriate network.

Below is a sample entry from a `/var/yp/securenets` file:

```
255.255.255.0    192.168.0.0
```



Warning

Never start a NIS server for the first time without creating the `/var/yp/securenets` file.

This technique does not provide protection from an IP spoofing attack, but it does at least place limits on what networks the NIS server services.

2.2.3.4. Assign Static Ports and Use iptables Rules

All of the servers related to NIS can be assigned specific ports except for `rpc.yppasswdd` — the

daemon that allows users to change their login passwords. Assigning ports to the other two NIS server daemons, **rpc.ypxfrd** and **ypserv**, allows for the creation of firewall rules to further protect the NIS server daemons from intruders.

To do this, add the following lines to **/etc/sysconfig/network**:

```
YPSERV_ARGS="-p 834"
YPXFRD_ARGS="-p 835"
```

The following iptables rules can then be used to enforce which network the server listens to for these ports:

```
~]# iptables -A INPUT -p ALL -s ! 192.168.0.0/24 --dport 834 -j DROP
~]# iptables -A INPUT -p ALL -s ! 192.168.0.0/24 --dport 835 -j DROP
```

This means that the server only allows connections to ports 834 and 835 if the requests come from the 192.168.0.0/24 network, regardless of the protocol.



Note

Refer to [Section 2.8, “Firewalls”](#) for more information about implementing firewalls with iptables commands.

2.2.3.5. Use Kerberos Authentication

One of the issues to consider when NIS is used for authentication is that whenever a user logs into a machine, a password hash from the **/etc/shadow** map is sent over the network. If an intruder gains access to a NIS domain and sniffs network traffic, they can collect usernames and password hashes. With enough time, a password cracking program can guess weak passwords, and an attacker can gain access to a valid account on the network.

Since Kerberos uses secret-key cryptography, no password hashes are ever sent over the network, making the system far more secure. Refer to *Managing Single Sign-On and Smart Cards* for more information about Kerberos.

2.2.4. Securing NFS



Important

The version of NFS included in Red Hat Enterprise Linux 6, NFSv4, no longer requires the **portmap** service as outlined in [Section 2.2.2, “Securing Portmap”](#). NFS traffic now utilizes TCP in all versions, rather than UDP, and requires it when using NFSv4. NFSv4 now includes Kerberos user and group authentication, as part of the **RPCSEC_GSS** kernel module. Information on **portmap** is still included, since Red Hat Enterprise Linux 6 supports NFSv2 and NFSv3, both of which utilize **portmap**.

2.2.4.1. Carefully Plan the Network

NFSv2 and NFSv3 traditionally passed data insecurely. All versions of NFS now have the ability to authenticate (and optionally encrypt) ordinary file system operations using Kerberos. Under NFSv4 all operations can use Kerberos; under v2 or v3, file locking and mounting still do not use it. When using NFSv4.0, delegations may be turned off if the clients are behind NAT or a firewall. Refer to the section

on pNFS in the Storage Administration Guide for information on the use of NFSv4.1 to allow delegations to operate through NAT and firewalls.

2.2.4.2. Securing NFS Mount Options

The use of the **mount** command in the **/etc/fstab** file is explained in the Storage Administration Guide. From a security administration point of view it is worthwhile to note that the NFS mount options can also be specified in **/etc/nfsmount.conf**, which can be used to set custom default options.

2.2.4.2.1. Review the NFS Server



Warning

Only export entire file systems. Exporting a subdirectory of a file system can be a security issue. It is possible in some cases for a client to "break out" of the exported part of the file system and get to unexported parts (see the section on subtree checking in the **exports(5)** man page).

Use the **ro** option to export the file system as read-only whenever possible to reduce the number of users able to write to the mounted file system. Only use the **rw** option when specifically required. Refer to the man **exports(5)** page for more information. Allowing write access increases the risk from symlink attacks for example. This includes temporary directories such as **/tmp** and **/usr/tmp**.

Where directories must be mounted with the **rw** option avoid making them world-writable whenever possible to reduce risk. Exporting home directories is also viewed as a risk as some applications store passwords in clear text or weakly encrypted. This risk is being reduced as application code is reviewed and improved. Some users do not set passwords on their SSH keys so this too means home directories present a risk. Enforcing the use of passwords or using Kerberos would mitigate that risk.

Restrict exports only to clients that need access. Use the **showmount -e** command on an NFS server to review what the server is exporting. Do not export anything that is not specifically required.

Do not use the **no_root_squash** option and review existing installations to make sure it is not used. Refer to [Section 2.2.4.4, "Do Not Use the no_root_squash Option"](#) for more information.

The **secure** option is the server-side export option used to restrict exports to "reserved" ports. By default, the server allows client communication only from "reserved" ports (ports numbered less than 1024), because traditionally clients have only allowed "trusted" code (such as in-kernel NFS clients) to use those ports. However, on many networks it is not difficult for anyone to become root on some client, so it is rarely safe for the server to assume that communication from a reserved port is privileged. Therefore the restriction to reserved ports is of limited value; it is better to rely on Kerberos, firewalls, and restriction of exports to particular clients.

Most clients still do use reserved ports when possible. However, reserved ports are a limited resource, so clients (especially those with a large number of NFS mounts) may choose to use higher-numbered ports as well. Linux clients may do this using the "noresvport" mount option. If you wish to allow this on an export, you may do so with the "insecure" export option.

It is good practice not to allow users to login to a server. While reviewing the above settings on an NFS server conduct a review of who and what can access the server.

2.2.4.2.2. Review the NFS Client

Use the **nosuid** option to disallow the use of a **setuid** program. The **nosuid** option disables the **set-user-identifier** or **set-group-identifier** bits. This prevents remote users from gaining higher privileges by running a **setuid** program. Use this option on the client and the server side.

The **noexec** option disables all executable files on the client. Use this to prevent users from inadvertently executing files placed in the file system being shared. The **nosuid** and **noexec** options are standard options for most, if not all, file systems.

Use the **nodelv** option to prevent “device-files” from being processed as a hardware device by the client.

The **resvport** option is a client-side mount option and **secure** is the corresponding server-side export option (see explanation above). It restricts communication to a “reserved port”. The reserved or “well known” ports are reserved for privileged users and processes such as the root user. Setting this option causes the client to use a reserved source port to communicate with the server.

All versions of NFS now support mounting with Kerberos authentication. The mount option to enable this is: **sec=krb5**.

NFSv4 supports mounting with Kerberos using **krb5i** for integrity and **krb5p** for privacy protection. These are used when mounting with **sec=krb5**, but need to be configured on the NFS server. Refer to the man page on exports (**man 5 exports**) for more information.

The NFS man page (**man 5 nfs**) has a “SECURITY CONSIDERATIONS” section which explains the security enhancements in NFSv4 and contains all the NFS specific mount options.

2.2.4.3. Beware of Syntax Errors

The NFS server determines which file systems to export and which hosts to export these directories to by consulting the **/etc/exports** file. Be careful not to add extraneous spaces when editing this file.

For instance, the following line in the **/etc/exports** file shares the directory **/tmp/nfs/** to the host **bob.example.com** with read/write permissions.

```
/tmp/nfs/      bob.example.com(rw)
```

The following line in the **/etc/exports** file, on the other hand, shares the same directory to the host **bob.example.com** with read-only permissions and shares it to the *world* with read/write permissions due to a single space character after the hostname.

```
/tmp/nfs/      bob.example.com (rw)
```

It is good practice to check any configured NFS shares by using the **showmount** command to verify what is being shared:

```
showmount -e <hostname>
```

2.2.4.4. Do Not Use the no_root_squash Option

By default, NFS shares change the root user to the **nfsnobody** user, an unprivileged user account. This changes the owner of all root-created files to **nfsnobody**, which prevents uploading of programs with the setuid bit set.

If **no_root_squash** is used, remote root users are able to change any file on the shared file system and leave applications infected by trojans for other users to inadvertently execute.

2.2.4.5. NFS Firewall Configuration

The ports used for NFS are assigned dynamically by rpcbind, which can cause problems when creating firewall rules. To simplify this process, use the **/etc/sysconfig/nfs** file to specify which ports are to be

used:

- **MOUNTD_PORT** — TCP and UDP port for mountd (rpc.mountd)
- **STATD_PORT** — TCP and UDP port for status (rpc.statd)
- **LOCKD_TCP** — TCP port for nlockmgr (rpc.lockd)
- **LOCKD_UDP** — UDP port nlockmgr (rpc.lockd)

Port numbers specified must not be used by any other service. Configure your firewall to allow the port numbers specified, as well as TCP and UDP port 2049 (NFS).

Run the **rpcinfo -p** command on the NFS server to see which ports and RPC programs are being used.

2.2.5. Securing the Apache HTTP Server

The Apache HTTP Server is one of the most stable and secure services that ships with Red Hat Enterprise Linux. A large number of options and techniques are available to secure the Apache HTTP Server — too numerous to delve into deeply here. The following section briefly explains good practices when running the Apache HTTP Server.

Always verify that any scripts running on the system work as intended *before* putting them into production. Also, ensure that only the root user has write permissions to any directory containing scripts or CGIs. To do this, run the following commands as the root user:

```
chown root <directory_name>
```

```
chmod 755 <directory_name>
```

System administrators should be careful when using the following configuration options (configured in **/etc/httpd/conf/httpd.conf**):

FollowSymLinks

This directive is enabled by default, so be sure to use caution when creating symbolic links to the document root of the Web server. For instance, it is a bad idea to provide a symbolic link to **/**.

Indexes

This directive is enabled by default, but may not be desirable. To prevent visitors from browsing files on the server, remove this directive.

UserDir

The **UserDir** directive is disabled by default because it can confirm the presence of a user account on the system. To enable user directory browsing on the server, use the following directives:

```
UserDir enabled  
UserDir disabled root
```

These directives activate user directory browsing for all user directories other than **/root/**. To add users to the list of disabled accounts, add a space-delimited list of users on the **UserDir disabled** line.

ServerTokens

The **ServerTokens** directive controls the server response header field which is sent back to clients. It includes various information which can be customized using the following parameters:

- ▶ **ServerTokens Full** (default option) — provides all available information (OS type and used modules), for example:

```
Apache/2.0.41 (Unix) PHP/4.2.2 MyMod/1.2
```

- ▶ **ServerTokens Prod** or **ServerTokens ProductOnly** — provides the following information:

```
Apache
```

- ▶ **ServerTokens Major** — provides the following information:

```
Apache/2
```

- ▶ **ServerTokens Minor** — provides the following information:

```
Apache/2.0
```

- ▶ **ServerTokens Min** or **ServerTokens Minimal** — provides the following information:

```
Apache/2.0.41
```

- ▶ **ServerTokens OS** — provides the following information:

```
Apache/2.0.41 (Unix)
```

It is recommended to use the **ServerTokens Prod** option so that a possible attacker does not gain any valuable information about your system.



Important

Do not remove the **IncludesNoExec** directive. By default, the *Server-Side Includes* (SSI) module cannot execute commands. It is recommended that you do not change this setting unless absolutely necessary, as it could, potentially, enable an attacker to execute commands on the system.

Removing httpd Modules

In certain scenarios, it is beneficial to remove certain **httpd** modules to limit the functionality of the HTTP Server. To do so, simply comment out the entire line which loads the module you wish to remove in the `/etc/httpd/conf/httpd.conf` file. For example, to remove the proxy module, comment out the following line by prepending it with a hash sign:

```
#LoadModule proxy_module modules/mod_proxy.so
```

Note that the `/etc/httpd/conf.d/` directory contains configuration files which are used to load

modules as well.

httpd and SELinux

For information regarding the Apache HTTP Server and SELinux, refer to [Chapter 3. The Apache HTTP Server](#) of the *Managing Confined Services* Guide.

2.2.6. Securing FTP

The *File Transfer Protocol* (FTP) is an older TCP protocol designed to transfer files over a network. Because all transactions with the server, including user authentication, are unencrypted, it is considered an insecure protocol and should be carefully configured.

Red Hat Enterprise Linux provides three FTP servers.

- ▶ **gssftpd** — A Kerberos-aware **xinetd**-based FTP daemon that does not transmit authentication information over the network.
- ▶ **Red Hat Content Accelerator (tux)** — A kernel-space Web server with FTP capabilities.
- ▶ **vsftpd** — A standalone, security oriented implementation of the FTP service.

The following security guidelines are for setting up the **vsftpd** FTP service.

2.2.6.1. FTP Greeting Banner

Before submitting a username and password, all users are presented with a greeting banner. By default, this banner includes version information useful to crackers trying to identify weaknesses in a system.

To change the greeting banner for **vsftpd**, add the following directive to the **/etc/vsftpd/vsftpd.conf** file:

```
ftpd_banner=<insert_greeting_here>
```

Replace **<insert_greeting_here>** in the above directive with the text of the greeting message.

For mutli-line banners, it is best to use a banner file. To simplify management of multiple banners, place all banners in a new directory called **/etc/banners/**. The banner file for FTP connections in this example is **/etc/banners/ftp.msg**. Below is an example of what such a file may look like:

```
##### Hello, all activity on ftp.example.com is logged. #####
```



Note

It is not necessary to begin each line of the file with **220** as specified in [Section 2.2.1.1.1, “TCP Wrappers and Connection Banners”](#).

To reference this greeting banner file for **vsftpd**, add the following directive to the **/etc/vsftpd/vsftpd.conf** file:

```
banner_file=/etc/banners/ftp.msg
```

It also is possible to send additional banners to incoming connections using TCP Wrappers as described in [Section 2.2.1.1.1, “TCP Wrappers and Connection Banners”](#).

The presence of the `/var/ftp/` directory activates the anonymous account.

The easiest way to create this directory is to install the **vsftpd** package. This package establishes a directory tree for anonymous users and configures the permissions on directories to read-only for anonymous users.

By default the anonymous user cannot write to any directories.



Warning

If enabling anonymous access to an FTP server, be aware of where sensitive data is stored.

2.2.6.2.1. Anonymous Upload

To allow anonymous users to upload files, it is recommended that a write-only directory be created within `/var/ftp/pub/`. To do this, run the following command as root:

```
~]# mkdir /var/ftp/pub/upload
```

Next, change the permissions so that anonymous users cannot view the contents of the directory:

```
~]# chmod 730 /var/ftp/pub/upload
```

A long format listing of the directory should look like this:

```
~]# ls -ld /var/ftp/pub/upload
drwx-wx---. 2 root ftp 4096 Nov 14 22:57 /var/ftp/pub/upload
```



Warning

Administrators who allow anonymous users to read and write in directories often find that their servers become a repository of stolen software.

Additionally, under **vsftpd**, add the following line to the `/etc/vsftpd/vsftpd.conf` file:

```
anon_upload_enable=YES
```

2.2.6.3. User Accounts

Because FTP transmits unencrypted usernames and passwords over insecure networks for authentication, it is a good idea to deny system users access to the server from their user accounts.

To disable all user accounts in **vsftpd**, add the following directive to `/etc/vsftpd/vsftpd.conf`:

```
local_enable=NO
```

2.2.6.3.1. Restricting User Accounts

To disable FTP access for specific accounts or specific groups of accounts, such as the root user and those with **sudo** privileges, the easiest way is to use a PAM list file as described in [Section 2.1.4.2, “Disallowing Root Access”](#). The PAM configuration file for **vsftpd** is `/etc/pam.d/vsftpd`.

It is also possible to disable user accounts within each service directly.

To disable specific user accounts in **vsftpd**, add the username to **/etc/vsftpd/ftpusers**

2.2.6.4. Use TCP Wrappers To Control Access

Use TCP Wrappers to control access to either FTP daemon as outlined in [Section 2.2.1.1, “Enhancing Security With TCP Wrappers”](#).

2.2.7. Securing Postfix

Postfix is a Mail Transfer Agent (MTA) that uses the Simple Mail Transfer Protocol (SMTP) to deliver electronic messages between other MTAs and to email clients or delivery agents. Although many MTAs are capable of encrypting traffic between one another, most do not, so sending email over any public networks is considered an inherently insecure form of communication.

It is recommended that anyone planning to implement a Postfix server address the following issues.

2.2.7.1. Limiting a Denial of Service Attack

Because of the nature of email, a determined attacker can flood the server with mail fairly easily and cause a denial of service. By setting limits to the following directives in **/etc/postfix/main.cf**, the effectiveness of such attacks is limited.

- ▶ **smtpd_client_connection_rate_limit** — The maximum number of connection attempts any client is allowed to make to this service per time unit (described below). The default value is 0, which means a client can make as many connections per time unit as Postfix can accept. By default, clients in trusted networks are excluded.
- ▶ **anvil_rate_time_unit** — This time unit is used for rate limit calculations. The default value is 60 seconds.
- ▶ **smtpd_client_event_limit_exceptions** — Clients that are excluded from the connection and rate limit commands. By default, clients in trusted networks are excluded.
- ▶ **smtpd_client_message_rate_limit** — The maximum number of message deliveries a client is allowed to request per time unit (regardless of whether or not Postfix actually accepts those messages).
- ▶ **default_process_limit** — The default maximum number of Postfix child processes that provide a given service. This limit can be overruled for specific services in the **master.cf** file. By default the value is 100.
- ▶ **queue_minfree** — The minimum amount of free space in bytes in the queue file system that is needed to receive mail. This is currently used by the Postfix SMTP server to decide if it will accept any mail at all. By default, the Postfix SMTP server rejects **MAIL FROM** commands when the amount of free space is less than 1.5 times the **message_size_limit**. To specify a higher minimum free space limit, specify a **queue_minfree** value that is at least 1.5 times the **message_size_limit**. By default the **queue_minfree** value is 0.
- ▶ **header_size_limit** — The maximum amount of memory in bytes for storing a message header. If a header is larger, the excess is discarded. By default the value is 102400.
- ▶ **message_size_limit** — The maximum size in bytes of a message, including envelope information. By default the value is 10240000.

2.2.7.2. NFS and Postfix

Never put the mail pool directory, **/var/spool/postfix/**, on an NFS shared volume.

Because NFSv2 and NFSv3 do not maintain control over user and group IDs, two or more users can have the same UID, and receive and read each other's mail.



Note

With NFSv4 using Kerberos, this is not the case, since the **SECRPC_GSS** kernel module does not utilize UID-based authentication. However, it is still considered good practice *not* to put the mail spool directory on NFS shared volumes.

2.2.7.3. Mail-only Users

To help prevent local user exploits on the Postfix server, it is best for mail users to only access the Postfix server using an email program. Shell accounts on the mail server should not be allowed and all user shells in the **/etc/passwd** file should be set to **/sbin/nologin** (with the possible exception of the root user).

2.2.7.4. Disable Postfix Network Listening

By default, Postfix is set up to only listen to the local loopback address. You can verify this by viewing the file **/etc/postfix/main.cf**.

View the file **/etc/postfix/main.cf** to ensure that only the following **inet_interfaces** line appears:

```
inet_interfaces = localhost
```

This ensures that Postfix only accepts mail messages (such as cron job reports) from the local system and not from the network. This is the default setting and protects Postfix from a network attack.

For removal of the localhost restriction and allowing Postfix to listen on all interfaces the **inet_interfaces = all** setting can be used.

2.2.8. Securing Sendmail

Sendmail is a Mail Transfer Agent (MTA) that uses the Simple Mail Transfer Protocol (SMTP) to deliver electronic messages between other MTAs and to email clients or delivery agents. Although many MTAs are capable of encrypting traffic between one another, most do not, so sending email over any public networks is considered an inherently insecure form of communication.

It is recommended that anyone planning to implement a Sendmail server address the following issues.

2.2.8.1. Limiting a Denial of Service Attack

Because of the nature of email, a determined attacker can flood the server with mail fairly easily and cause a denial of service. By setting limits to the following directives in **/etc/mail/sendmail.mc**, the effectiveness of such attacks is limited.

- **confCONNECTION_RATE_THROTTLE** — The number of connections the server can receive per second. By default, Sendmail does not limit the number of connections. If a limit is set and reached, further connections are delayed.
- **confMAX_DAEMON_CHILDREN** — The maximum number of child processes that can be spawned by the server. By default, Sendmail does not assign a limit to the number of child processes. If a limit is set and reached, further connections are delayed.
- **confMIN_FREE_BLOCKS** — The minimum number of free blocks which must be available for the server to accept mail. The default is 100 blocks.
- **confMAX_HEADERS_LENGTH** — The maximum acceptable size (in bytes) for a message header.

- **confMAX_MESSAGE_SIZE** — The maximum acceptable size (in bytes) for a single message.

2.2.8.2. NFS and Sendmail

Never put the mail spool directory, **/var/spool/mail/**, on an NFS shared volume.

Because NFSv2 and NFSv3 do not maintain control over user and group IDs, two or more users can have the same UID, and receive and read each other's mail.



Note

With NFSv4 using Kerberos, this is not the case, since the **SECRPC_GSS** kernel module does not utilize UID-based authentication. However, it is still considered good practice *not* to put the mail spool directory on NFS shared volumes.

2.2.8.3. Mail-only Users

To help prevent local user exploits on the Sendmail server, it is best for mail users to only access the Sendmail server using an email program. Shell accounts on the mail server should not be allowed and all user shells in the **/etc/passwd** file should be set to **/sbin/nologin** (with the possible exception of the root user).

2.2.8.4. Disable Sendmail Network Listening

By default, Sendmail is set up to only listen to the local loopback address. You can verify this by viewing the file **/etc/mail/sendmail.mc** to ensure that the following line appears:

```
DAEMON_OPTIONS(`Port=smtp,Addr=127.0.0.1, Name=MTA')dnl
```

This ensures that Sendmail only accepts mail messages (such as cron job reports) from the local system and not from the network. This is the default setting and protects Sendmail from a network attack.

For removal of the localhost restriction, the **Addr=127.0.0.1** string needs to be removed. Changing Sendmail's configuration requires installing the **sendmail-cf** package, then editing the **.mc** file, running **/etc/mail/make** and finally restarting **sendmail**. The **.cf** configuration file will be regenerated. Note that the system clock must be correct and working and that there must not be any system clock time shifts between these actions in order for the configuration file to be automatically regenerated.

2.2.9. Verifying Which Ports Are Listening

Unnecessary open ports should be avoided because it increases the attack surface of your system. If after the system has been in service you find unexpected open ports in listening state, that might be signs of intrusion and it should be investigated.

Issue the following command, as root, from the console to determine which ports are listening for connections from the network:


```

~]# netstat -tanp | grep LISTEN
tcp        0      0 0.0.0.0:45876          0.0.0.0:*
LISTEN     1193/rpc.statd
tcp        0      0 192.168.122.1:53      0.0.0.0:*
LISTEN     1241/dnsmasq
tcp        0      0 127.0.0.1:631         0.0.0.0:*
LISTEN     1783/cupsd
tcp        0      0 127.0.0.1:25          0.0.0.0:*
LISTEN     7696/sendmail
tcp        0      0 0.0.0.0:111           0.0.0.0:*
LISTEN     1167/rpcbind
tcp        0      0 127.0.0.1:30003       0.0.0.0:*
LISTEN     1118/tcsd
tcp        0      0 :::631                :::*
LISTEN     1/init
tcp        0      0 :::35018              :::*
LISTEN     1193/rpc.statd
tcp        0      0 :::111                :::*
LISTEN     1167/rpcbind

```

Review the output of the command with the services needed on the system, turn off what is not specifically required or authorized, repeat the check. Proceed then to make external checks using **nmap** from another system connected via the network to the first system. This can be used to verify the rules in **iptables**. Make a scan for every IP address shown in the **netstat** output (except for localhost 127.0.0.0 or ::1 range) from an external system. Use the **-6** option for scanning an IPv6 address. See **man nmap(1)** for more information.

The following is an example of the command to be issued from the console of another system to determine which ports are listening for TCP connections from the network:

```
~]# nmap -sT -O 192.168.122.1
```

Refer to the man pages for **netstat**, **nmap**, and **services** for more information.

2.2.10. Disable Source Routing

Source routing is an Internet Protocol mechanism that allows an IP packet to carry information, a list of addresses, that tells a router the path the packet must take. There is also an option to record the hops as the route is traversed. The list of hops taken, the "route record", provides the destination with a return path to the source. This allows the source (that is to say, the sending host) to specify the route, loosely or strictly, ignoring the routing tables of some or all of the routers. It can allow a user to redirect network traffic for malicious purposes. Therefore, source-based routing should be disabled.

The **accept_source_route** option causes network interfaces to accept packets with the *Strict Source Route* (SSR) or *Loose Source Routing* (LSR) option set. The acceptance of source routed packets is controlled by sysctl settings. Issue the following command as root to drop packets with the SSR or LSR option set:

```
~]# /sbin/sysctl -w net.ipv4.conf.all.accept_source_route=0
```

Disabling the forwarding of packets should also be done in conjunction with the above when possible (disabling forwarding may interfere with virtualization). Issue the commands listed below as root:

These commands disable forwarding of IPv4 and IPv6 packets on all interfaces.

```
~]# /sbin/sysctl -w net.ipv4.conf.all.forwarding=0
```

```
~]# /sbin/sysctl -w net.ipv6.conf.all.forwarding=0
```

These commands disable forwarding of all multicast packets on all interfaces.

```
~]# /sbin/sysctl -w net.ipv4.conf.all.mc_forwarding=0
```

```
~]# /sbin/sysctl -w net.ipv6.conf.all.mc_forwarding=0
```

Accepting ICMP redirects has few legitimate uses. Disable the acceptance and sending of ICMP redirected packets unless specifically required.

These commands disable acceptance of all ICMP redirected packets on all interfaces.

```
~]# /sbin/sysctl -w net.ipv4.conf.all.accept_redirects=0
```

```
~]# /sbin/sysctl -w net.ipv6.conf.all.accept_redirects=0
```

This command disables acceptance of secure ICMP redirected packets on all interfaces.

```
~]# /sbin/sysctl -w net.ipv4.conf.all.secure_redirects=0
```

This command disables acceptance of all IPv4 ICMP redirected packets on all interfaces.

```
~]# /sbin/sysctl -w net.ipv4.conf.all.send_redirects=0
```

There is only a directive to disable sending of IPv4 redirected packets. Refer to [RFC4294](#) for an explanation of “IPv6 Node Requirements” which resulted in this difference between IPv4 and IPv6.

In order to make the settings permanent they must be added to `/etc/sysctl.conf`.

Refer to the sysctl man page, **sysctl(8)**, for more information. Refer to [RFC791](#) for an explanation of the Internet options related to source based routing and its variants.



Warning

Ethernet networks provide additional ways to redirect traffic, such as ARP or MAC address spoofing, unauthorized DHCP servers, and IPv6 router or neighbor advertisements. In addition, unicast traffic is occasionally broadcast, causing information leaks. These weaknesses can only be addressed by specific countermeasures implemented by the network operator. Host-based countermeasures are not fully effective.

2.2.11. Reverse Path Filtering

Reverse path filtering is used to prevent packets which arrived via one interface from leaving via a different interface. When outgoing routes and incoming routes are different it is sometimes referred to as “asymmetric routing”. Routers often route packets this way but most hosts should not need to do this. Exceptions are such applications as sending traffic out over one link and receiving traffic over another link from a different service provider. For example, using leased lines in combination with xDSL, or Satellite links with 3G modems. If such a scenario is applicable to you then turning off reverse path filtering on the incoming interface is necessary. In short, unless you know that it is required, it is best disabled as it prevents users spoofing IP addresses from local subnets and reduces the opportunity for DDoS attacks.

**Note**

Red Hat Enterprise Linux 6 (unlike Red Hat Enterprise Linux 5) defaults to using Strict Reverse Path filtering. Red Hat Enterprise Linux 6 follows the “Strict Reverse Path” recommendation from RFC 3704, Ingress Filtering for Multihomed Networks. This currently only applies to IPv4 in Red Hat Enterprise Linux 6.

**Warning**

If forwarding is enabled, then Reverse Path Filtering should only be disabled if there are other means for source address validation (such as iptables rules for example).

rp_filter

Reverse Path Filter is enabled by means of the **rp_filter** directive. The **rp_filter** option is used to direct the kernel to select from one of three modes.

It takes the following form when setting the default behavior:

```
~]# /sbin/sysctl -w net.ipv4.conf.default.rp_filter=INTEGER
```

where **INTEGER** is one of the following:

- ▶ **0** — No source validation.
- ▶ **1** — Strict mode as defined in RFC3704.
- ▶ **2** — Loose mode as defined in RFC3704.

The setting can be overridden per network interface using **net.ipv4.interface.rp_filter**. To make these settings persistent across reboot, modify the **/etc/sysctl.conf** file.

2.2.11.1. Additional Resources

The following are resources which explain more about Reverse Path Filtering.

2.2.11.1.1. Installed Documentation

usr/share/doc/kernel-doc-version/Documentation/networking/ip-sysctl.txt

This contains a complete list of files and options available in the **/proc/sys/net/ipv4/** directory.

2.2.11.1.2. Useful Websites

<https://access.redhat.com/knowledge/solutions/53031>

The Red Hat Knowledgebase article about **rp_filter**.

Refer to [RFC3704](#) for an explanation of Ingress Filtering for Multihomed Networks.

2.3. Single Sign-on (SSO)

The Red Hat Enterprise Linux SSO functionality reduces the number of times Red Hat Enterprise Linux desktop users have to enter their passwords. Several major applications leverage the same underlying authentication and authorization mechanisms so that users can log in to Red Hat Enterprise Linux from the log-in screen, and then not need to re-enter their passwords. These applications are detailed below.

For more information on Pluggable Authentication Modules, refer to [Red Hat Enterprise Linux 6 Managing Single Sign-On and Smart Cards](#) guide.

2.4. Pluggable Authentication Modules (PAM)

Pluggable authentication modules are a common framework for authentication and security. Both of Red Hat Enterprise Linux's single sign-on methods — Kerberos and smart cards — depend on underlying PAM configuration.

For more information on Pluggable Authentication Modules, refer to [Chapter 2, "Using Pluggable Authentication Modules \(PAM\)"](#) of the *Red Hat Enterprise Linux 6 Managing Single Sign-On and Smart Cards* guide.

2.5. Kerberos

Maintaining system security and integrity within a network is critical, and it encompasses every user, application, service, and server within the network infrastructure. It requires an understanding of everything that is running on the network and the manner in which these services are used. At the core of maintaining this security is maintaining access to these applications and services and enforcing that access.

Kerberos provides a mechanism that allows both users and machines to identify themselves to network and receive defined, limited access to the areas and services that the administrator configured. Kerberos authenticates entities by verifying their identity, and Kerberos also secures this authenticating data so that it cannot be accessed and used or tampered with by an outsider.

For more information on Pluggable Authentication Modules, refer to [Chapter 3, "Using Kerberos"](#) of the *Red Hat Enterprise Linux 6 Managing Single Sign-On and Smart Cards* guide.

2.6. TCP Wrappers and xinetd

Controlling access to network services is one of the most important security tasks facing a server administrator. Red Hat Enterprise Linux provides several tools for this purpose. For example, an **iptables**-based firewall filters out unwelcome network packets within the kernel's network stack. For network services that utilize it, *TCP Wrappers* add an additional layer of protection by defining which hosts are or are not allowed to connect to "wrapped" network services. One such wrapped network service is the **xinetd** super server. This service is called a super server because it controls connections to a subset of network services and further refines access control.

[Figure 2.4, "Access Control to Network Services"](#) is a basic illustration of how these tools work together to protect network services.

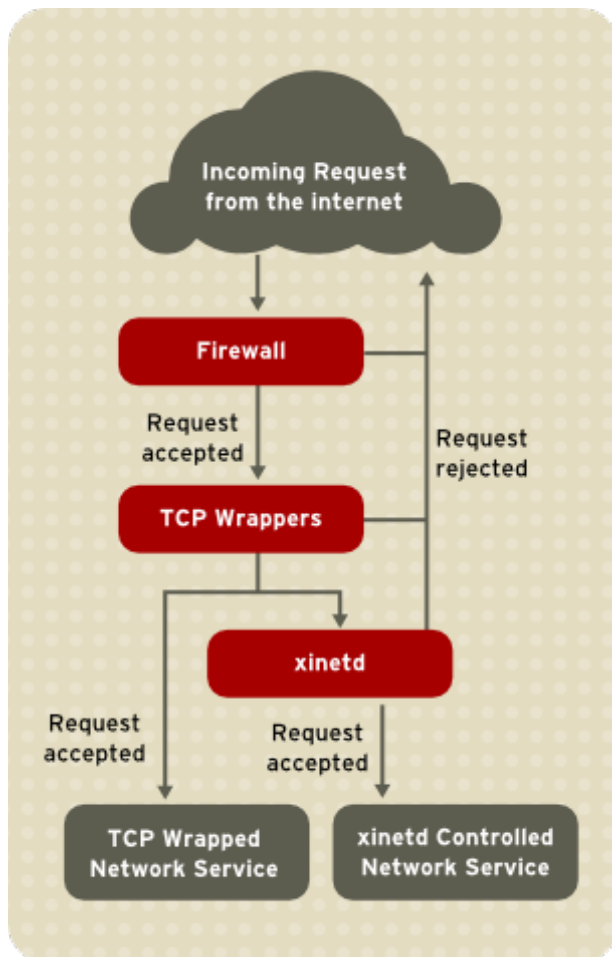


Figure 2.4. Access Control to Network Services

This chapter focuses on the role of TCP Wrappers and **xinetd** in controlling access to network services and reviews how these tools can be used to enhance both logging and utilization management. Refer to [Section 2.8.9, “IPTables”](#) for information about using firewalls with **iptables**.

2.6.1. TCP Wrappers

The TCP Wrappers packages (*tcp_wrappers* and *tcp_wrappers-libs*) are installed by default and provide host-based access control to network services. The most important component within the package is the **/lib/libwrap.so** or **/lib64/libwrap.so** library. In general terms, a TCP-wrapped service is one that has been compiled against the **libwrap.so** library.

When a connection attempt is made to a TCP-wrapped service, the service first references the host's access files (**/etc/hosts.allow** and **/etc/hosts.deny**) to determine whether or not the client is allowed to connect. In most cases, it then uses the syslog daemon (**syslogd**) to write the name of the requesting client and the requested service to **/var/log/secure** or **/var/log/messages**.

If a client is allowed to connect, TCP Wrappers release control of the connection to the requested service and take no further part in the communication between the client and the server.

In addition to access control and logging, TCP Wrappers can execute commands to interact with the client before denying or releasing control of the connection to the requested network service.

Because TCP Wrappers are a valuable addition to any server administrator's arsenal of security tools, most network services within Red Hat Enterprise Linux are linked to the **libwrap.so** library. Such

applications include `/usr/sbin/sshd`, `/usr/sbin/sendmail`, and `/usr/sbin/xinetd`.



Note

To determine if a network service binary is linked to **libwrap.so**, type the following command as the root user:

```
ldd <binary-name> | grep libwrap
```

Replace **<binary-name>** with the name of the network service binary. If the command returns straight to the prompt with no output, then the network service is *not* linked to **libwrap.so**. The following example indicates that `/usr/sbin/sshd` is linked to **libwrap.so**:

```
~]# ldd /usr/sbin/sshd | grep libwrap
libwrap.so.0 => /lib/libwrap.so.0 (0x00655000)
```

2.6.1.1. Advantages of TCP Wrappers

TCP Wrappers provide the following advantages over other network service control techniques:

- ▶ *Transparency to both the client and the wrapped network service* — Both the connecting client and the wrapped network service are unaware that TCP Wrappers are in use. Legitimate users are logged and connected to the requested service while connections from banned clients fail.
- ▶ *Centralized management of multiple protocols* — TCP Wrappers operate separately from the network services they protect, allowing many server applications to share a common set of access control configuration files, making for simpler management.

2.6.2. TCP Wrappers Configuration Files

To determine if a client is allowed to connect to a service, TCP Wrappers reference the following two files, which are commonly referred to as *hosts access files*:

- ▶ **/etc/hosts.allow**
- ▶ **/etc/hosts.deny**

When a TCP-wrapped service receives a client request, it performs the following steps:

1. *It references **/etc/hosts.allow*** — The TCP-wrapped service sequentially parses the **/etc/hosts.allow** file and applies the first rule specified for that service. If it finds a matching rule, it allows the connection. If not, it moves on to the next step.
2. *It references **/etc/hosts.deny*** — The TCP-wrapped service sequentially parses the **/etc/hosts.deny** file. If it finds a matching rule, it denies the connection. If not, it grants access to the service.

The following are important points to consider when using TCP Wrappers to protect network services:

- ▶ Because access rules in **hosts.allow** are applied first, they take precedence over rules specified in **hosts.deny**. Therefore, if access to a service is allowed in **hosts.allow**, a rule denying access to that same service in **hosts.deny** is ignored.
- ▶ The rules in each file are read from the top down and the first matching rule for a given service is the only one applied. The order of the rules is extremely important.
- ▶ If no rules for the service are found in either file, or if neither file exists, access to the service is

granted.

- » TCP-wrapped services do not cache the rules from the hosts access files, so any changes to **hosts.allow** or **hosts.deny** take effect immediately, without restarting network services.



Warning

If the last line of a hosts access file is not a newline character (created by pressing the **Enter** key), the last rule in the file fails and an error is logged to either **/var/log/messages** or **/var/log/secure**. This is also the case for a rule that spans multiple lines without using the backslash character. The following example illustrates the relevant portion of a log message for a rule failure due to either of these circumstances:

```
warning: /etc/hosts.allow, line 20: missing newline or line too long
```

2.6.2.1. Formatting Access Rules

The format for both **/etc/hosts.allow** and **/etc/hosts.deny** is identical. Each rule must be on its own line. Blank lines or lines that start with a hash (#) are ignored.

Each rule uses the following basic format to control access to network services:

```
<daemon list> : <client list> [: <option> : <option> : ...]
```

- » **<daemon list>** — A comma-separated list of process names (*not* service names) or the **ALL** wildcard. The daemon list also accepts operators (refer to [Section 2.6.2.1.4, “Operators”](#)) to allow greater flexibility.
- » **<client list>** — A comma-separated list of hostnames, host IP addresses, special patterns, or wildcards which identify the hosts affected by the rule. The client list also accepts operators listed in [Section 2.6.2.1.4, “Operators”](#) to allow greater flexibility.
- » **<option>** — An optional action or colon-separated list of actions performed when the rule is triggered. Option fields support expansions, launch shell commands, allow or deny access, and alter logging behavior.



Note

More information on some of the terms above can be found elsewhere in this guide:

- » [Section 2.6.2.1.1, “Wildcards”](#)
- » [Section 2.6.2.1.2, “Patterns”](#)
- » [Section 2.6.2.2.4, “Expansions”](#)
- » [Section 2.6.2.2, “Option Fields”](#)

The following is a basic sample hosts access rule:

```
vsftpd : .example.com
```

This rule instructs TCP Wrappers to watch for connections to the FTP daemon (**vsftpd**) from any host in the **example.com** domain. If this rule appears in **hosts.allow**, the connection is accepted. If this

rule appears in **hosts.deny**, the connection is rejected.

The next sample hosts access rule is more complex and uses two option fields:

```
sshd : .example.com \
: spawn /bin/echo `/bin/date` access denied>>/var/log/sshd.log \
: deny
```

Note that each option field is preceded by the backslash (\). Use of the backslash prevents failure of the rule due to length.

This sample rule states that if a connection to the SSH daemon (**sshd**) is attempted from a host in the **example.com** domain, execute the **echo** command to append the attempt to a special log file, and deny the connection. Because the optional **deny** directive is used, this line denies access even if it appears in the **hosts.allow** file. Refer to [Section 2.6.2.2, “Option Fields”](#) for a more detailed look at available options.

2.6.2.1.1. Wildcards

Wildcards allow TCP Wrappers to more easily match groups of daemons or hosts. They are used most frequently in the client list field of access rules.

The following wildcards are available:

- **ALL** — Matches everything. It can be used for both the daemon list and the client list.
- **LOCAL** — Matches any host that does not contain a period (.), such as localhost.
- **KNOWN** — Matches any host where the hostname and host address are known or where the user is known.
- **UNKNOWN** — Matches any host where the hostname or host address are unknown or where the user is unknown.
- **PARANOID** — A reverse DNS lookup is done on the source IP address to obtain the host name. Then a DNS lookup is performed to resolve the IP address. If the two IP addresses do not match the connection is dropped and the logs are updated



Important

The **KNOWN**, **UNKNOWN**, and **PARANOID** wildcards should be used with care, because they rely on a functioning DNS server for correct operation. Any disruption to name resolution may prevent legitimate users from gaining access to a service.

2.6.2.1.2. Patterns

Patterns can be used in the client field of access rules to more precisely specify groups of client hosts.

The following is a list of common patterns for entries in the client field:

- *Hostname beginning with a period (.)* — Placing a period at the beginning of a hostname matches all hosts sharing the listed components of the name. The following example applies to any host within the **example.com** domain:

```
ALL : .example.com
```

- *IP address ending with a period (.)* — Placing a period at the end of an IP address matches all hosts

sharing the initial numeric groups of an IP address. The following example applies to any host within the **192.168.x.x** network:

```
ALL : 192.168.
```

- *IP address/netmask pair* — Netmask expressions can also be used as a pattern to control access to a particular group of IP addresses. The following example applies to any host with an address range of **192.168.0.0** through **192.168.1.255**:

```
ALL : 192.168.0.0/255.255.254.0
```



Important

When working in the IPv4 address space, the address/prefix length (*prefixlen*) pair declarations (CIDR notation) are not supported. Only IPv6 rules can use this format.

- *[IPv6 address]/prefixlen pair* — [net]/prefixlen pairs can also be used as a pattern to control access to a particular group of IPv6 addresses. The following example would apply to any host with an address range of **3ffe:505:2:1::** through **3ffe:505:2:1:ffff:ffff:ffff:ffff**:

```
ALL : [3ffe:505:2:1::]/64
```

- *The asterisk (*)* — Asterisks can be used to match entire groups of hostnames or IP addresses, as long as they are not mixed in a client list containing other types of patterns. The following example would apply to any host within the **example.com** domain:

```
ALL : *.example.com
```

- *The slash (/)* — If a client list begins with a slash, it is treated as a file name. This is useful if rules specifying large numbers of hosts are necessary. The following example refers TCP Wrappers to the **/etc/telnet.hosts** file for all Telnet connections:

```
in.telnetd : /etc/telnet.hosts
```

Other, less used patterns are also accepted by TCP Wrappers. Refer to the **hosts_access** man 5 page for more information.



Warning

Be very careful when using hostnames and domain names. Attackers can use a variety of tricks to circumvent accurate name resolution. In addition, disruption to DNS service prevents even authorized users from using network services. It is, therefore, best to use IP addresses whenever possible.

2.6.2.1.3. Portmap and TCP Wrappers

Portmap's implementation of TCP Wrappers does not support host look-ups, which means **portmap** can not use hostnames to identify hosts. Consequently, access control rules for portmap in **hosts.allow** or **hosts.deny** must use IP addresses, or the keyword **ALL**, for specifying hosts.

Changes to **portmap** access control rules may not take effect immediately. You may need to restart the

portmap service.

Widely used services, such as NIS and NFS, depend on **portmap** to operate, so be aware of these limitations.

2.6.2.1.4. Operators

At present, access control rules accept one operator, **EXCEPT**. It can be used in both the daemon list and the client list of a rule.

The **EXCEPT** operator allows specific exceptions to broader matches within the same rule.

In the following example from a **hosts.allow** file, all **example.com** hosts are allowed to connect to all services except **cracker.example.com**:

```
ALL : .example.com EXCEPT cracker.example.com
```

In another example from a **hosts.allow** file, clients from the **192.168.0.x** network can use all services except for FTP:

```
ALL EXCEPT vsftpd : 192.168.0.
```



Note

Organizationally, it is often easier to avoid using **EXCEPT** operators. This allows other administrators to quickly scan the appropriate files to see what hosts are allowed or denied access to services, without having to sort through **EXCEPT** operators.

2.6.2.2. Option Fields

In addition to basic rules that allow and deny access, the Red Hat Enterprise Linux implementation of TCP Wrappers supports extensions to the access control language through *option fields*. By using option fields in hosts access rules, administrators can accomplish a variety of tasks such as altering log behavior, consolidating access control, and launching shell commands.

2.6.2.2.1. Logging

Option fields let administrators easily change the log facility and priority level for a rule by using the **severity** directive.

In the following example, connections to the SSH daemon from any host in the **example.com** domain are logged to the default **authpriv syslog** facility (because no facility value is specified) with a priority of **emerg**:

```
sshd : .example.com : severity emerg
```

It is also possible to specify a facility using the **severity** option. The following example logs any SSH connection attempts by hosts from the **example.com** domain to the **local0** facility with a priority of **alert**:

```
sshd : .example.com : severity local0.alert
```

**Note**

In practice, this example does not work until the syslog daemon (**syslogd**) is configured to log to the **local0** facility. Refer to the **syslog.conf** man page for information about configuring custom log facilities.

2.6.2.2.2. Access Control

Option fields also allow administrators to explicitly allow or deny hosts in a single rule by adding the **allow** or **deny** directive as the final option.

For example, the following two rules allow SSH connections from **client-1.example.com**, but deny connections from **client-2.example.com**:

```
sshd : client-1.example.com : allow
sshd : client-2.example.com : deny
```

By allowing access control on a per-rule basis, the option field allows administrators to consolidate all access rules into a single file: either **hosts.allow** or **hosts.deny**. Some administrators consider this an easier way of organizing access rules.

2.6.2.2.3. Shell Commands

Option fields allow access rules to launch shell commands through the following two directives:

- **spawn** — Launches a shell command as a child process. This directive can perform tasks like using **/usr/sbin/safe_finger** to get more information about the requesting client or create special log files using the **echo** command.

In the following example, clients attempting to access Telnet services from the **example.com** domain are quietly logged to a special file:

```
in.telnetd : .example.com \
: spawn /bin/echo `/bin/date` from %h>>/var/log/telnet.log \
: allow
```

- **twist** — Replaces the requested service with the specified command. This directive is often used to set up traps for intruders (also called "honey pots"). It can also be used to send messages to connecting clients. The **twist** directive must occur at the end of the rule line.

In the following example, clients attempting to access FTP services from the **example.com** domain are sent a message using the **echo** command:

```
vsftpd : .example.com \
: twist /bin/echo "421 This domain has been black-listed. Access denied!"
```

For more information about shell command options, refer to the **hosts_options** man page.

2.6.2.2.4. Expansions

Expansions, when used in conjunction with the **spawn** and **twist** directives, provide information about the client, server, and processes involved.

The following is a list of supported expansions:

- **%a** — Returns the client's IP address.
- **%A** — Returns the server's IP address.
- **%c** — Returns a variety of client information, such as the username and hostname, or the username and IP address.
- **%d** — Returns the daemon process name.
- **%h** — Returns the client's hostname (or IP address, if the hostname is unavailable).
- **%H** — Returns the server's hostname (or IP address, if the hostname is unavailable).
- **%n** — Returns the client's hostname. If unavailable, **unknown** is printed. If the client's hostname and host address do not match, **paranoid** is printed.
- **%N** — Returns the server's hostname. If unavailable, **unknown** is printed. If the server's hostname and host address do not match, **paranoid** is printed.
- **%p** — Returns the daemon's process ID.
- **%s** — Returns various types of server information, such as the daemon process and the host or IP address of the server.
- **%u** — Returns the client's username. If unavailable, **unknown** is printed.

The following sample rule uses an expansion in conjunction with the **spawn** command to identify the client host in a customized log file.

When connections to the SSH daemon (**sshd**) are attempted from a host in the **example.com** domain, execute the **echo** command to log the attempt, including the client hostname (by using the **%h** expansion), to a special file:

```
sshd : .example.com \  
      : spawn /bin/echo `/bin/date` access denied to %h>>/var/log/sshd.log \  
      : deny
```

Similarly, expansions can be used to personalize messages back to the client. In the following example, clients attempting to access FTP services from the **example.com** domain are informed that they have been banned from the server:

```
vsftpd : .example.com \  
        : twist /bin/echo "421 %h has been banned from this server!"
```

For a full explanation of available expansions, as well as additional access control options, refer to section 5 of the man pages for **hosts_access** (**man 5 hosts_access**) and the man page for **hosts_options**.

Refer to [Section 2.6.5, “Additional Resources”](#) for more information about TCP Wrappers.

2.6.3. xinetd

The **xinetd** daemon is a TCP-wrapped *super service* which controls access to a subset of popular network services, including FTP, IMAP, and Telnet. It also provides service-specific configuration options for access control, enhanced logging, binding, redirection, and resource utilization control.

When a client attempts to connect to a network service controlled by **xinetd**, the super service receives the request and checks for any TCP Wrappers access control rules.

If access is allowed, **xinetd** verifies that the connection is allowed under its own access rules for that service. It also checks that the service is able to have more resources assigned to it and that it is not in breach of any defined rules.

If all these conditions are met (that is, access is allowed to the service; the service has not reached its resource limit; and the service is not in breach of any defined rule), **xinetd** then starts an instance of the requested service and passes control of the connection to it. After the connection has been established, **xinetd** takes no further part in the communication between the client and the server.

2.6.4. xinetd Configuration Files

The configuration files for **xinetd** are as follows:

- ▶ **/etc/xinetd.conf** — The global **xinetd** configuration file.
- ▶ **/etc/xinetd.d/** — The directory containing all service-specific files.

2.6.4.1. The /etc/xinetd.conf File

The **/etc/xinetd.conf** file contains general configuration settings which affect every service under **xinetd**'s control. It is read when the **xinetd** service is first started, so for configuration changes to take effect, you need to restart the **xinetd** service. The following is a sample **/etc/xinetd.conf** file:

```
defaults
{
    instances            = 60
    log_type              = SYSLOG authpriv
    log_on_success        = HOST PID
    log_on_failure        = HOST
    cps                  = 25 30
}
includedir /etc/xinetd.d
```

These lines control the following aspects of **xinetd**:

- ▶ **instances** — Specifies the maximum number of simultaneous requests that **xinetd** can process.
- ▶ **log_type** — Configures **xinetd** to use the **authpriv** log facility, which writes log entries to the **/var/log/secure** file. Adding a directive such as **FILE /var/log/xinetdlog** would create a custom log file called **xinetdlog** in the **/var/log/** directory.
- ▶ **log_on_success** — Configures **xinetd** to log successful connection attempts. By default, the remote host's IP address and the process ID of the server processing the request are recorded.
- ▶ **log_on_failure** — Configures **xinetd** to log failed connection attempts or if the connection was denied.
- ▶ **cps** — Configures **xinetd** to allow no more than 25 connections per second to any given service. If this limit is exceeded, the service is retired for 30 seconds.
- ▶ **includedir /etc/xinetd.d/** — Includes options declared in the service-specific configuration files located in the **/etc/xinetd.d/** directory. Refer to [Section 2.6.4.2, “The /etc/xinetd.d/ Directory”](#) for more information.



Note

Often, both the **log_on_success** and **log_on_failure** settings in **/etc/xinetd.conf** are further modified in the service-specific configuration files. More information may therefore appear in a given service's log file than the **/etc/xinetd.conf** file may indicate. Refer to [Section 2.6.4.3.1, “Logging Options”](#) for further information.

2.6.4.2. The `/etc/xinetd.d/` Directory

The `/etc/xinetd.d/` directory contains the configuration files for each service managed by **xinetd** and the names of the files are correlated to the service. As with **xinetd.conf**, this directory is read only when the **xinetd** service is started. For any changes to take effect, the administrator must restart the **xinetd** service.

The format of files in the `/etc/xinetd.d/` directory use the same conventions as `/etc/xinetd.conf`. The primary reason the configuration for each service is stored in a separate file is to make customization easier and less likely to affect other services.

To gain an understanding of how these files are structured, consider the `/etc/xinetd.d/krb5-telnet` file:

```
service telnet
{
    flags                = REUSE
    socket_type          = stream
    wait                 = no
    user                  = root
    server                = /usr/kerberos/sbin/telnetd
    log_on_failure        += USERID
    disable               = yes
}
```

These lines control various aspects of the **telnet** service:

- ▶ **service** — Specifies the service name, usually one of those listed in the `/etc/services` file.
- ▶ **flags** — Sets any of a number of attributes for the connection. **REUSE** instructs **xinetd** to reuse the socket for a Telnet connection.



Note

The **REUSE** flag is deprecated. All services now implicitly use the **REUSE** flag.

- ▶ **socket_type** — Sets the network socket type to **stream**.
- ▶ **wait** — Specifies whether the service is single-threaded (**yes**) or multi-threaded (**no**).
- ▶ **user** — Specifies which user ID the process runs under.
- ▶ **server** — Specifies which binary executable to launch.
- ▶ **log_on_failure** — Specifies logging parameters for **log_on_failure** in addition to those already defined in **xinetd.conf**.
- ▶ **disable** — Specifies whether the service is disabled (**yes**) or enabled (**no**).

Refer to the **xinetd.conf** man page for more information about these options and their usage.

2.6.4.3. Altering **xinetd** Configuration Files

A range of directives is available for services protected by **xinetd**. This section highlights some of the more commonly used options.

2.6.4.3.1. Logging Options

The following logging options are available for both `/etc/xinetd.conf` and the service-specific configuration files within the `/etc/xinetd.d/` directory.

The following is a list of some of the more commonly used logging options:

- ▶ **ATTEMPT** — Logs the fact that a failed attempt was made (**log_on_failure**).
- ▶ **DURATION** — Logs the length of time the service is used by a remote system (**log_on_success**).
- ▶ **EXIT** — Logs the exit status or termination signal of the service (**log_on_success**).
- ▶ **HOST** — Logs the remote host's IP address (**log_on_failure** and **log_on_success**).
- ▶ **PID** — Logs the process ID of the server receiving the request (**log_on_success**).
- ▶ **USERID** — Logs the remote user using the method defined in RFC 1413 for all multi-threaded stream services (**log_on_failure** and **log_on_success**).

For a complete list of logging options, refer to the **xinetd.conf** man page.

2.6.4.3.2. Access Control Options

Users of **xinetd** services can choose to use the TCP Wrappers hosts access rules, provide access control via the **xinetd** configuration files, or a mixture of both. Refer to [Section 2.6.2, “TCP Wrappers Configuration Files”](#) for more information about TCP Wrappers hosts access control files.

This section discusses using **xinetd** to control access to services.



Note

Unlike TCP Wrappers, changes to access control only take effect if the **xinetd** administrator restarts the **xinetd** service.
Also, unlike TCP Wrappers, access control through **xinetd** only affects services controlled by **xinetd**.

The **xinetd** hosts access control differs from the method used by TCP Wrappers. While TCP Wrappers places all of the access configuration within two files, **/etc/hosts.allow** and **/etc/hosts.deny**, **xinetd**'s access control is found in each service's configuration file in the **/etc/xinetd.d/** directory.

The following hosts access options are supported by **xinetd**:

- ▶ **only_from** — Allows only the specified hosts to use the service.
- ▶ **no_access** — Blocks listed hosts from using the service.
- ▶ **access_times** — Specifies the time range when a particular service may be used. The time range must be stated in 24-hour format notation, HH:MM-HH:MM.

The **only_from** and **no_access** options can use a list of IP addresses or host names, or can specify an entire network. Like TCP Wrappers, combining **xinetd** access control with the enhanced logging configuration can increase security by blocking requests from banned hosts while verbosely recording each connection attempt.

For example, the following **/etc/xinetd.d/telnet** file can be used to block Telnet access from a particular network group and restrict the overall time range that even allowed users can log in:

```

service telnet
{
    disable            = no
    flags              = REUSE
    socket_type        = stream
    wait               = no
    user               = root
    server             = /usr/kerberos/sbin/telnetd
    log_on_failure     += USERID
    no_access          = 172.16.45.0/24
    log_on_success     += PID HOST EXIT
    access_times       = 09:45-16:15
}

```

In this example, when a client system from the **172.16.45.0/24** network, such as **172.16.45.2**, tries to access the Telnet service, it receives the following message:

Connection closed by foreign host.

In addition, their login attempts are logged in **/var/log/messages** as follows:

```

Sep  7 14:58:33 localhost xinetd[5285]: FAIL: telnet address from=172.16.45.107
Sep  7 14:58:33 localhost xinetd[5283]: START: telnet pid=5285 from=172.16.45.107
Sep  7 14:58:33 localhost xinetd[5283]: EXIT: telnet status=0 pid=5285
duration=0(sec)

```

When using TCP Wrappers in conjunction with **xinetd** access controls, it is important to understand the relationship between the two access control mechanisms.

The following is the sequence of events followed by **xinetd** when a client requests a connection:

1. The **xinetd** daemon accesses the TCP Wrappers hosts access rules using a **libwrap.so** library call. If a deny rule matches the client, the connection is dropped. If an allow rule matches the client, the connection is passed to **xinetd**.
2. The **xinetd** daemon checks its own access control rules both for the **xinetd** service and the requested service. If a deny rule matches the client, the connection is dropped. Otherwise, **xinetd** starts an instance of the requested service and passes control of the connection to that service.



Important

Care should be taken when using TCP Wrappers access controls in conjunction with **xinetd** access controls. Misconfiguration can cause undesirable effects.

2.6.4.3.3. Binding and Redirection Options

The service configuration files for **xinetd** support binding the service to an IP address and redirecting incoming requests for that service to another IP address, hostname, or port.

Binding is controlled with the **bind** option in the service-specific configuration files and links the service to one IP address on the system. When this is configured, the **bind** option only allows requests to the correct IP address to access the service. You can use this method to bind different services to different network interfaces based on requirements.

This is particularly useful for systems with multiple network adapters or with multiple IP addresses. On such a system, insecure services (for example, Telnet), can be configured to listen only on the interface connected to a private network and not to the interface connected to the Internet.

The **redirect** option accepts an IP address or hostname followed by a port number. It configures the service to redirect any requests for this service to the specified host and port number. This feature can be used to point to another port number on the same system, redirect the request to a different IP address on the same machine, shift the request to a totally different system and port number, or any combination of these options. A user connecting to a certain service on a system may therefore be rerouted to another system without disruption.

The **xinetd** daemon is able to accomplish this redirection by spawning a process that stays alive for the duration of the connection between the requesting client machine and the host actually providing the service, transferring data between the two systems.

The advantages of the **bind** and **redirect** options are most clearly evident when they are used together. By binding a service to a particular IP address on a system and then redirecting requests for this service to a second machine that only the first machine can see, an internal system can be used to provide services for a totally different network. Alternatively, these options can be used to limit the exposure of a particular service on a multi-homed machine to a known IP address, as well as redirect any requests for that service to another machine especially configured for that purpose.

For example, consider a system that is used as a firewall with this setting for its Telnet service:

```
service telnet
{
    socket_type = stream
    wait       = no
    server      = /usr/kerberos/sbin/telnetd
    log_on_success += DURATION USERID
    log_on_failure += USERID
    bind                = 123.123.123.123
    redirect             = 10.0.1.13 23
}
```

The **bind** and **redirect** options in this file ensure that the Telnet service on the machine is bound to the external IP address (**123.123.123.123**), the one facing the Internet. In addition, any requests for Telnet service sent to **123.123.123.123** are redirected via a second network adapter to an internal IP address (**10.0.1.13**) that only the firewall and internal systems can access. The firewall then sends the communication between the two systems, and the connecting system thinks it is connected to **123.123.123.123** when it is actually connected to a different machine.

This feature is particularly useful for users with broadband connections and only one fixed IP address. When using Network Address Translation (NAT), the systems behind the gateway machine, which are using internal-only IP addresses, are not available from outside the gateway system. However, when certain services controlled by **xinetd** are configured with the **bind** and **redirect** options, the gateway machine can act as a proxy between outside systems and a particular internal machine configured to provide the service. In addition, the various **xinetd** access control and logging options are also available for additional protection.

2.6.4.3.4. Resource Management Options

The **xinetd** daemon can add a basic level of protection from Denial of Service (DoS) attacks. The following is a list of directives which can aid in limiting the effectiveness of such attacks:

- **per_source** — Defines the maximum number of instances for a service per source IP address. It

accepts only integers as an argument and can be used in both **xinetd.conf** and in the service-specific configuration files in the **xinetd.d/** directory.

- **cps** — Defines the maximum number of connections per second. This directive takes two integer arguments separated by white space. The first argument is the maximum number of connections allowed to the service per second. The second argument is the number of seconds that **xinetd** must wait before re-enabling the service. It accepts only integers as arguments and can be used in either the **xinetd.conf** file or the service-specific configuration files in the **xinetd.d/** directory.
- **max_load** — Defines the CPU usage or load average threshold for a service. It accepts a floating point number argument.

The load average is a rough measure of how many processes are active at a given time. See the **uptime**, **who**, and **procinfo** commands for more information about load average.

There are more resource management options available for **xinetd**. Refer to the **xinetd.conf** man page for more information.

2.6.5. Additional Resources

More information about TCP Wrappers and **xinetd** is available from system documentation and on the Internet.

2.6.5.1. Installed TCP Wrappers Documentation

The documentation on your system is a good place to start looking for additional configuration options for TCP Wrappers, **xinetd**, and access control.

- **/usr/share/doc/tcp_wrappers-<version>/** — This directory contains a **README** file that discusses how TCP Wrappers work and the various hostname and host address spoofing risks that exist.
- **/usr/share/doc/xinetd-<version>/** — This directory contains a **README** file that discusses aspects of access control and a **sample.conf** file with various ideas for modifying service-specific configuration files in the **/etc/xinetd.d/** directory.
- TCP Wrappers and **xinetd**-related man pages — A number of man pages exist for the various applications and configuration files involved with TCP Wrappers and **xinetd**. The following are some of the more important man pages:

Server Applications

- **man xinetd** — The man page for **xinetd**.

Configuration Files

- **man 5 hosts_access** — The man page for the TCP Wrappers hosts access control files.
- **man hosts_options** — The man page for the TCP Wrappers options fields.
- **man xinetd.conf** — The man page listing **xinetd** configuration options.

2.6.5.2. Useful TCP Wrappers Websites

- <http://www.docstoc.com/docs/2133633/An-Unofficial-Xinetd-Tutorial> — A thorough tutorial that discusses many different ways to optimize default **xinetd** configuration files to meet specific security goals.

2.6.5.3. Related Books

- *Hacking Linux Exposed* by Brian Hatch, James Lee, and George Kurtz; Osbourne/McGraw-Hill — An

excellent security resource with information about TCP Wrappers and **xinetd**.

2.7. Virtual Private Networks (VPNs)

Organizations with several satellite offices often connect to each other with dedicated lines for efficiency and protection of sensitive data in transit. For example, many businesses use frame relay or *Asynchronous Transfer Mode* (ATM) lines as an end-to-end networking solution to link one office with others. This can be an expensive proposition, especially for small to medium sized businesses (SMBs) that want to expand without paying the high costs associated with enterprise-level, dedicated digital circuits.

To address this need, *Virtual Private Networks* (VPNs) were developed. Following the same functional principles as dedicated circuits, VPNs allow for secured digital communication between two parties (or networks), creating a *Wide Area Network* (WAN) from existing *Local Area Networks* (LANs). Where it differs from frame relay or ATM is in its transport medium. VPNs transmit over IP using datagrams as the transport layer, making it a secure conduit through the Internet to an intended destination. Most free software VPN implementations incorporate open standard encryption methods to further mask data in transit.

Some organizations employ hardware VPN solutions to augment security, while others use software or protocol-based implementations. Several vendors provide hardware VPN solutions, such as Cisco, Nortel, IBM, and Checkpoint. There is a free software-based VPN solution for Linux called FreeS/Wan that utilizes a standardized *Internet Protocol Security* (IPsec) implementation. These VPN solutions, irrespective of whether they are hardware or software based, act as specialized routers that exist between the IP connection from one office to another.

2.7.1. How Does a VPN Work?

When a packet is transmitted from a client, it sends it through the VPN router or gateway, which adds an *Authentication Header* (AH) for routing and authentication. The data is then encrypted and, finally, enclosed with an *Encapsulating Security Payload* (ESP). This latter constitutes the decryption and handling instructions.

The receiving VPN router strips the header information, decrypts the data, and routes it to its intended destination (either a workstation or other node on a network). Using a network-to-network connection, the receiving node on the local network receives the packets already decrypted and ready for processing. The encryption/decryption process in a network-to-network VPN connection is transparent to a local node.

With such a heightened level of security, an attacker must not only intercept a packet, but decrypt the packet as well. Intruders who employ a man-in-the-middle attack between a server and client must also have access to at least one of the private keys for authenticating sessions. Because they employ several layers of authentication and encryption, VPNs are a secure and effective means of connecting multiple remote nodes to act as a unified intranet.

2.7.2. Openswan

2.7.2.1. Overview

Overview

Openswan is an open source, kernel-level IPsec implementation available in Red Hat Enterprise Linux. It employs key establishment protocols IKE (Internet Key Exchange) v1 and v2, implemented as user-level daemons. Manual key establishment is also possible via **ip xfrm** commands, however this is not recommended.

Cryptographic Support

Openswan uses the NSS (Network Security Services) cryptographic library, which is required for FIPS security compliance. More information on the FIPS (Federal Information Processing Standard) can be found in [Section 7.2, “Federal Information Processing Standard \(FIPS\)”](#).

Installation

Run the **yum install openswan** command to install Openswan.

2.7.2.2. Configuration

Locations

This section lists and explains important directories and files used for configuring Openswan.

- ▶ **/etc/ipsec.d** - main directory. Stores Openswan related files.
- ▶ **/etc/ipsec.conf** - master configuration file. Further ***.conf** configuration files can be created in **/etc/ipsec.d** for individual configurations.
- ▶ **/etc/ipsec.secrets** - master secrets file. Further ***.secrets** files can be created in **/etc/ipsec.d** for individual configurations.
- ▶ **/etc/ipsec.d/cert*.db** - Certificate database files. The old default NSS database file is **cert8.db**. From Red Hat Enterprise Linux 6 onwards, NSS sqlite databases are used in the **cert9.db** file.
- ▶ **/etc/ipsec.d/key*.db** - Key database files. The old default NSS database file is **key3.db**. From Red Hat Enterprise Linux 6 onwards, NSS sqlite databases are used in the **key4.db** file.
- ▶ **/etc/ipsec.d/cacerts** - Location for Certificate Authority (CA) certificates.
- ▶ **/etc/ipsec.d/certs** - Location for user certificates. Not needed when using NSS.
- ▶ **/etc/ipsec.d/policies** - Groups policies. Policies can be defined as *block*, *clear*, *clear-or-private*, *private*, *private-or-clear*.
- ▶ **/etc/ipsec.d/nsspassword** - NSS password file. This file does not exist by default, and is required if the NSS database in use is created with a password.

Configuration Parameters

This section lists some of the configuration options available, mostly written to **/etc/ipsec.conf**.

- ▶ **protostack** - defines which protocol stack is used. The default option in Red Hat Enterprise Linux 6 is *netkey*. Other valid values are *auto*, *klips* and *mast*.
- ▶ **nat_traversal** - defines if NAT workaround for connections is accepted. Default is no.
- ▶ **dumpdir** - defines the location for core dump files.
- ▶ **nhelpers** - When using NSS, defines the number of threads used for cryptographic operations. When not using NSS, defines the number of processes used for cryptographic operations.
- ▶ **virtual_private** - subnets allowed for the client connection. Ranges that may exist behind a NAT router through which a client connects.
- ▶ **plutorestartoncrash** - set to yes by default.
- ▶ **plutostderr** - path for pluto error log. Points to syslog location by default.
- ▶ **connaddrfamily** - can be set to either *ipv4* or *ipv6*.

Further details about Openswan configuration can be found in the **ipsec.conf(5)** manual page.

2.7.2.3. Commands

This section explains and gives examples of some of the commands used for Openswan.



Note

As shown in the following example, using **service ipsec start/stop** is the recommended method of changing the state of the ipsec service. This is also the recommended technique for starting and stopping all other services in Red Hat Enterprise Linux 6.

- ▶ Starting and Stopping Openswan:
 - **ipsec setup start/stop**
 - **service ipsec start/stop**
- ▶ Adding/deleting a connection:
 - **ipsec auto --add/delete <connection name>**
- ▶ Connection establishment/breaking:
 - **ipsec auto --up/down <connection-name>**
- ▶ Generating RSA keys:
 - **ipsec newhostkey --configdir /etc/ipsec.d --password password --output /etc/ipsec.d/<name-of-file>**
- ▶ Checking ipsec policies in Kernel:
 - **ip xfrm policy**
 - **ip xfrm state**
- ▶ Creating self-signed certificate:
 - **certutil -S -k rsa -n <ca-cert-nickname> -s "CN=ca-cert-common-name" -w 12 -t "C,C,C" -x -d /etc/ipsec.d**
- ▶ Creating user certificate signed by the previous CA:
 - **certutil -S -k rsa -c <ca-cert-nickname> -n <user-cert-nickname> -s "CN=user-cert-common-name" -w 12 -t "u,u,u" -d /etc/ipsec.d**

2.7.2.4. Openswan Resources

- ▶ <http://www.openswan.org>
- ▶ <http://lists.openswan.org/pipermail/users/>
- ▶ <http://lists.openswan.org/pipermail/dev/>
- ▶ <http://www.mozilla.org/projects/security/pki/nss/>
- ▶ The *Openswan-doc* package: HTML, examples, README.*
- ▶ README.nss

2.8. Firewalls

Information security is commonly thought of as a process and not a product. However, standard security implementations usually employ some form of dedicated mechanism to control access privileges and restrict network resources to users who are authorized, identifiable, and traceable. Red Hat Enterprise Linux includes several tools to assist administrators and security engineers with network-level access control issues.

Firewalls are one of the core components of a network security implementation. Several vendors market firewall solutions catering to all levels of the marketplace: from home users protecting one PC to data center solutions safeguarding vital enterprise information. Firewalls can be stand-alone hardware solutions, such as firewall appliances by Cisco, Nokia, and Sonicwall. Vendors such as Checkpoint, McAfee, and Symantec have also developed proprietary software firewall solutions for home and business markets.

Apart from the differences between hardware and software firewalls, there are also differences in the way firewalls function that separate one solution from another. [Table 2.6, “Firewall Types”](#) details three common types of firewalls and how they function:

Table 2.6. Firewall Types

Method	Description	Advantages	Disadvantages
NAT	<i>Network Address Translation</i> (NAT) places private IP subnetworks behind one or a small pool of public IP addresses, masquerading all requests to one source rather than several. The Linux kernel has built-in NAT functionality through the Netfilter kernel subsystem.	· Can be configured transparently to machines on a LAN. · Protection of many machines and services behind one or more external IP addresses simplifies administration duties. · Restriction of user access to and from the LAN can be configured by opening and closing ports on the NAT firewall/gateway.	· Cannot prevent malicious activity once users connect to a service outside of the firewall.
Packet Filter	A packet filtering firewall reads each data packet that passes through a LAN. It can read and process packets by header information and filters the packet based on sets of programmable rules implemented by the firewall administrator. The Linux kernel has built-in packet filtering functionality through the Netfilter kernel subsystem.	· Customizable through the iptables front-end utility. · Does not require any customization on the client side, as all network activity is filtered at the router level rather than the application level. · Since packets are not transmitted through a proxy, network performance is faster due to direct connection from client to remote host.	· Cannot filter packets for content like proxy firewalls. · Processes packets at the protocol layer, but cannot filter packets at an application layer. · Complex network architectures can make establishing packet filtering rules difficult, especially if coupled with <i>IP masquerading</i> or local subnets and DMZ networks.
Proxy	Proxy firewalls filter all requests of a certain protocol or type from LAN clients to a proxy machine, which then makes those requests to the Internet on behalf of the local client. A proxy machine acts as a buffer between malicious remote users and the internal network client machines.	· Gives administrators control over what applications and protocols function outside of the LAN. · Some proxy servers can cache frequently-accessed data locally rather than having to use the Internet connection to request it. This helps to reduce bandwidth consumption. · Proxy services can be logged and monitored closely, allowing tighter	· Proxies are often application-specific (HTTP, Telnet, etc.), or protocol-restricted (most proxies work with TCP-connected services only). · Application services cannot run behind a proxy, so your application servers must use a separate form of network security. · Proxies can become a network bottleneck, as all requests and transmissions

control over resource
utilization on the network.

are passed through one
source rather than directly
from a client to a remote
service.

2.8.1. Netfilter and IPTables

The Linux kernel features a powerful networking subsystem called *Netfilter*. The Netfilter subsystem provides stateful or stateless packet filtering as well as NAT and IP masquerading services. Netfilter also has the ability to *mangle* IP header information for advanced routing and connection state management. Netfilter is controlled using the **iptables** tool.

2.8.1.1. IPTables Overview

The power and flexibility of Netfilter is implemented using the **iptables** administration tool, a command line tool similar in syntax to its predecessor, **ipchains**, which Netfilter/iptables replaced in the Linux kernel 2.4 and above.

iptables uses the Netfilter subsystem to enhance network connection, inspection, and processing. **iptables** features advanced logging, pre- and post-routing actions, network address translation, and port forwarding, all in one command line interface.

This section provides an overview of **iptables**. For more detailed information, refer to [Section 2.8.9, “IPTables”](#).

2.8.2. Basic Firewall Configuration

Just as a firewall in a building attempts to prevent a fire from spreading, a computer firewall attempts to prevent malicious software from spreading to your computer. It also helps to prevent unauthorized users from accessing your computer.

In a default Red Hat Enterprise Linux installation, a firewall exists between your computer or network and any untrusted networks, for example the Internet. It determines which services on your computer remote users can access. A properly configured firewall can greatly increase the security of your system. It is recommended that you configure a firewall for any Red Hat Enterprise Linux system with an Internet connection.

2.8.2.1. Firewall Configuration Tool

During the **Firewall Configuration** screen of the Red Hat Enterprise Linux installation, you were given the option to enable a basic firewall as well as to allow specific devices, incoming services, and ports.

After installation, you can change this preference by using the **Firewall Configuration Tool**.

To start this application, either select **System** → **Administration** → **Firewall** from the panel, or type **system-config-firewall** at a shell prompt.

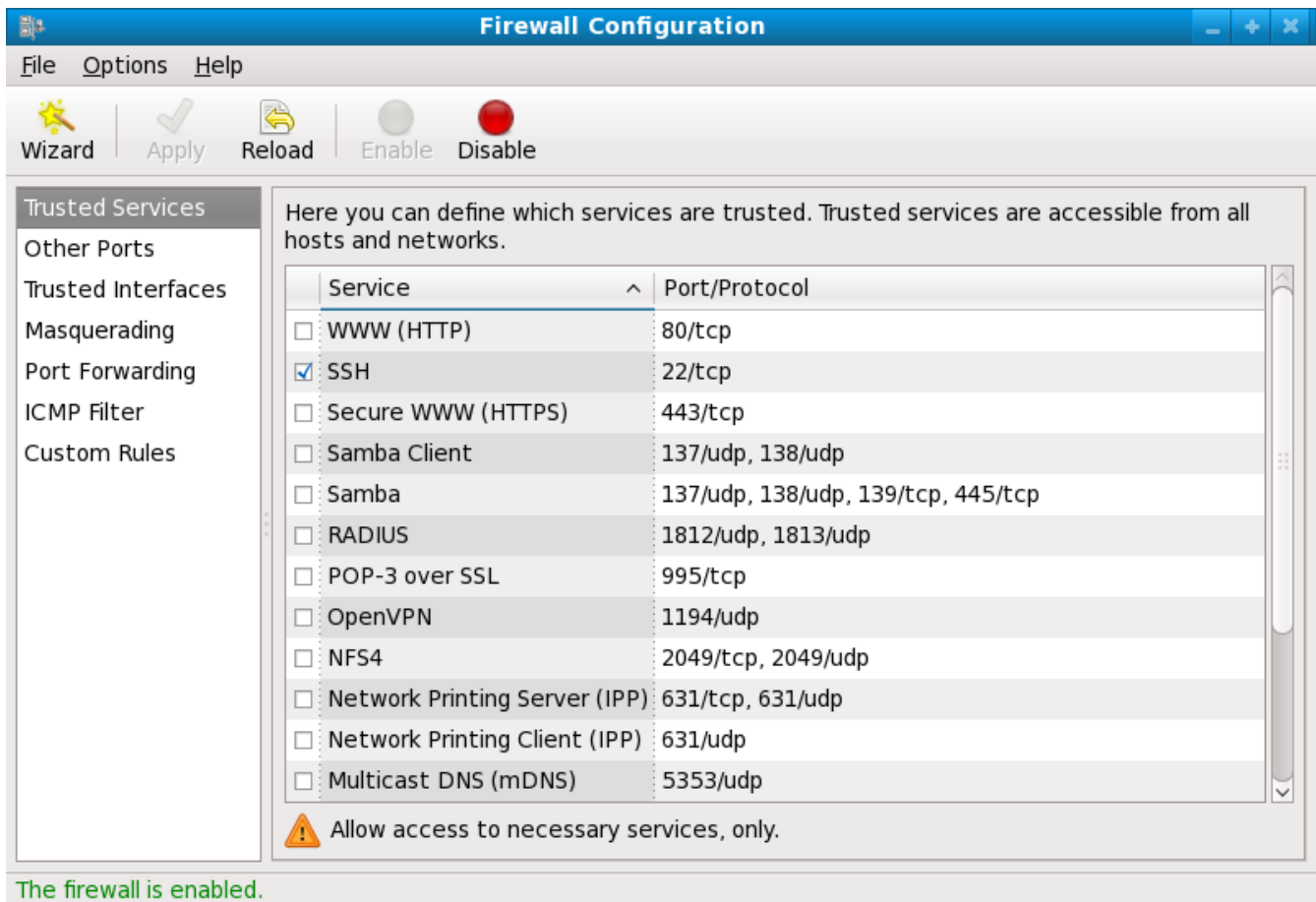


Figure 2.5. Firewall Configuration Tool

Note

The **Firewall Configuration Tool** only configures a basic firewall. If the system needs more complex rules, refer to [Section 2.8.9, “IPTables”](#) for details on configuring specific **iptables** rules.

2.8.2.2. Enabling and Disabling the Firewall

Select one of the following options for the firewall:

- **Disabled** — Disabling the firewall provides complete access to your system and does no security checking. This should only be selected if you are running on a trusted network (not the Internet) or need to configure a custom firewall using the **iptables** command line tool.



Warning

Firewall configurations and any customized firewall rules are stored in the **/etc/sysconfig/iptables** file. If you choose **Disabled** and click **OK**, these configurations and firewall rules will be lost.

- **Enabled** — This option configures the system to reject incoming connections that are not in response to outbound requests, such as DNS replies or DHCP requests. If access to services

running on this machine is needed, you can choose to allow specific services through the firewall. If you are connecting your system to the Internet, but do not plan to run a server, this is the safest choice.

2.8.2.3. Trusted Services

Enabling options in the **Trusted services** list allows the specified service to pass through the firewall.

WWW (HTTP)

The HTTP protocol is used by Apache (and by other Web servers) to serve web pages. If you plan on making your Web server publicly available, select this check box. This option is not required for viewing pages locally or for developing web pages. This service requires that the **httpd** package be installed.

Enabling **WWW (HTTP)** will not open a port for HTTPS, the SSL version of HTTP. If this service is required, select the **Secure WWW (HTTPS)** check box.

FTP

The FTP protocol is used to transfer files between machines on a network. If you plan on making your FTP server publicly available, select this check box. This service requires that the **vsftpd** package be installed.

SSH

Secure Shell (SSH) is a suite of tools for logging into and executing commands on a remote machine. To allow remote access to the machine via ssh, select this check box. This service requires that the **openssh-server** package be installed.

Telnet

Telnet is a protocol for logging into remote machines. Telnet communications are unencrypted and provide no security from network snooping. Allowing incoming Telnet access is not recommended. To allow remote access to the machine via telnet, select this check box. This service requires that the **telnet-server** package be installed.

Mail (SMTP)

SMTP is a protocol that allows remote hosts to connect directly to your machine to deliver mail. You do not need to enable this service if you collect your mail from your ISP's server using POP3 or IMAP, or if you use a tool such as **fetchmail**. To allow delivery of mail to your machine, select this check box. Note that an improperly configured SMTP server can allow remote machines to use your server to send spam.

NFS4

The Network File System (NFS) is a file sharing protocol commonly used on *NIX systems. Version 4 of this protocol is more secure than its predecessors. If you want to share files or directories on your system with other network users, select this check box.

Samba

Samba is an implementation of Microsoft's proprietary SMB networking protocol. If you need to share files, directories, or locally-connected printers with Microsoft Windows machines, select

this check box.

2.8.2.4. Other Ports

The **Firewall Configuration Tool** includes an **Other ports** section for specifying custom IP ports as being trusted by **iptables**. For example, to allow IRC and Internet printing protocol (IPP) to pass through the firewall, add the following to the **Other ports** section:

```
194:tcp,631:tcp
```

2.8.2.5. Saving the Settings

Click **OK** to save the changes and enable or disable the firewall. If **Enable firewall** was selected, the options selected are translated to **iptables** commands and written to the `/etc/sysconfig/iptables` file. The **iptables** service is also started so that the firewall is activated immediately after saving the selected options. If **Disable firewall** was selected, the `/etc/sysconfig/iptables` file is removed and the **iptables** service is stopped immediately.

The selected options are also written to the `/etc/sysconfig/system-config-firewall` file so that the settings can be restored the next time the application is started. Do not edit this file by hand.

Even though the firewall is activated immediately, the **iptables** service is not configured to start automatically at boot time. Refer to [Section 2.8.2.6, “Activating the IPTables Service”](#) for more information.

2.8.2.6. Activating the IPTables Service

The firewall rules are only active if the **iptables** service is running. To manually start the service, use the following command as the root user:

```
~]# service iptables restart
iptables: Applying firewall rules: [ OK ]
```

To ensure that **iptables** starts when the system is booted, use the following command:

```
~]# chkconfig --level 345 iptables on
```

2.8.3. Using IPTables

The first step in using **iptables** is to start the **iptables** service. Use the following command as the root user to start the **iptables** service:

```
~]# service iptables restart
iptables: Applying firewall rules: [ OK ]
```



Note

The **ip6tables** service can be turned off if you intend to use the **iptables** service only. If you deactivate the **ip6tables** service, remember to deactivate the IPv6 network also. Never leave a network device active without the matching firewall.

To force **iptables** to start by default when the system is booted, use the following command as the root user:

```
~]# chkconfig --level 345 iptables on
```

This forces **iptables** to start whenever the system is booted into runlevel 3, 4, or 5.

2.8.3.1. IPTables Command Syntax

The following sample **iptables** command illustrates the basic command syntax:

```
iptables -A <chain> -j <target>
```

The **-A** option specifies that the rule be appended to **<chain>**. Each chain is comprised of one or more *rules*, and is therefore also known as a *ruleset*.

The three built-in chains are INPUT, OUTPUT, and FORWARD. These chains are permanent and cannot be deleted. The chain specifies the point at which a packet is manipulated.

The **-j <target>** option specifies the target of the rule; i.e., what to do if the packet matches the rule. Examples of built-in targets are ACCEPT, DROP, and REJECT.

Refer to the **iptables** man page for more information on the available chains, options, and targets.

2.8.3.2. Basic Firewall Policies

Establishing basic firewall policies creates a foundation for building more detailed, user-defined rules.

Each **iptables** chain is comprised of a default policy, and zero or more rules which work in concert with the default policy to define the overall ruleset for the firewall.

The default policy for a chain can be either DROP or ACCEPT. Security-minded administrators typically implement a default policy of DROP, and only allow specific packets on a case-by-case basis. For example, the following policies block all incoming and outgoing packets on a network gateway:

```
~]# iptables -P INPUT DROP
~]# iptables -P OUTPUT DROP
```

It is also recommended that any *forwarded packets* — network traffic that is to be routed from the firewall to its destination node — be denied as well, to restrict internal clients from inadvertent exposure to the Internet. To do this, use the following rule:

```
~]# iptables -P FORWARD DROP
```

When you have established the default policies for each chain, you can create and save further rules for your particular network and security requirements.

The following sections describe how to save iptables rules and outline some of the rules you might implement in the course of building your iptables firewall.

2.8.3.3. Saving and Restoring IPTables Rules

Changes to **iptables** are transitory; if the system is rebooted or if the **iptables** service is restarted, the rules are automatically flushed and reset. To save the rules so that they are loaded when the **iptables** service is started, use the following command as the root user:

```
~]# service iptables save
iptables: Saving firewall rules to /etc/sysconfig/iptables:[ OK ]
```

The rules are stored in the file `/etc/sysconfig/iptables` and are applied whenever the service is started or the machine is rebooted.

2.8.4. Common IPTables Filtering

Preventing remote attackers from accessing a LAN is one of the most important aspects of network security. The integrity of a LAN should be protected from malicious remote users through the use of stringent firewall rules.

However, with a default policy set to block all incoming, outgoing, and forwarded packets, it is impossible for the firewall/gateway and internal LAN users to communicate with each other or with external resources.

To allow users to perform network-related functions and to use networking applications, administrators must open certain ports for communication.

For example, to allow access to port 80 *on the firewall*, append the following rule:

```
~]# iptables -A INPUT -p tcp -m tcp --dport 80 -j ACCEPT
```

This allows users to browse websites that communicate using the standard port 80. To allow access to secure websites (for example, <https://www.example.com/>), you also need to provide access to port 443, as follows:

```
~]# iptables -A INPUT -p tcp -m tcp --dport 443 -j ACCEPT
```



Important

When creating an **iptables** ruleset, order is important.

If a rule specifies that any packets from the 192.168.100.0/24 subnet be dropped, and this is followed by a rule that allows packets from 192.168.100.13 (which is within the dropped subnet), then the second rule is ignored.

The rule to allow packets from 192.168.100.13 must precede the rule that drops the remainder of the subnet.

To insert a rule in a specific location in an existing chain, use the **-I** option. For example:

```
~]# iptables -I INPUT 1 -i lo -p all -j ACCEPT
```

This rule is inserted as the first rule in the INPUT chain to allow local loopback device traffic.

There may be times when you require remote access to the LAN. Secure services, for example SSH, can be used for encrypted remote connection to LAN services.

Administrators with PPP-based resources (such as modem banks or bulk ISP accounts), dial-up access can be used to securely circumvent firewall barriers. Because they are direct connections, modem connections are typically behind a firewall/gateway.

For remote users with broadband connections, however, special cases can be made. You can configure **iptables** to accept connections from remote SSH clients. For example, the following rules allow remote SSH access:

```
~]# iptables -A INPUT -p tcp --dport 22 -j ACCEPT
~]# iptables -A OUTPUT -p tcp --sport 22 -j ACCEPT
```

These rules allow incoming and outbound access for an individual system, such as a single PC directly connected to the Internet or a firewall/gateway. However, they do not allow nodes behind the firewall/gateway to access these services. To allow LAN access to these services, you can use *Network Address Translation* (NAT) with **iptables** filtering rules.

2.8.5. FORWARD and NAT Rules

Most ISPs provide only a limited number of publicly routable IP addresses to the organizations they serve.

Administrators must, therefore, find alternative ways to share access to Internet services without giving public IP addresses to every node on the LAN. Using private IP addresses is the most common way of allowing all nodes on a LAN to properly access internal and external network services.

Edge routers (such as firewalls) can receive incoming transmissions from the Internet and route the packets to the intended LAN node. At the same time, firewalls/gateways can also route outgoing requests from a LAN node to the remote Internet service.

This forwarding of network traffic can become dangerous at times, especially with the availability of modern cracking tools that can spoof *internal* IP addresses and make the remote attacker's machine act as a node on your LAN.

To prevent this, **iptables** provides routing and forwarding policies that can be implemented to prevent abnormal usage of network resources.

The **FORWARD** chain allows an administrator to control where packets can be routed within a LAN. For example, to allow forwarding for the entire LAN (assuming the firewall/gateway is assigned an internal IP address on `eth1`), use the following rules:

```
~]# iptables -A FORWARD -i eth1 -j ACCEPT
~]# iptables -A FORWARD -o eth1 -j ACCEPT
```

This rule gives systems behind the firewall/gateway access to the internal network. The gateway routes packets from one LAN node to its intended destination node, passing all packets through its **eth1** device.



Note

By default, the IPv4 policy in Red Hat Enterprise Linux kernels disables support for IP forwarding. This prevents machines that run Red Hat Enterprise Linux from functioning as dedicated edge routers. To enable IP forwarding, use the following command as the root user:

```
~]# sysctl -w net.ipv4.ip_forward=1
net.ipv4.ip_forward = 1
```

This configuration change is only valid for the current session; it does not persist beyond a reboot or network service restart. To permanently set IP forwarding, edit the `/etc/sysctl.conf` file as follows:

Locate the following line:

```
net.ipv4.ip_forward = 0
```

Edit it to read as follows:

```
net.ipv4.ip_forward = 1
```

As the root user, run the following command to enable the change to the `sysctl.conf` file:

```
~]# sysctl -p /etc/sysctl.conf
net.ipv4.ip_forward = 1
net.ipv4.conf.default.rp_filter = 1
net.ipv4.conf.default.accept_source_route = 0
[output truncated]
```

2.8.5.1. Postrouting and IP Masquerading

Accepting forwarded packets via the firewall's internal IP device allows LAN nodes to communicate with each other; however they still cannot communicate externally to the Internet.

To allow LAN nodes with private IP addresses to communicate with external public networks, configure the firewall for *IP masquerading*, which masks requests from LAN nodes with the IP address of the firewall's external device (in this case, `eth0`):

```
~]# iptables -t nat -A POSTROUTING -o eth0 -j MASQUERADE
```

This rule uses the NAT packet matching table (`-t nat`) and specifies the built-in POSTROUTING chain for NAT (`-A POSTROUTING`) on the firewall's external networking device (`-o eth0`).

POSTROUTING allows packets to be altered as they are leaving the firewall's external device.

The `-j MASQUERADE` target is specified to mask the private IP address of a node with the external IP address of the firewall/gateway.

2.8.5.2. Prerouting

If you have a server on your internal network that you want make available externally, you can use the `-j DNAT` target of the PREROUTING chain in NAT to specify a destination IP address and port where incoming packets requesting a connection to your internal service can be forwarded.

For example, if you want to forward incoming HTTP requests to your dedicated Apache HTTP Server at 172.31.0.23, use the following command as the root user:

```
~]# iptables -t nat -A PREROUTING -i eth0 -p tcp --dport 80 -j DNAT --to 172.31.0.23:80
```

This rule specifies that the nat table use the built-in PREROUTING chain to forward incoming HTTP requests exclusively to the listed destination IP address of 172.31.0.23.



Note

If you have a default policy of DROP in your FORWARD chain, you must append a rule to forward all incoming HTTP requests so that destination NAT routing is possible. To do this, use the following command as the root user:

```
~]# iptables -A FORWARD -i eth0 -p tcp --dport 80 -d 172.31.0.23 -j ACCEPT
```

This rule forwards all incoming HTTP requests from the firewall to the intended destination; the Apache HTTP Server behind the firewall.

2.8.5.3. DMZs and IPTables

You can create **iptables** rules to route traffic to certain machines, such as a dedicated HTTP or FTP server, in a *demilitarized zone* (DMZ). A DMZ is a special local subnetwork dedicated to providing services on a public carrier, such as the Internet.

For example, to set a rule for routing incoming HTTP requests to a dedicated HTTP server at 10.0.4.2 (outside of the 192.168.1.0/24 range of the LAN), NAT uses the **PREROUTING** table to forward the packets to the appropriate destination:

```
~]# iptables -t nat -A PREROUTING -i eth0 -p tcp --dport 80 -j DNAT \
--to-destination 10.0.4.2:80
```

With this command, all HTTP connections to port 80 from outside of the LAN are routed to the HTTP server on a network separate from the rest of the internal network. This form of network segmentation can prove safer than allowing HTTP connections to a machine on the network.

If the HTTP server is configured to accept secure connections, then port 443 must be forwarded as well.

2.8.6. Malicious Software and Spoofed IP Addresses

More elaborate rules can be created that control access to specific subnets, or even specific nodes, within a LAN. You can also restrict certain dubious applications or programs such as trojans, worms, and other client/server viruses from contacting their server.

For example, some trojans scan networks for services on ports from 31337 to 31340 (called the *elite* ports in cracking terminology).

Since there are no legitimate services that communicate via these non-standard ports, blocking them can effectively diminish the chances that potentially infected nodes on your network independently communicate with their remote master servers.

The following rules drop all TCP traffic that attempts to use port 31337:


```
~]# iptables -A OUTPUT -o eth0 -p tcp --dport 31337 --sport 31337 -j DROP
~]# iptables -A FORWARD -o eth0 -p tcp --dport 31337 --sport 31337 -j DROP
```

You can also block outside connections that attempt to spoof private IP address ranges to infiltrate your LAN.

For example, if your LAN uses the 192.168.1.0/24 range, you can design a rule that instructs the Internet-facing network device (for example, eth0) to drop any packets to that device with an address in your LAN IP range.

Because it is recommended to reject forwarded packets as a default policy, any other spoofed IP address to the external-facing device (eth0) is rejected automatically.

```
~]# iptables -A FORWARD -s 192.168.1.0/24 -i eth0 -j DROP
```



Note

There is a distinction between the **DROP** and **REJECT** targets when dealing with *appended* rules. The **REJECT** target denies access and returns a **connection refused** error to users who attempt to connect to the service. The **DROP** target, as the name implies, drops the packet without any warning. Administrators can use their own discretion when using these targets.

2.8.7. IPTables and Connection Tracking

You can inspect and restrict connections to services based on their *connection state*. A module within **iptables** uses a method called *connection tracking* to store information about incoming connections. You can allow or deny access based on the following connection states:

- ▶ **NEW** — A packet requesting a new connection, such as an HTTP request.
- ▶ **ESTABLISHED** — A packet that is part of an existing connection.
- ▶ **RELATED** — A packet that is requesting a new connection but is part of an existing connection. For example, FTP uses port 21 to establish a connection, but data is transferred on a different port (typically port 20).
- ▶ **INVALID** — A packet that is not part of any connections in the connection tracking table.

You can use the stateful functionality of **iptables** connection tracking with any network protocol, even if the protocol itself is stateless (such as UDP). The following example shows a rule that uses connection tracking to forward only the packets that are associated with an established connection:

```
~]# iptables -A FORWARD -m state --state ESTABLISHED,RELATED -j ACCEPT
```

2.8.8. IPv6

The introduction of the next-generation Internet Protocol, called IPv6, expands beyond the 32-bit address limit of IPv4 (or IP). IPv6 supports 128-bit addresses, and carrier networks that are IPv6 aware are therefore able to address a larger number of routable addresses than IPv4.

Red Hat Enterprise Linux supports IPv6 firewall rules using the Netfilter 6 subsystem and the **ip6tables** command. In Red Hat Enterprise Linux 6, both IPv4 and IPv6 services are enabled by default.

The **ip6tables** command syntax is identical to **iptables** in every aspect except that it supports 128-bit addresses. For example, use the following command to enable SSH connections on an IPv6-aware network server:

```
~]# ip6tables -A INPUT -i eth0 -p tcp -s 3ffe:ffff:100::1/128 --dport 22 -j ACCEPT
```

For more information about IPv6 networking, refer to the IPv6 Information Page at <http://www.ipv6.org/>.

2.8.9. IPTables

Included with Red Hat Enterprise Linux are advanced tools for network *packet filtering* — the process of controlling network packets as they enter, move through, and exit the network stack within the kernel. Kernel versions prior to 2.4 relied on **ipchains** for packet filtering and used lists of rules applied to packets at each step of the filtering process. The 2.4 kernel introduced **iptables** (also called *netfilter*), which is similar to **ipchains** but greatly expands the scope and control available for filtering network packets.

This chapter focuses on packet filtering basics, explains various options available with **iptables** commands, and explains how filtering rules can be preserved between system reboots.



Important

The default firewall mechanism in the 2.4 and later kernels is **iptables**, but **iptables** cannot be used if **ipchains** is already running. If **ipchains** is present at boot time, the kernel issues an error and fails to start **iptables**.

The functionality of **ipchains** is not affected by these errors.

2.8.9.1. Packet Filtering

The Linux kernel uses the **Netfilter** facility to filter packets, allowing some of them to be received by or pass through the system while stopping others. This facility is built in to the Linux kernel, and has three built-in *tables* or *rules lists*, as follows:

- ▶ **filter** — The default table for handling network packets.
- ▶ **nat** — Used to alter packets that create a new connection and used for *Network Address Translation (NAT)*.
- ▶ **mangle** — Used for specific types of packet alteration.

Each table has a group of built-in *chains*, which correspond to the actions performed on the packet by **netfilter**.

The built-in chains for the **filter** table are as follows:

- ▶ **INPUT** — Applies to network packets that are targeted for the host.
- ▶ **OUTPUT** — Applies to locally-generated network packets.
- ▶ **FORWARD** — Applies to network packets routed through the host.

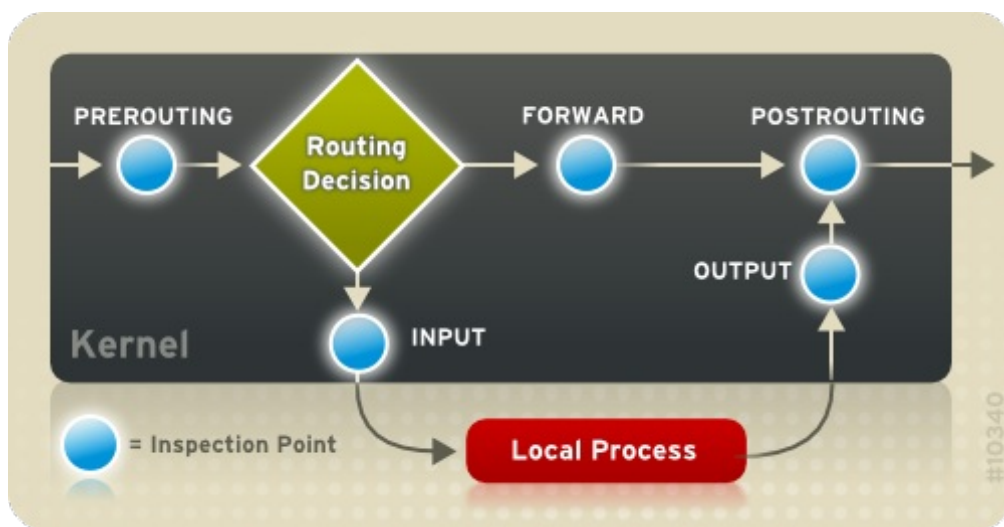
The built-in chains for the **nat** table are as follows:

- ▶ **PREROUTING** — Alters network packets when they arrive.
- ▶ **OUTPUT** — Alters locally-generated network packets before they are sent out.
- ▶ **POSTROUTING** — Alters network packets before they are sent out.

The built-in chains for the **mangle** table are as follows:

- **INPUT** — Alters network packets targeted for the host.
- **OUTPUT** — Alters locally-generated network packets before they are sent out.
- **FORWARD** — Alters network packets routed through the host.
- **PREROUTING** — Alters incoming network packets before they are routed.
- **POSTROUTING** — Alters network packets before they are sent out.

Every network packet received by or sent from a Linux system is subject to at least one table. However, a packet may be subjected to multiple rules within each table before emerging at the end of the chain. The structure and purpose of these rules may vary, but they usually seek to identify a packet coming from or going to a particular IP address, or set of addresses, when using a particular protocol and network service. The following image outlines how the flow of packets is examined by the iptables subsystem:



Important

By default, firewall rules are saved in the **/etc/sysconfig/iptables** or **/etc/sysconfig/ip6tables** files.

The **iptables** service starts before any DNS-related services when a Linux system is booted. This means that firewall rules can only reference numeric IP addresses (for example, 192.168.0.1). Domain names (for example, host.example.com) in such rules produce errors.

Regardless of their destination, when packets match a particular rule in one of the tables, a *target* or action is applied to them. If the rule specifies an **ACCEPT** target for a matching packet, the packet skips the rest of the rule checks and is allowed to continue to its destination. If a rule specifies a **DROP** target, that packet is refused access to the system and nothing is sent back to the host that sent the packet. If a rule specifies a **QUEUE** target, the packet is passed to user-space. If a rule specifies the optional **REJECT** target, the packet is dropped, but an error packet is sent to the packet's originator.

Every chain has a default policy to **ACCEPT**, **DROP**, **REJECT**, or **QUEUE**. If none of the rules in the chain apply to the packet, then the packet is dealt with in accordance with the default policy.

The **iptables** command configures these tables, as well as sets up new tables if necessary.

2.8.9.2. Command Options for IPTables

Rules for filtering packets are created using the **iptables** command. The following aspects of the packet are most often used as criteria:

- » *Packet Type* — Specifies the type of packets the command filters.
- » *Packet Source/Destination* — Specifies which packets the command filters based on the source or destination of the packet.
- » *Target* — Specifies what action is taken on packets matching the above criteria.

Refer to [Section 2.8.9.2.4, “IPTables Match Options”](#) and [Section 2.8.9.2.5, “Target Options”](#) for more information about specific options that address these aspects of a packet.

The options used with specific **iptables** rules must be grouped logically, based on the purpose and conditions of the overall rule, for the rule to be valid. The remainder of this section explains commonly-used options for the **iptables** command.

2.8.9.2.1. Structure of IPTables Command Options

Many **iptables** commands have the following structure:

```
iptables [-t <table-name>] <command> <chain-name> \
  <parameter-1> <option-1> \
  <parameter-n> <option-n>
```

<table-name> — Specifies which table the rule applies to. If omitted, the **filter** table is used.

<command> — Specifies the action to perform, such as appending or deleting a rule.

<chain-name> — Specifies the chain to edit, create, or delete.

<parameter>-<option> pairs — Parameters and associated options that specify how to process a packet that matches the rule.

The length and complexity of an **iptables** command can change significantly, based on its purpose.

For example, a command to remove a rule from a chain can be very short:

```
iptables -D <chain-name> <line-number>
```

In contrast, a command that adds a rule which filters packets from a particular subnet using a variety of specific parameters and options can be rather long. When constructing **iptables** commands, it is important to remember that some parameters and options require further parameters and options to construct a valid rule. This can produce a cascading effect, with the further parameters requiring yet more parameters. Until every parameter and option that requires another set of options is satisfied, the rule is not valid.

Type **iptables -h** to view a comprehensive list of **iptables** command structures.

2.8.9.2.2. Command Options

Command options instruct **iptables** to perform a specific action. Only one command option is allowed per **iptables** command. With the exception of the help command, all commands are written in upper-case characters.

The **iptables** commands are as follows:

- » **-A** — Appends the rule to the end of the specified chain. Unlike the **-I** option described below, it

does not take an integer argument. It always appends the rule to the end of the specified chain.

- **-D <integer> | <rule>** — Deletes a rule in a particular chain by number (such as **5** for the fifth rule in a chain), or by rule specification. The rule specification must exactly match an existing rule.
- **-E** — Renames a user-defined chain. A user-defined chain is any chain other than the default, pre-existing chains. (Refer to the **-N** option, below, for information on creating user-defined chains.) This is a cosmetic change and does not affect the structure of the table.



Note

If you attempt to rename one of the default chains, the system reports a **Match not found** error. You cannot rename the default chains.

- **-F** — Flushes the selected chain, which effectively deletes every rule in the chain. If no chain is specified, this command flushes every rule from every chain.
- **-h** — Provides a list of command structures, as well as a quick summary of command parameters and options.
- **-I [<integer>]** — Inserts the rule in the specified chain at a point specified by a user-defined integer argument. If no argument is specified, the rule is inserted at the top of the chain.



Important

As noted above, the order of rules in a chain determines which rules apply to which packets. This is important to remember when adding rules using either the **-A** or **-I** option. This is especially important when adding rules using the **-I** with an integer argument. If you specify an existing number when adding a rule to a chain, **iptables** adds the new rule *before* (or above) the existing rule.

- **-L** — Lists all of the rules in the chain specified after the command. To list all rules in all chains in the default **filter** table, do not specify a chain or table. Otherwise, the following syntax should be used to list the rules in a specific chain in a particular table:

```
iptables -L <chain-name> -t <table-name>
```

Additional options for the **-L** command option, which provide rule numbers and allow more verbose rule descriptions, are described in [Section 2.8.9.2.6, “Listing Options”](#).

- **-N** — Creates a new chain with a user-specified name. The chain name must be unique, otherwise an error message is displayed.
- **-P** — Sets the default policy for the specified chain, so that when packets traverse an entire chain without matching a rule, they are sent to the specified target, such as **ACCEPT** or **DROP**.
- **-R** — Replaces a rule in the specified chain. The rule's number must be specified after the chain's name. The first rule in a chain corresponds to rule number one.
- **-X** — Deletes a user-specified chain. You cannot delete a built-in chain.
- **-Z** — Sets the byte and packet counters in all chains for a table to zero.

2.8.9.2.3. IPTables Parameter Options

Certain **iptables** commands, including those used to add, append, delete, insert, or replace rules within a particular chain, require various parameters to construct a packet filtering rule.

- **-c** — Resets the counters for a particular rule. This parameter accepts the **PKTS** and **BYTES** options to specify which counter to reset.
- **-d** — Sets the destination hostname, IP address, or network of a packet that matches the rule. When matching a network, the following IP address/netmask formats are supported:
 - **N.N.N.N/M.M.M.M** — Where **N.N.N.N** is the IP address range and **M.M.M.M** is the netmask.
 - **N.N.N.N/M** — Where **N.N.N.N** is the IP address range and **M** is the bitmask.
- **-f** — Applies this rule only to fragmented packets.

You can use the exclamation point character (!) option before this parameter to specify that only unfragmented packets are matched.



Note

Distinguishing between fragmented and unfragmented packets is desirable, despite fragmented packets being a standard part of the IP protocol. Originally designed to allow IP packets to travel over networks with differing frame sizes, these days fragmentation is more commonly used to generate DoS attacks using malformed packets. It's also worth noting that IPv6 disallows fragmentation entirely.

- **-i** — Sets the incoming network interface, such as **eth0** or **ppp0**. With **iptables**, this optional parameter may only be used with the INPUT and FORWARD chains when used with the **filter** table and the PREROUTING chain with the **nat** and **mangle** tables.

This parameter also supports the following special options:

- Exclamation point character (!) — Reverses the directive, meaning any specified interfaces are excluded from this rule.
- Plus character (+) — A wildcard character used to match all interfaces that match the specified string. For example, the parameter **-i eth+** would apply this rule to any Ethernet interfaces but exclude any other interfaces, such as **ppp0**.

If the **-i** parameter is used but no interface is specified, then every interface is affected by the rule.

- **-j** — Jumps to the specified target when a packet matches a particular rule.

The standard targets are **ACCEPT**, **DROP**, **QUEUE**, and **RETURN**.

Extended options are also available through modules loaded by default with the Red Hat Enterprise Linux **iptables** RPM package. Valid targets in these modules include **LOG**, **MARK**, and **REJECT**, among others. Refer to the **iptables** man page for more information about these and other targets.

This option can also be used to direct a packet matching a particular rule to a user-defined chain outside of the current chain so that other rules can be applied to the packet.

If no target is specified, the packet moves past the rule with no action taken. The counter for this rule, however, increases by one.

- **-o** — Sets the outgoing network interface for a rule. This option is only valid for the OUTPUT and FORWARD chains in the **filter** table, and the POSTROUTING chain in the **nat** and **mangle** tables. This parameter accepts the same options as the incoming network interface parameter (**-i**).
- **-p <protocol>** — Sets the IP protocol affected by the rule. This can be either **icmp**, **tcp**, **udp**, or **all**, or it can be a numeric value, representing one of these or a different protocol. You can also use any protocols listed in the **/etc/protocols** file.

The "**all**" protocol means the rule applies to every supported protocol. If no protocol is listed with this rule, it defaults to "**all**".

- **-s** — Sets the source for a particular packet using the same syntax as the destination (**-d**) parameter.

2.8.9.2.4. IPTables Match Options

Different network protocols provide specialized matching options which can be configured to match a particular packet using that protocol. However, the protocol must first be specified in the **iptables** command. For example, **-p <protocol-name>** enables options for the specified protocol. Note that you can also use the protocol ID, instead of the protocol name. Refer to the following examples, each of which have the same effect:

```
~]# iptables -A INPUT -p icmp --icmp-type any -j ACCEPT
~]# iptables -A INPUT -p 5813 --icmp-type any -j ACCEPT
```

Service definitions are provided in the **/etc/services** file. For readability, it is recommended that you use the service names rather than the port numbers.



Warning

Secure the **/etc/services** file to prevent unauthorized editing. If this file is editable, crackers can use it to enable ports on your machine you have otherwise closed. To secure this file, run the following commands as root:

```
~]# chown root.root /etc/services
~]# chmod 0644 /etc/services
~]# chattr +i /etc/services
```

This prevents the file from being renamed, deleted or having links made to it.

2.8.9.2.4.1. TCP Protocol

These match options are available for the TCP protocol (**-p tcp**):

- **--dport** — Sets the destination port for the packet.

To configure this option, use a network service name (such as **www** or **smtp**); a port number; or a range of port numbers.

To specify a range of port numbers, separate the two numbers with a colon (:). For example: **-p tcp --dport 3000:3200**. The largest acceptable valid range is **0:65535**.

Use an exclamation point character (!) after the **--dport** option to match all packets that *do not* use that network service or port.

To browse the names and aliases of network services and the port numbers they use, view the **/etc/services** file.

The **--destination-port** match option is synonymous with **--dport**.

- **--sport** — Sets the source port of the packet using the same options as **--dport**. The **--source-port** match option is synonymous with **--sport**.

- **--syn** — Applies to all TCP packets designed to initiate communication, commonly called *SYN packets*. Any packets that carry a data payload are not touched.

Use an exclamation point character (!) before the **--syn** option to match all non-SYN packets.

- **--tcp-flags <tested flag list> <set flag list>** — Allows TCP packets that have specific bits (flags) set, to match a rule.

The **--tcp-flags** match option accepts two parameters. The first parameter is the mask; a comma-separated list of flags to be examined in the packet. The second parameter is a comma-separated list of flags that must be set for the rule to match.

The possible flags are:

- **ACK**
- **FIN**
- **PSH**
- **RST**
- **SYN**
- **URG**
- **ALL**
- **NONE**

For example, an **iptables** rule that contains the following specification only matches TCP packets that have the SYN flag set and the ACK and FIN flags not set:

--tcp-flags ACK,FIN,SYN SYN

Use the exclamation point character (!) after the **--tcp-flags** to reverse the effect of the match option.

- **--tcp-option** — Attempts to match with TCP-specific options that can be set within a particular packet. This match option can also be reversed by using the exclamation point character (!) after the option.

2.8.9.2.4.2. UDP Protocol

These match options are available for the UDP protocol (**-p udp**):

- **--dport** — Specifies the destination port of the UDP packet, using the service name, port number, or range of port numbers. The **--destination-port** match option is synonymous with **--dport**.
- **--sport** — Specifies the source port of the UDP packet, using the service name, port number, or range of port numbers. The **--source-port** match option is synonymous with **--sport**.

For the **--dport** and **--sport** options, to specify a range of port numbers, separate the two numbers with a colon (:). For example: **-p tcp --dport 3000:3200**. The largest acceptable valid range is 0:65535.

2.8.9.2.4.3. ICMP Protocol

The following match options are available for the Internet Control Message Protocol (ICMP) (**-p icmp**):

- **--icmp-type** — Sets the name or number of the ICMP type to match with the rule. A list of valid ICMP names can be retrieved by typing the **iptables -p icmp -h** command.

2.8.9.2.4.4. Additional Match Option Modules

Additional match options are available through modules loaded by the **iptables** command.

To use a match option module, load the module by name using the **-m <module-name>**, where **<module-name>** is the name of the module.

Many modules are available by default. You can also create modules to provide additional functionality.

The following is a partial list of the most commonly used modules:

- **limit** module — Places limits on how many packets are matched to a particular rule.
When used in conjunction with the **LOG** target, the **limit** module can prevent a flood of matching packets from filling up the system log with repetitive messages or using up system resources.

Refer to [Section 2.8.9.2.5, “Target Options”](#) for more information about the **LOG** target.

The **limit** module enables the following options:

- **--limit** — Sets the maximum number of matches for a particular time period, specified as a **<value>/<period>** pair. For example, using **--limit 5/hour** allows five rule matches per hour.

Periods can be specified in seconds, minutes, hours, or days.

If a number and time modifier are not used, the default value of **3/hour** is assumed.

- **--limit-burst** — Sets a limit on the number of packets able to match a rule at one time. This option is specified as an integer and should be used in conjunction with the **--limit** option.

If no value is specified, the default value of five (5) is assumed.

- **state** module — Enables state matching.

The **state** module enables the following options:

- **--state** — match a packet with the following connection states:
 - **ESTABLISHED** — The matching packet is associated with other packets in an established connection. You need to accept this state if you want to maintain a connection between a client and a server.
 - **INVALID** — The matching packet cannot be tied to a known connection.
 - **NEW** — The matching packet is either creating a new connection or is part of a two-way connection not previously seen. You need to accept this state if you want to allow new connections to a service.
 - **RELATED** — The matching packet is starting a new connection related in some way to an existing connection. An example of this is FTP, which uses one connection for control traffic (port 21), and a separate connection for data transfer (port 20).

These connection states can be used in combination with one another by separating them with commas, such as **-m state --state INVALID,NEW**.

- **mac** module — Enables hardware MAC address matching.

The **mac** module enables the following option:

- **--mac-source** — Matches a MAC address of the network interface card that sent the packet. To exclude a MAC address from a rule, place an exclamation point character (!) after the **--mac-source** match option.

Refer to the **iptables** man page for more match options available through modules.

2.8.9.2.5. Target Options

When a packet has matched a particular rule, the rule can direct the packet to a number of different targets which determine the appropriate action. Each chain has a default target, which is used if none of the rules on that chain match a packet or if none of the rules which match the packet specify a target.

The following are the standard targets:

- **<user-defined-chain>** — A user-defined chain within the table. User-defined chain names must be unique. This target passes the packet to the specified chain.
- **ACCEPT** — Allows the packet through to its destination or to another chain.
- **DROP** — Drops the packet without responding to the requester. The system that sent the packet is not notified of the failure.
- **QUEUE** — The packet is queued for handling by a user-space application.

- **RETURN** — Stops checking the packet against rules in the current chain. If the packet with a **RETURN** target matches a rule in a chain called from another chain, the packet is returned to the first chain to resume rule checking where it left off. If the **RETURN** rule is used on a built-in chain and the packet cannot move up to its previous chain, the default target for the current chain is used.

In addition, extensions are available which allow other targets to be specified. These extensions are called target modules or match option modules and most only apply to specific tables and situations. Refer to [Section 2.8.9.2.4.4, “Additional Match Option Modules”](#) for more information about match option modules.

Many extended target modules exist, most of which only apply to specific tables or situations. Some of the most popular target modules included by default in Red Hat Enterprise Linux are:

- **LOG** — Logs all packets that match this rule. Because the packets are logged by the kernel, the `/etc/syslog.conf` file determines where these log entries are written. By default, they are placed in the `/var/log/messages` file.

Additional options can be used after the **LOG** target to specify the way in which logging occurs:

- **--log-level** — Sets the priority level of a logging event. Refer to the `syslog.conf` man page for a list of priority levels.
- **--log-ip-options** — Logs any options set in the header of an IP packet.
- **--log-prefix** — Places a string of up to 29 characters before the log line when it is written. This is useful for writing syslog filters for use in conjunction with packet logging.



Note

Due to an issue with this option, you should add a trailing space to the **log-prefix** value.

- **--log-tcp-options** — Logs any options set in the header of a TCP packet.
- **--log-tcp-sequence** — Writes the TCP sequence number for the packet in the log.
- **REJECT** — Sends an error packet back to the remote system and drops the packet.

The **REJECT** target accepts **--reject-with <type>** (where **<type>** is the rejection type) allowing more detailed information to be returned with the error packet. The message **port-unreachable** is the default error type given if no other option is used. Refer to the `iptables` man page for a full list of **<type>** options.

Other target extensions, including several that are useful for IP masquerading using the **nat** table, or with packet alteration using the **mangle** table, can be found in the `iptables` man page.

2.8.9.2.6. Listing Options

The default list command, `iptables -L [<chain-name>]`, provides a very basic overview of the default filter table's current chains. Additional options provide more information:

- **-v** — Displays verbose output, such as the number of packets and bytes each chain has processed, the number of packets and bytes each rule has matched, and which interfaces apply to a particular rule.
- **-x** — Expands numbers into their exact values. On a busy system, the number of packets and bytes processed by a particular chain or rule may be abbreviated to **Kilobytes**, **Megabytes**, or **Gigabytes**. This option forces the full number to be displayed.
- **-n** — Displays IP addresses and port numbers in numeric format, rather than the default hostname and network service format.

- **--line-numbers** — Lists rules in each chain next to their numeric order in the chain. This option is useful when attempting to delete the specific rule in a chain or to locate where to insert a rule within a chain.
- **-t <table-name>** — Specifies a table name. If omitted, defaults to the filter table.

2.8.9.3. Saving IPTables Rules

Rules created with the **iptables** command are stored in memory. If the system is restarted before saving the **iptables** rule set, all rules are lost. For netfilter rules to persist through a system reboot, they need to be saved. To save netfilter rules, type the following command as root:

```
~]# /sbin/service iptables save
iptables: Saving firewall rules to /etc/sysconfig/iptables:[ OK ]
```

This executes the **iptables** init script, which runs the **/sbin/iptables-save** program and writes the current **iptables** configuration to **/etc/sysconfig/iptables**. The existing **/etc/sysconfig/iptables** file is saved as **/etc/sysconfig/iptables.save**.

The next time the system boots, the **iptables** init script reapplies the rules saved in **/etc/sysconfig/iptables** by using the **/sbin/iptables-restore** command.

While it is always a good idea to test a new **iptables** rule before committing it to the **/etc/sysconfig/iptables** file, it is possible to copy **iptables** rules into this file from another system's version of this file. This provides a quick way to distribute sets of **iptables** rules to multiple machines.

You can also save the **iptables** rules to a separate file for distribution, backup, or other purposes. To do so, run the following command as root:

```
iptables-save > <filename>
```

... where **<filename>** is a user-defined name for your ruleset.



Important

If distributing the **/etc/sysconfig/iptables** file to other machines, type **/sbin/service iptables restart** for the new rules to take effect.



Note

Note the difference between the **iptables** *command* (**/sbin/iptables**), which is used to manipulate the tables and chains that constitute the **iptables** functionality, and the **iptables** *service* (**/sbin/service iptables**), which is used to enable and disable the **iptables** service itself.

2.8.9.4. IPTables Control Scripts

There are two basic methods for controlling **iptables** in Red Hat Enterprise Linux:

- **Firewall Configuration Tool (system-config-firewall)** — A graphical interface for creating, activating, and saving basic firewall rules. Refer to [Section 2.8.2, “Basic Firewall Configuration”](#) for

more information.

- » **/sbin/service iptables <option>** — Used to manipulate various functions of **iptables** using its initscript. The following options are available:

- **start** — If a firewall is configured (that is, **/etc/sysconfig/iptables** exists), all running **iptables** are stopped completely and then started using the **/sbin/iptables-restore** command. This option only works if the **ipchains** kernel module is *not* loaded. To check if this module is loaded, type the following command as root:

```
~]# lsmod | grep ipchains
```

If this command returns no output, it means the module is not loaded. If necessary, use the **/sbin/rmmod** command to remove the module.

- **stop** — If a firewall is running, the firewall rules in memory are flushed, and all **iptables** modules and helpers are unloaded.

If the **IPTABLES_SAVE_ON_STOP** directive in the **/etc/sysconfig/iptables-config** configuration file is changed from its default value to **yes**, current rules are saved to **/etc/sysconfig/iptables** and any existing rules are moved to the file **/etc/sysconfig/iptables.save**.

Refer to [Section 2.8.9.4.1, “IPTables Control Scripts Configuration File”](#) for more information about the **iptables-config** file.

- **restart** — If a firewall is running, the firewall rules in memory are flushed, and the firewall is started again if it is configured in **/etc/sysconfig/iptables**. This option only works if the **ipchains** kernel module is not loaded.

If the **IPTABLES_SAVE_ON_RESTART** directive in the **/etc/sysconfig/iptables-config** configuration file is changed from its default value to **yes**, current rules are saved to **/etc/sysconfig/iptables** and any existing rules are moved to the file **/etc/sysconfig/iptables.save**.

Refer to [Section 2.8.9.4.1, “IPTables Control Scripts Configuration File”](#) for more information about the **iptables-config** file.

- **status** — Displays the status of the firewall and lists all active rules.

The default configuration for this option displays IP addresses in each rule. To display domain and hostname information, edit the **/etc/sysconfig/iptables-config** file and change the value of **IPTABLES_STATUS_NUMERIC** to **no**. Refer to [Section 2.8.9.4.1, “IPTables Control Scripts Configuration File”](#) for more information about the **iptables-config** file.

- **panic** — Flushes all firewall rules. The policy of all configured tables is set to **DROP**.

This option could be useful if a server is known to be compromised. Rather than physically disconnecting from the network or shutting down the system, you can use this option to stop all further network traffic but leave the machine in a state ready for analysis or other forensics.

- **save** — Saves firewall rules to **/etc/sysconfig/iptables** using **iptables-save**. Refer to [Section 2.8.9.3, “Saving IPTables Rules”](#) for more information.



Note

To use the same initscript commands to control netfilter for IPv6, substitute **ip6tables** for **iptables** in the **/sbin/service** commands listed in this section. For more information about IPv6 and netfilter, refer to [Section 2.8.9.5, “IPTables and IPv6”](#).

2.8.9.4.1. IPTables Control Scripts Configuration File

The behavior of the **iptables** initscripts is controlled by the **/etc/sysconfig/iptables-config** configuration file. The following is a list of directives contained in this file:

- ▶ **IPTABLES_MODULES** — Specifies a space-separated list of additional **iptables** modules to load when a firewall is activated. These can include connection tracking and NAT helpers.
- ▶ **IPTABLES_MODULES_UNLOAD** — Unloads modules on restart and stop. This directive accepts the following values:
 - **yes** — The default value. This option must be set to achieve a correct state for a firewall restart or stop.
 - **no** — This option should only be set if there are problems unloading the netfilter modules.
- ▶ **IPTABLES_SAVE_ON_STOP** — Saves current firewall rules to **/etc/sysconfig/iptables** when the firewall is stopped. This directive accepts the following values:
 - **yes** — Saves existing rules to **/etc/sysconfig/iptables** when the firewall is stopped, moving the previous version to the **/etc/sysconfig/iptables.save** file.
 - **no** — The default value. Does not save existing rules when the firewall is stopped.
- ▶ **IPTABLES_SAVE_ON_RESTART** — Saves current firewall rules when the firewall is restarted. This directive accepts the following values:
 - **yes** — Saves existing rules to **/etc/sysconfig/iptables** when the firewall is restarted, moving the previous version to the **/etc/sysconfig/iptables.save** file.
 - **no** — The default value. Does not save existing rules when the firewall is restarted.
- ▶ **IPTABLES_SAVE_COUNTER** — Saves and restores all packet and byte counters in all chains and rules. This directive accepts the following values:
 - **yes** — Saves the counter values.
 - **no** — The default value. Does not save the counter values.
- ▶ **IPTABLES_STATUS_NUMERIC** — Outputs IP addresses in numeric form instead of domain or hostnames. This directive accepts the following values:
 - **yes** — The default value. Returns only IP addresses within a status output.
 - **no** — Returns domain or hostnames within a status output.

2.8.9.5. IPTables and IPv6

If the **iptables-ipv6** package is installed, netfilter in Red Hat Enterprise Linux can filter the next-generation IPv6 Internet protocol. The command used to manipulate the IPv6 netfilter is **ip6tables**.

Most directives for this command are identical to those used for **iptables**, except the **nat** table is not yet supported. This means that it is not yet possible to perform IPv6 network address translation tasks, such as masquerading and port forwarding.

Rules for **ip6tables** are saved in the **/etc/sysconfig/ip6tables** file. Previous rules saved by the **ip6tables** initscripts are saved in the **/etc/sysconfig/ip6tables.save** file.

Configuration options for the **ip6tables** init script are stored in **/etc/sysconfig/ip6tables-config**, and the names for each directive vary slightly from their **iptables** counterparts.

For example, the **iptables-config** directive **IPTABLES_MODULES**:the equivalent in the **ip6tables-config** file is **IP6TABLES_MODULES**.

2.8.9.6. Additional Resources

There are several aspects to firewalls and the Linux Netfilter subsystem that could not be covered in

this chapter. For more information, refer to the following resources.

2.8.9.6.1. Useful Firewall Websites

- ▶ <http://www.netfilter.org/> — The home of the netfilter/iptables project. Contains assorted information about **iptables**, including a FAQ addressing specific problems and various helpful guides by Rusty Russell, the Linux IP firewall maintainer. The HOWTO documents on the site cover subjects such as basic networking concepts, kernel packet filtering, and NAT configurations.
- ▶ <http://www.tldp.org/> — The Linux Documentation Project contains several useful guides relating to firewall creation and administration.
- ▶ <http://www.iana.org/assignments/port-numbers> — The official list of registered and common service ports as assigned by the Internet Assigned Numbers Authority.

2.8.9.6.2. Related Documentation

- ▶ *Red Hat Linux Firewalls*, by Bill McCarty; Red Hat Press — a comprehensive reference to building network and server firewalls using open source packet filtering technology such as Netfilter and **iptables**. It includes topics that cover analyzing firewall logs, developing firewall rules, and customizing your firewall using various graphical tools.
- ▶ *Linux Firewalls*, by Robert Ziegler; New Riders Press — contains a wealth of information on building firewalls using both 2.2 kernel **ipchains** as well as Netfilter and **iptables**. Additional security topics such as remote access issues and intrusion detection systems are also covered.

2.8.9.6.3. Installed IP Tables Documentation

- ▶ **man iptables** — Contains a description of **iptables** as well as a comprehensive list of targets, options, and match extensions.

2.8.9.6.4. Useful IP Tables Websites

- ▶ http://www.linuxnewbie.org/nhf/Security/IPtables_Basics.html — An introduction to the way packets move through the Linux kernel, plus an introduction to constructing basic **iptables** commands.

[11] Since system BIOSes differ between manufacturers, some may not support password protection of either type, while others may support one type but not the other.

[12] GRUB also accepts unencrypted passwords, but it is recommended that an MD5 hash be used for added security.

Chapter 3. Encryption

There are two main types of data that must be protected: data at rest and data in motion. These different types of data are protected in similar ways using similar technology but the implementations can be completely different. No single protective implementation can prevent all possible methods of compromise as the same information may be at rest and in motion at different points in time.

3.1. Data at Rest

Data at rest is data that is stored on a hard drive, tape, CD, DVD, disk, or other media. This information's biggest threat comes from being physically stolen. Laptops in airports, CDs going through the mail, and backup tapes that get left in the wrong places are all examples of events where data can be compromised through theft. If the data was encrypted on the media then you wouldn't have to worry as much about the data being compromised.

3.1.1. Full Disk Encryption

Full disk or partition encryption is one of the best ways of protecting your data. Not only is each file protected but also the temporary storage that may contain parts of these files is also protected. Full disk encryption will protect all of your files so you don't have to worry about selecting what you want to protect and possibly missing a file.

Red Hat Enterprise Linux 6 natively supports LUKS Encryption. LUKS will bulk encrypt your hard drive partitions so that while your computer is off your data is protected. This will also protect your computer from attackers attempting to use single-user-mode to login to your computer or otherwise gain access.

Full disk encryption solutions like LUKS only protect the data when your computer is off. Once the computer is on and LUKS has decrypted the disk, the files on that disk are available to anyone who would normally have access to them. To protect your files when the computer is on, use full disk encryption in combination with another solution such as file based encryption. Also remember to lock your computer whenever you are away from it. A passphrase protected screen saver set to activate after a few minutes of inactivity is a good way to keep intruders out. For more information on LUKS, refer to [Section 3.2.4, "LUKS Disk Encryption"](#).

3.1.2. File Based Encryption

GnuPG (GPG) is an open source version of PGP that allows you to sign and/or encrypt a file or an email message. This is useful to maintain integrity of the message or file and also protects the confidentiality of the information contained within the file or email. In the case of email, GPG provides dual protection. Not only can it provide Data at Rest protection but also Data In Motion protection once the message has been sent across the network.

File based encryption is intended to protect a file after it has left your computer, such as when you send a CD through the mail. Some file based encryption solutions will leave remnants of the encrypted files that an attacker who has physical access to your computer can recover under some circumstances. To protect the contents of those files from attackers who may have access to your computer, use file based encryption combined with another solution such as full disk encryption.

3.2. Data in Motion

Data in motion is data that is being transmitted over a network. The biggest threats to data in motion are interception and alteration. Your username and password should never be transmitted over a network without protection as it could be intercepted and used by someone else to impersonate you or gain access to sensitive information. Other private information such as bank account information should also be protected when transmitted across a network. If the network session was encrypted then you would

not have to worry as much about the data being compromised while it is being transmitted.

Data in motion is particularly vulnerable to attackers because the attacker does not have to be near the computer in which the data is being stored rather they only have to be somewhere along the path. Encryption tunnels can protect data along the path of communications.

3.2.1. Virtual Private Networks

Virtual Private Networks (VPN) provide encrypted tunnels between computers or networks of computers across all ports. With a VPN in place, all network traffic from the client is forwarded to the server through the encrypted tunnel. This means that the client is logically on the same network as the server it is connected to via the VPN. VPNs are very common and are simple to use and setup.

3.2.2. Secure Shell

Secure Shell (SSH) is a powerful network protocol used to communicate with another system over a secure channel. The transmissions over SSH are encrypted and protected from interception. Cryptographic log-on can also be utilized to provide a better authentication method over traditional usernames and passwords.

SSH is very easy to activate. By simply starting the `sshd` service, the system will begin to accept connections and will allow access to the system when a correct username and password is provided during the connection process. The standard TCP port for the SSH service is 22, however this can be changed by modifying the configuration file `/etc/ssh/sshd_config` and restarting the service. This file also contains other configuration options for SSH.

Secure Shell (SSH) also provides encrypted tunnels between computers but only using a single port. [Port forwarding can be done over an SSH tunnel](#) and traffic will be encrypted as it passes over that tunnel but using port forwarding is not as fluid as a VPN.

3.2.2.1. SSH Cryptographic Login

SSH supports the use of cryptographic keys to login to a computer. This is much more secure than using a password and if setup properly could be considered multifactor authentication.

A configuration change must occur before cryptographic login can occur. In the file `/etc/ssh/sshd_config` uncomment and modify the following lines so that appear as such:

```
PubkeyAuthentication yes
AuthorizedKeysFile .ssh/authorized_keys
```

The first line tells the SSH program to allow public key authentication. The second line points to a file in the home directory where the public key of authorized key pairs exists on the system.

The next thing to do is to generate the SSH key pairs on the client you will use to connect to the system. The command **ssh-keygen** will generate an RSA 2048-bit key set for logging into the system. The keys are stored, by default, in the `~/.ssh` directory. You can utilize the switch **-b** to modify the bit-strength of the key. Using 2048-bit is normally sufficient.

In your `~/.ssh` directory you should see the two keys you just created. If you accepted the defaults when running the **ssh-keygen** then your keys are named `id_rsa` and `id_rsa.pub`, the private and public keys. You should always protect the private key from exposure. The public key, however, needs to be transferred over to the system you are going to login to. Once you have it on your system the easiest way to add the key to the approved list is by:

```
$ cat id_rsa.pub >> ~/.ssh/authorized_keys
```


This will append the public key to the `authorized_key` file. The SSH application will check this file when you attempt to login to the computer.

Similarly to passwords and any other authentication mechanism, you should change your SSH keys regularly. When you do, make sure you clean out any unused keys from the **authorized_key** file.

3.2.3. OpenSSL Intel AES-NI Engine

The Intel Advanced Encryption Standard (AES) New Instructions (AES-NI) engine is available for certain Intel processors, and allows for extremely fast hardware encryption and decryption.



Note

For a list of Intel processors that support the AES-NI engine, refer to: [Intel's ARK](https://ark.intel.com/).

The AES-NI engine is automatically enabled if the detected processor is among the supported ones.

To check if the engine is enabled, run the following command as the root user:

```
openssl engine -c -tt
```

To test its speed, run the following command as root:

```
openssl speed aes-128-cbc
```

To test the speed of OpenSSH you can run a command like the following:

```
~]# dd if=/dev/zero count=100 bs=1M | ssh -c aes128-cbc localhost "cat  
>/dev/null"  
root@localhost's password:  
100+0 records in  
100+0 records out  
104857600 bytes (105 MB) copied, 4.81868 s, 21.8 MB/s
```

You can find out more about the AES-NI engine at the following URL: <http://software.intel.com/en-us/articles/intel-advanced-encryption-standard-instructions-aes-ni/>.

3.2.4. LUKS Disk Encryption

Linux Unified Key Setup-on-disk-format (or LUKS) allows you to encrypt partitions on your Linux computer. This is particularly important when it comes to mobile computers and removable media. LUKS allows multiple user keys to decrypt a master key which is used for the bulk encryption of the partition.

Overview of LUKS

What LUKS does

- ▶ LUKS encrypts entire block devices and is therefore well-suited for protecting the contents of mobile devices such as removable storage media or laptop disk drives.
- ▶ The underlying contents of the encrypted block device are arbitrary. This makes it useful for encrypting **swap** devices. This can also be useful with certain databases that use specially formatted block devices for data storage.
- ▶ LUKS uses the existing device mapper kernel subsystem.
- ▶ LUKS provides passphrase strengthening which protects against dictionary attacks.

- ▶ LUKS devices contain multiple key slots, allowing users to add backup keys/passphrases.

What LUKS does *not* do:

- ▶ LUKS is not well-suited for applications requiring many (more than eight) users to have distinct access keys to the same device.
- ▶ LUKS is not well-suited for applications requiring file-level encryption.

3.2.4.1. LUKS Implementation in Red Hat Enterprise Linux

Red Hat Enterprise Linux 6 utilizes LUKS to perform file system encryption. By default, the option to encrypt the file system is unchecked during the installation. If you select the option to encrypt your hard drive, you will be prompted for a passphrase that will be asked every time you boot the computer. This passphrase "unlocks" the bulk encryption key that is used to decrypt your partition. If you choose to modify the default partition table you can choose which partitions you want to encrypt. This is set in the partition table settings.

The default cipher used for LUKS (refer to **cryptsetup --help**) is aes-cbc-essiv:sha256 (ESSIV - Encrypted Salt-Sector Initialization Vector). Note that the installation program, **Anaconda**, uses by default XTS mode (aes-xts-plain64). The default key size for LUKS is 256 bits. The default key size for LUKS with **Anaconda** (XTS mode) is 512 bits. Ciphers that are available are:

- ▶ AES - Advanced Encryption Standard - [FIPS PUB 197](#)
- ▶ Twofish (A 128-bit Block Cipher)
- ▶ Serpent
- ▶ cast5 - [RFC 2144](#)
- ▶ cast6 - [RFC 2612](#)

3.2.4.2. Manually Encrypting Directories



Warning

Following this procedure will remove all data on the partition that you are encrypting. You **WILL** lose all your information! Make sure you backup your data to an external source before beginning this procedure!

1. Enter runlevel 1 by typing the following at a shell prompt as root:

```
telinit 1
```

2. Unmount your existing **/home**:

```
umount /home
```

3. If the command in the previous step fails, use **fuser** to find processes hogging **/home** and kill them:

```
fuser -mvk /home
```

4. Verify **/home** is no longer mounted:

```
grep home /proc/mounts
```

5. Fill your partition with random data:

```
shred -v --iterations=1 /dev/VG00/LV_home
```

This command proceeds at the sequential write speed of your device and may take some time to complete. It is an important step to ensure no unencrypted data is left on a used device, and to obfuscate the parts of the device that contain encrypted data as opposed to just random data.

6. Initialize your partition:

```
cryptsetup --verbose --verify-passphrase luksFormat /dev/VG00/LV_home
```

7. Open the newly encrypted device:

```
cryptsetup luksOpen /dev/VG00/LV_home home
```

8. Make sure the device is present:

```
ls -l /dev/mapper | grep home
```

9. Create a file system:

```
mkfs.ext3 /dev/mapper/home
```

10. Mount the file system:

```
mount /dev/mapper/home /home
```

11. Make sure the file system is visible:

```
df -h | grep home
```

12. Add the following to the **/etc/crypttab** file:

```
home /dev/VG00/LV_home none
```

13. Edit the **/etc/fstab** file, removing the old entry for **/home** and adding the following line:

```
/dev/mapper/home /home ext3 defaults 1 2
```

14. Restore default SELinux security contexts:

```
/sbin/restorecon -v -R /home
```

15. Reboot the machine:

```
shutdown -r now
```

16. The entry in the **/etc/crypttab** makes your computer ask your **luks** passphrase on boot.

17. Log in as root and restore your backup.

You now have an encrypted partition for all of your data to safely rest while the computer is off.

3.2.4.3. Add a new passphrase to an existing device

Use the following command to add a new passphrase to an existing device:

```
cryptsetup luksAddKey <device>
```

After being prompted for any one of the existing passphrases for authentication, you will be prompted to enter the new passphrase.

3.2.4.4. Remove a passphrase from an existing device

Use the following command to remove a passphrase from an existing device:

```
cryptsetup luksRemoveKey <device>
```

You will be prompted for the passphrase you wish to remove and then for any one of the remaining passphrases for authentication.

3.2.4.5. Creating Encrypted Block Devices in Anaconda

You can create encrypted devices during system installation. This allows you to easily configure a system with encrypted partitions.

To enable block device encryption, check the **Encrypt System** checkbox when selecting automatic partitioning or the **Encrypt** checkbox when creating an individual partition, software RAID array, or logical volume. After you finish partitioning, you will be prompted for an encryption passphrase. This passphrase will be required to access the encrypted devices. If you have pre-existing LUKS devices and provided correct passphrases for them earlier in the install process the passphrase entry dialog will also contain a checkbox. Checking this checkbox indicates that you would like the new passphrase to be added to an available slot in each of the pre-existing encrypted block devices.



Note

Checking the **Encrypt System** checkbox on the **Automatic Partitioning** screen and then choosing **Create custom layout** does not cause any block devices to be encrypted automatically.



Note

You can use **kickstart** to set a separate passphrase for each new encrypted block device.

3.2.4.6. Links of Interest

For additional information on LUKS or encrypting hard drives under Red Hat Enterprise Linux please visit one of the following links:

- » [LUKS home page](#)
- » [LUKS/cryptsetup FAQ](#)
- » [LUKS - Linux Unified Key Setup](#)
- » [HOWTO: Creating an encrypted Physical Volume \(PV\) using a second hard drive and pvmove](#)

3.2.5. Using GNU Privacy Guard (GnuPG)

GPG is used to identify yourself and authenticate your communications, including those with people you do not know. GPG allows anyone reading a GPG-signed email to verify its authenticity. In other words, GPG allows someone to be reasonably certain that communications signed by you actually are from you. GPG is useful because it helps prevent third parties from altering code or intercepting conversations and altering the message.

3.2.5.1. Creating GPG Keys in GNOME

Install the **Seahorse** utility, which makes GPG key management easier:

```
~]# yum install seahorse
```

To create a key, from the **Applications** → **Accessories** menu select **Passwords and Encryption Keys**, which starts the application **Seahorse**. From the **File** menu select **New** and then **PGP Key**. Then click **Continue**. Type your full name, email address, and an optional comment describing who are you (for example: John C. Smith, jsmith@example.com, Software Engineer). Click **Create**. A dialog is displayed asking for a passphrase for the key. Choose a strong passphrase but also easy to remember. Click **OK** and the key is created.



Warning

If you forget your passphrase, the key cannot be used and any data encrypted using that key will be lost.

To find your GPG key ID, look in the **Key ID** column next to the newly created key. In most cases, if you are asked for the key ID, you should prepend **0x** to the key ID, as in **0x6789ABCD**. You should make a backup of your private key and store it somewhere secure.

3.2.5.2. Creating GPG Keys in KDE

Start the KGpg program from the main menu by selecting **Applications** → **Utilities** → **Encryption Tool**. If you have never used KGpg before, the program walks you through the process of creating your own GPG keypair. A dialog box appears prompting you to create a new key pair. Enter your name, email address, and an optional comment. You can also choose an expiration time for your key, as well as the key strength (number of bits) and algorithms. The next dialog box prompts you for your passphrase. At this point, your key appears in the main **KGpg** window.



Warning

If you forget your passphrase, the key cannot be used and any data encrypted using that key will be lost.

To find your GPG key ID, look in the **Key ID** column next to the newly created key. In most cases, if you are asked for the key ID, you should prepend **0x** to the key ID, as in **0x6789ABCD**. You should make a backup of your private key and store it somewhere secure.

3.2.5.3. Creating GPG Keys Using the Command Line

Use the following shell command: **gpg2 --gen-key**

This command generates a key pair that consists of a public and a private key. Other people use your public key to authenticate and/or decrypt your communications. Distribute your public key as widely as

possible, especially to people who you know will want to receive authentic communications from you, such as a mailing list.

A series of prompts directs you through the process. Press the **Enter** key to assign a default value if desired. The first prompt asks you to select what kind of key you prefer:

```
Please select what kind of key you want:
(1) RSA and RSA (default)
(2) DSA and Elgamal
(3) DSA (sign only)
(4) RSA (sign only)
Your selection?
```

In almost all cases, the default is the correct choice. A RSA/RSA key allows you not only to sign communications, but also to encrypt files.

Next, choose the key size:

```
RSA keys may be between 1024 and 4096 bits long.
What keysize do you want? (2048)
```

Again, the default, 2048, is sufficient for almost all users, and represents an extremely strong level of security.

Next, choose when the key will expire. It is a good idea to choose an expiration date instead of using the default, which is **none**. If, for example, the email address on the key becomes invalid, an expiration date will remind others to stop using that public key.

Please specify how long the key should be valid. 0 = key does not expire d = key expires in n days w = key expires in n weeks m = key expires in n months y = key expires in n years Key is valid for? (0)

Entering a value of **1y**, for example, makes the key valid for one year. (You may change this expiration date after the key is generated, if you change your mind.)

Before the **gpg2** application asks for signature information, the following prompt appears:

```
Is this correct (y/n)?
```

Enter **y** to finish the process.

Next, enter your name and email address. Remember this process is about authenticating you as a real individual. For this reason, include your real name. Do not use aliases or handles, since these disguise or obfuscate your identity.

Enter your real email address for your GPG key. If you choose a bogus email address, it will be more difficult for others to find your public key. This makes authenticating your communications difficult. If you are using this GPG key for self-introduction on a mailing list, for example, enter the email address you use on that list.

Use the comment field to include aliases or other information. (Some people use different keys for different purposes and identify each key with a comment, such as "Office" or "Open Source Projects.")

At the confirmation prompt, enter the letter **0** to continue if all entries are correct, or use the other options to fix any problems. Finally, enter a passphrase for your secret key. The **gpg2** program asks you to enter your passphrase twice to ensure you made no typing errors.

Finally, **gpg2** generates random data to make your key as unique as possible. Move your mouse, type random keys, or perform other tasks on the system during this step to speed up the process. Once this step is finished, your keys are complete and ready to use:

```
pub 1024D/1B2AFA1C 2005-03-31 John Q. Doe <jqdoe@example.com>
Key fingerprint = 117C FE83 22EA B843 3E86 6486 4320 545E 1B2A FA1C
sub 1024g/CEA4B22E 2005-03-31 [expires: 2006-03-31]
```

The key fingerprint is a shorthand "signature" for your key. It allows you to confirm to others that they have received your actual public key without any tampering. You do not need to write this fingerprint down. To display the fingerprint at any time, use this command, substituting your email address:

```
~]$ gpg2 --fingerprint jqdoe@example.com
```

Your "GPG key ID" consists of 8 hex digits identifying the public key. In the example above, the GPG key ID is **1B2AFA1C**. In most cases, if you are asked for the key ID, you should prepend **0x** to the key ID, as in **0x6789ABCD**.



Warning

If you forget your passphrase, the key cannot be used and any data encrypted using that key will be lost.

3.2.5.4. About Public Key Encryption

1. [Wikipedia - Public Key Cryptography](#)
2. [HowStuffWorks - Encryption](#)

Chapter 4. General Principles of Information Security

The following general principals provide an overview of good security practices:

- ▶ Encrypt all data transmitted over networks to help prevent man-in-the-middle attacks and eavesdropping. It is important to encrypt authentication information, such as passwords.
- ▶ Minimize the amount of software installed and running services.
- ▶ Use security-enhancing software and tools, for example, Security-Enhanced Linux (SELinux) for Mandatory Access Control (MAC), Netfilter iptables for packet filtering (firewall), and the GNU Privacy Guard (GnuPG) for encrypting files.
- ▶ If possible, run each network service on a separate system to minimize the risk of one compromised service being used to compromise other services.
- ▶ Maintain user accounts: create and enforce a strong password policy; delete unused user accounts.
- ▶ Routinely review system and application logs. By default, security-relevant system logs are written to **/var/log/secure** and **/var/log/audit/audit.log**. Note: sending logs to a dedicated log server helps prevent attackers from easily modifying local logs to avoid detection.
- ▶ Never log in as the root user unless absolutely necessary. It is recommended that administrators use **sudo** to execute commands as root when required. Users capable of running **sudo** are specified in **/etc/sudoers**. Use the **visudo** utility to edit **/etc/sudoers**.

4.1. Tips, Guides, and Tools

The United States' [National Security Agency \(NSA\)](#) provides hardening guides and tips for many different operating systems, to help government agencies, businesses, and individuals secure their systems against attack. The following guides (in PDF format) provide guidance for Red Hat Enterprise Linux 6:

- ▶ [Hardening Tips for the Red Hat Enterprise Linux 5](#)
- ▶ [Guide to the Secure Configuration of Red Hat Enterprise Linux 5](#)



Note

References to Red Hat Enterprise Linux 5 hardening guides are provided in this document until hardening guides for Red Hat Enterprise Linux 6 can be made available. In the meantime, please note that the hardening guides for Red Hat Enterprise Linux 5 may not apply completely to Red Hat Enterprise Linux 6.

The [Defense Information Systems Agency \(DISA\)](#) provides documentation, checklists, and tests to help secure your system ([Information Assurance Support Environment](#)).

Chapter 5. Secure Installation

Security begins with the first time you put that CD or DVD into your disk drive to install Red Hat Enterprise Linux. Configuring your system securely from the beginning makes it easier to implement additional security settings later.

5.1. Disk Partitions

The NSA recommends creating separate partitions for `/boot`, `/`, `/home`, `/tmp`, and `/var/tmp`. The reasons for each are different and we will address each partition.

`/boot` - This partition is the first partition that is read by the system during boot up. The boot loader and kernel images that are used to boot your system into Red Hat Enterprise Linux are stored in this partition. This partition should not be encrypted. If this partition is included in `/` and that partition is encrypted or otherwise becomes unavailable then your system will not be able to boot.

`/home` - When user data (`/home`) is stored in `/` instead of in a separate partition, the partition can fill up causing the operating system to become unstable. Also, when upgrading your system to the next version of Red Hat Enterprise Linux it is a lot easier when you can keep your data in the `/home` partition as it will not be overwritten during installation. If the root partition (`/`) becomes corrupt your data could be lost forever. By using a separate partition there is slightly more protection against data loss. You can also target this partition for frequent backups.

`/tmp` and `/var/tmp` - Both the `/tmp` and the `/var/tmp` directories are used to store data that doesn't need to be stored for a long period of time. However if a lot of data floods one of these directories it can consume all of your storage space. If this happens and these directories are stored within `/` then your system could become unstable and crash. For this reason, moving these directories into their own partitions is a good idea.

5.2. Utilize LUKS Partition Encryption

During the installation process an option to encrypt your partitions will be presented to the user. The user must supply a passphrase that will be the key to unlock the bulk encryption key that will be used to secure the partition's data.

Chapter 6. Software Maintenance

Software maintenance is extremely important to maintaining a secure system. It is vital to patch software as soon as it becomes available in order to prevent attackers from using known holes to infiltrate your system.

6.1. Install Minimal Software

It is best practice to install only the packages you will use because each piece of software on your computer could possibly contain a vulnerability. If you are installing from the DVD media take the opportunity to select exactly what packages you want to install during the installation. When you find you need another package, you can always add it to the system later.

For more information on minimal installation, refer to *Section 9.17.*, “*Package Group Selection*” of the *Red Hat Enterprise Linux 6 Installation Guide*. A minimal installation can also be performed via a kickstart file using the **-nobase** option. For more information, refer to *Section 32.5.*, “*Package Selection*” of the *Red Hat Enterprise Linux 6 Installation Guide*.

6.2. Plan and Configure Security Updates

All software contains bugs. Often, these bugs can result in a vulnerability that can expose your system to malicious users. Unpatched systems are a common cause of computer intrusions. You should have a plan to install security patches in a timely manner to close those vulnerabilities so they can not be exploited.

For home users, security updates should be installed as soon as possible. Configuring automatic installation of security updates is one way to avoid having to remember, but does carry a slight risk that something can cause a conflict with your configuration or with other software on the system.

For business or advanced home users, security updates should be tested and scheduled for installation. Additional controls will need to be used to protect the system during the time between the patch release and its installation on the system. These controls would depend on the exact vulnerability, but could include additional firewall rules, the use of external firewalls, or changes in software settings.

6.3. Adjusting Automatic Updates

Red Hat Enterprise Linux is configured to apply all updates on a daily schedule. If you want to change how your system installs updates, you must do so via **Software Update Preferences**. You can change the schedule, the type of updates to apply, or to notify you of available updates.

In GNOME, you can find controls for your updates at: **System** → **Preferences** → **Software Updates**. In KDE, it is located at: **Applications** → **Settings** → **Software Updates**.

6.4. Install Signed Packages from Well Known Repositories

Software packages are published through repositories. All well known repositories support package signing. Package signing uses public key technology to prove that the package that was published by the repository has not been changed since the signature was applied. This provides some protection against installing software that may have been maliciously altered after the package was created but before you downloaded it.

Using too many repositories, untrustworthy repositories, or repositories with unsigned packages has a higher risk of introducing malicious or vulnerable code into your system. Use caution when adding repositories to yum/software update.

Chapter 7. Federal Standards and Regulations

7.1. Introduction

In order to maintain security levels, it is possible for your organization to make efforts to comply with federal and industry security specifications, standards and regulations. This chapter describes some of these standards and regulations.

7.2. Federal Information Processing Standard (FIPS)

The Federal Information Processing Standard (FIPS) Publication 140-2, is a computer security standard, developed by a U.S. Government and industry working group to validate the quality of cryptographic modules. FIPS publications (including 140-2) can be found at the following URL:

<http://csrc.nist.gov/publications/PubsFIPS.html>. Note that at the time of writing, Publication 140-3 is at Draft status, and may not represent the completed standard. The FIPS standard provides four (4) security *levels*, to ensure adequate coverage of different industries, implementations of cryptographic modules and organizational sizes and requirements. These levels are described below:

- ▶ Level 1 - Security Level 1 provides the lowest level of security. Basic security requirements are specified for a cryptographic module (e.g., at least one Approved algorithm or Approved security function shall be used). No specific physical security mechanisms are required in a Security Level 1 cryptographic module beyond the basic requirement for production-grade components. An example of a Security Level 1 cryptographic module is a personal computer (PC) encryption board.
- ▶ Level 2 - Security Level 2 enhances the physical security mechanisms of a Security Level 1 cryptographic module by adding the requirement for tamper-evidence, which includes the use of tamper-evident coatings or seals or for pick-resistant locks on removable covers or doors of the module. Tamper-evident coatings or seals are placed on a cryptographic module so that the coating or seal must be broken to attain physical access to the plaintext cryptographic keys and critical security parameters (CSPs) within the module. Tamper-evident seals or pick-resistant locks are placed on covers or doors to protect against unauthorized physical access.
- ▶ Level 3 - In addition to the tamper-evident physical security mechanisms required at Security Level 2, Security Level 3 attempts to prevent the intruder from gaining access to CSPs held within the cryptographic module. Physical security mechanisms required at Security Level 3 are intended to have a high probability of detecting and responding to attempts at physical access, use or modification of the cryptographic module. The physical security mechanisms may include the use of strong enclosures and tamper detection/response circuitry that zeroes all plaintext CSPs when the removable covers/doors of the cryptographic module are opened.
- ▶ Level 4 - Security Level 4 provides the highest level of security defined in this standard. At this security level, the physical security mechanisms provide a complete envelope of protection around the cryptographic module with the intent of detecting and responding to all unauthorized attempts at physical access. Penetration of the cryptographic module enclosure from any direction has a very high probability of being detected, resulting in the immediate zeroization of all plaintext CSPs. Security Level 4 cryptographic modules are useful for operation in physically unprotected environments.

Refer to the full FIPS 140-2 standard at <http://csrc.nist.gov/publications/fips/fips140-2/fips1402.pdf> for further details on these levels and the other specifications of the FIPS standard.

7.2.1. Enabling FIPS Mode

To make Red Hat Enterprise Linux 6 compliant with the Federal Information Processing Standard (FIPS) Publication 140-2 you need to make several changes to ensure that accredited cryptographic modules are used. To turn your system (kernel and user space) into FIPS mode, follow these steps:

1. For proper operation of the in-module integrity verification, the prelink has to be disabled. This can be done by setting configuring **PRELINKING=no** in the **/etc/sysconfig/prelink** configuration file. Existing prelinking, if any, should be undone on all system files using the **prelink -u -a** command.
2. Next, install the *dracut-fips* package:

```
~]# yum install dracut-fips
```

3. Recreate the **initramfs** file:

```
~]# dracut -f
```



Warning

This operation will overwrite the existing **initramfs** file.

4. Modify the kernel command line of the current kernel in the **/boot/grub/grub.conf** file by adding the following option:

```
fips=1
```



Note

If **/boot** or **/boot/efi** reside on separate partitions, the kernel parameter **boot=<partition of /boot or /boot/efi>** must be added to the kernel command line. Partitions can be identified with the **df /boot** or **df /boot/efi** command respectively. For example:

```
~]$ df /boot
Filesystem            1K-blocks    Used Available Use% Mounted on
/dev/sda1              495844      53780    416464   12% /boot
```

In the example above, the **/boot** partition is located on **/dev/sda1**. Therefore, the following string needs to be appended to the kernel command line:

```
boot=/dev/sda1
```

5. Reboot your system.

Should you require strict FIPS compliance, the **fips=1** kernel option needs to be added to the kernel command line during system installation so that key generation is done with FIPS approved algorithms and continuous monitoring tests in place. Users should also ensure that the system has plenty of entropy during the installation process by moving the mouse around, or if no mouse is available, ensuring that many keystrokes are typed. The recommended amount of keystrokes is 256 and more. Less than 256 keystrokes may generate a non-unique key.

7.3. National Industrial Security Program Operating Manual (NISPOM)

The NISPOM (also called DoD 5220.22-M), as a component of the National Industrial Security Program (NISP), establishes a series of procedures and requirements for all government contractors with regard to classified information. The current NISPOM is dated February 28, 2006. The NISPOM document can be downloaded from the following URL:

https://www.dss.mil/GW/ShowBinary/DSS/isp/fac_clear/download_nispom.html.

7.4. Payment Card Industry Data Security Standard (PCI DSS)

From <https://www.pcisecuritystandards.org/about/index.shtml>: *The PCI Security Standards Council is an open global forum, launched in 2006, that is responsible for the development, management, education, and awareness of the PCI Security Standards, including the Data Security Standard (DSS).*

You can download the PCI DSS standard from

https://www.pcisecuritystandards.org/security_standards/pci_dss.shtml.

7.5. Security Technical Implementation Guide

A Security Technical Implementation Guide or STIG is a methodology for standardized secure installation and maintenance of computer software and hardware.

Refer to the following URL for more information on STIG: <http://iase.disa.mil/stigs/index.html> .

Chapter 8. References

The following references are pointers to additional information that is relevant to SELinux and Red Hat Enterprise Linux but beyond the scope of this guide. Note that due to the rapid development of SELinux, some of this material may only apply to specific releases of Red Hat Enterprise Linux.

Books

SELinux by Example

Mayer, MacMillan, and Caplan

Prentice Hall, 2007

Tutorials and Help

Understanding and Customizing the Apache HTTP SELinux Policy

<http://docs.fedoraproject.org/selinux-apache-fc3/>

Tutorials and talks from Russell Coker

<http://www.coker.com.au/selinux/talks/ibmtu-2004/>

Generic Writing SELinux policy HOWTO

<http://www.lurking-grue.org/writingselinuxpolicyHOWTO.html>

Red Hat Knowledgebase

<http://kbase.redhat.com/>

General Information

NSA SELinux main website

<http://www.nsa.gov/selinux/>

NSA SELinux FAQ

<http://www.nsa.gov/selinux/info/faq.cfm>

Fedora SELinux FAQ

<http://docs.fedoraproject.org/selinux-faq/>

SELinux NSA's Open Source Security Enhanced Linux

<http://www.oreilly.com/catalog/selinux/>

Technology

An Overview of Object Classes and Permissions

http://www.tresys.com/selinux/obj_perms_help.html

Integrating Flexible Support for Security Policies into the Linux Operating System (a history of Flask implementation in Linux)

<http://www.nsa.gov/research/files/selinux/papers/selsymp2005.pdf>

Implementing SELinux as a Linux Security Module

http://www.nsa.gov/research/files/publications/implementing_selinux.pdf

A Security Policy Configuration for the Security-Enhanced Linux

<http://www.nsa.gov/research/files/selinux/papers/policy/policy.shtml>

Community

Fedora SELinux User Guide

http://docs.fedoraproject.org/en-US/Fedora/13/html/Security-Enhanced_Linux/

Fedora SELinux Managing Confined Services Guide

http://docs.fedoraproject.org/en-US/Fedora/13/html/Managing_Confined_Services/

SELinux community page

<http://selinuxproject.org/>

IRC

irc.freenode.net, #selinux, #fedora-selinux, #security

History

Quick history of Flask

<http://www.cs.utah.edu/flux/fluke/html/flask.html>

Full background on Fluke

<http://www.cs.utah.edu/flux/fluke/html/index.html>

Encryption Standards

A.1. Synchronous Encryption

A.1.1. Advanced Encryption Standard - AES

In cryptography, the Advanced Encryption Standard (AES) is an encryption standard adopted by the U.S. government. The standard comprises three block ciphers, AES-128, AES-192 and AES-256, adopted from a larger collection originally published as Rijndael. Each AES cipher has a 128-bit block size, with key sizes of 128, 192 and 256 bits, respectively. The AES ciphers have been analyzed extensively and are now used worldwide, as was the case with its predecessor, the Data Encryption Standard (DES). ^[13]

A.1.1.1. AES History

AES was announced by National Institute of Standards and Technology (NIST) as U.S. FIPS PUB 197 (FIPS 197) on November 26, 2001 after a 5-year standardization process in which fifteen competing designs were presented and evaluated before Rijndael was selected as the most suitable (see Advanced Encryption Standard process for more details). It became effective as a standard May 26, 2002. It is available in many different encryption packages. AES is the first publicly accessible and open cipher approved by the NSA for top secret information (see Security of AES, below). ^[14]

The Rijndael cipher was developed by two Belgian cryptographers, Joan Daemen and Vincent Rijmen, and submitted by them to the AES selection process. Rijndael (pronounced [reinda:l]) is a portmanteau of the names of the two inventors. ^[15]

A.1.2. Data Encryption Standard - DES

The Data Encryption Standard (DES) is a block cipher (a form of shared secret encryption) that was selected by the National Bureau of Standards as an official Federal Information Processing Standard (FIPS) for the United States in 1976 and which has subsequently enjoyed widespread use internationally. It is based on a symmetric-key algorithm that uses a 56-bit key. The algorithm was initially controversial with classified design elements, a relatively short key length, and suspicions about a National Security Agency (NSA) backdoor. DES consequently came under intense academic scrutiny which motivated the modern understanding of block ciphers and their cryptanalysis. ^[16]

A.1.2.1. DES History

DES is now considered to be insecure for many applications. This is chiefly due to the 56-bit key size being too small; in January, 1999, distributed.net and the Electronic Frontier Foundation collaborated to publicly break a DES key in 22 hours and 15 minutes (see chronology). There are also some analytical results which demonstrate theoretical weaknesses in the cipher, although they are unfeasible to mount in practice. The algorithm is believed to be practically secure in the form of Triple DES, although there are theoretical attacks. In recent years, the cipher has been superseded by the Advanced Encryption Standard (AES). ^[17]

In some documentation, a distinction is made between DES as a standard and DES the algorithm which is referred to as the DEA (the Data Encryption Algorithm). When spoken, "DES" is either spelled out as an abbreviation (/ˌdiːˌjɪːˈes/), or pronounced as a one-syllable acronym (/ˈdez/). ^[18]

A.2. Public-key Encryption

Public-key cryptography is a cryptographic approach, employed by many cryptographic algorithms and cryptosystems, whose distinguishing characteristic is the use of asymmetric key algorithms instead of or in addition to symmetric key algorithms. Using the techniques of public key-private key cryptography,

many methods of protecting communications or authenticating messages formerly unknown have become practical. They do not require a secure initial exchange of one or more secret keys as is required when using symmetric key algorithms. It can also be used to create digital signatures. ^[19]

Public key cryptography is a fundamental and widely used technology around the world, and is the approach which underlies such Internet standards as Transport Layer Security (TLS) (successor to SSL), PGP and GPG. ^[20]

The distinguishing technique used in public key cryptography is the use of asymmetric key algorithms, where the key used to encrypt a message is not the same as the key used to decrypt it. Each user has a pair of cryptographic keys — a public key and a private key. The private key is kept secret, whilst the public key may be widely distributed. Messages are encrypted with the recipient's public key and can only be decrypted with the corresponding private key. The keys are related mathematically, but the private key cannot be feasibly (ie, in actual or projected practice) derived from the public key. It was the discovery of such algorithms which revolutionized the practice of cryptography beginning in the middle 1970s. ^[21]

In contrast, Symmetric-key algorithms, variations of which have been used for some thousands of years, use a single secret key shared by sender and receiver (which must also be kept private, thus accounting for the ambiguity of the common terminology) for both encryption and decryption. To use a symmetric encryption scheme, the sender and receiver must securely share a key in advance. ^[22]

Because symmetric key algorithms are nearly always much less computationally intensive, it is common to exchange a key using a key-exchange algorithm and transmit data using that key and a symmetric key algorithm. PGP, and the SSL/TLS family of schemes do this, for instance, and are called hybrid cryptosystems in consequence. ^[23]

A.2.1. Diffie-Hellman

Diffie–Hellman key exchange (D–H) is a cryptographic protocol that allows two parties that have no prior knowledge of each other to jointly establish a shared secret key over an insecure communications channel. This key can then be used to encrypt subsequent communications using a symmetric key cipher. ^[24]

A.2.1.1. Diffie-Hellman History

The scheme was first published by Whitfield Diffie and Martin Hellman in 1976, although it later emerged that it had been separately invented a few years earlier within GCHQ, the British signals intelligence agency, by Malcolm J. Williamson but was kept classified. In 2002, Hellman suggested the algorithm be called Diffie–Hellman–Merkle key exchange in recognition of Ralph Merkle's contribution to the invention of public-key cryptography (Hellman, 2002). ^[25]

Although Diffie–Hellman key agreement itself is an anonymous (non-authenticated) key-agreement protocol, it provides the basis for a variety of authenticated protocols, and is used to provide perfect forward secrecy in Transport Layer Security's ephemeral modes (referred to as EDH or DHE depending on the cipher suite). ^[26]

U.S. Patent 4,200,770, now expired, describes the algorithm and credits Hellman, Diffie, and Merkle as inventors. ^[27]

A.2.2. RSA

In cryptography, RSA (which stands for Rivest, Shamir and Adleman who first publicly described it; see below) is an algorithm for public-key cryptography. It is the first algorithm known to be suitable for signing as well as encryption, and was one of the first great advances in public key cryptography. RSA is widely

used in electronic commerce protocols, and is believed to be secure given sufficiently long keys and the use of up-to-date implementations.

A.2.3. DSA

DSA (Digital Signature Algorithm) is a standard for digital signatures, a United States federal government standard for digital signatures. DSA is for signatures only and is not an encryption algorithm. [28]

A.2.4. SSL/TLS

Transport Layer Security (TLS) and its predecessor, Secure Sockets Layer (SSL), are cryptographic protocols that provide security for communications over networks such as the Internet. TLS and SSL encrypt the segments of network connections at the Transport Layer end-to-end.

Several versions of the protocols are in widespread use in applications like web browsing, electronic mail, Internet faxing, instant messaging and voice-over-IP (VoIP). [29]

A.2.5. Cramer-Shoup Cryptosystem

The Cramer–Shoup system is an asymmetric key encryption algorithm, and was the first efficient scheme proven to be secure against adaptive chosen ciphertext attack using standard cryptographic assumptions. Its security is based on the computational intractability (widely assumed, but not proved) of the decisional Diffie–Hellman assumption. Developed by Ronald Cramer and Victor Shoup in 1998, it is an extension of the ElGamal cryptosystem. In contrast to ElGamal, which is extremely malleable, Cramer–Shoup adds additional elements to ensure non-malleability even against a resourceful attacker. This non-malleability is achieved through the use of a collision-resistant hash function and additional computations, resulting in a ciphertext which is twice as large as in ElGamal. [30]

A.2.6. ElGamal Encryption

In cryptography, the ElGamal encryption system is an asymmetric key encryption algorithm for public-key cryptography which is based on the Diffie-Hellman key agreement. It was described by Taher ElGamal in 1985. ElGamal encryption is used in the free GNU Privacy Guard software, recent versions of PGP, and other cryptosystems. [31]

[13]"Advanced Encryption Standard." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Advanced_Encryption_Standard

[14]"Advanced Encryption Standard." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Advanced_Encryption_Standard

[15]"Advanced Encryption Standard." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Advanced_Encryption_Standard

[16]"Data Encryption Standard." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Data_Encryption_Standard

[17]"Data Encryption Standard." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Data_Encryption_Standard

[18]"Data Encryption Standard." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Data_Encryption_Standard

[19]"Public-key Encryption." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Public-key_cryptography

[20]"Public-key Encryption." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Public-key_cryptography

[21]"Public-key Encryption." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Public-key_cryptography

[22]"Public-key Encryption." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Public-key_cryptography

[23]"Public-key Encryption." *Wikipedia*. 14 November 2009 http://en.wikipedia.org/wiki/Public-key_cryptography

[24]"Diffie-Hellman." *Wikipedia*. 14 November 2009 <http://en.wikipedia.org/wiki/Diffie-Hellman>

[25]"Diffie-Hellman." *Wikipedia*. 14 November 2009 <http://en.wikipedia.org/wiki/Diffie-Hellman>

[26]"Diffie-Hellman." *Wikipedia*. 14 November 2009 <http://en.wikipedia.org/wiki/Diffie-Hellman>

[27]"Diffie-Hellman." *Wikipedia*. 14 November 2009 <http://en.wikipedia.org/wiki/Diffie-Hellman>

[28]"DSA." *Wikipedia*. 24 February 2010 http://en.wikipedia.org/wiki/Digital_Signature_Algorithm

[29] "TLS/SSL." *Wikipedia*. 24 February 2010 http://en.wikipedia.org/wiki/Transport_Layer_Security

[30] "Cramer-Shoup cryptosystem." *Wikipedia*. 24 February 2010 http://en.wikipedia.org/wiki/Cramer-Shoup_cryptosystem

[31] "ElGamal encryption" *Wikipedia*. 24 February 2010 http://en.wikipedia.org/wiki/ElGamal_encryption

Revision History

Revision 1-9	Feb 21 2013	Martin Prpič
Release of the Security Guide for Red Hat Enterprise Linux 6.4		
Revision 1-8.25	Jun 20 2012	Martin Prpič
Release of the Security Guide for Red Hat Enterprise Linux 6.3		
Revision 1-7	Dec 6 2011	Martin Prpič
Release of the Security Guide for Red Hat Enterprise Linux 6.2		
Revision 1-6	Aug 05 2011	Paul Kennedy
Resolves BZ#702249, broken link.		
Revision 1-5	Apr 19 2010	Scott Radvan
Minor fixes, final build for Beta		
Revision 1-4	Mar 5 2010	Scott Radvan
QE Review and Updates		
Revision 1-3	Feb 19 2010	Scott Radvan
Push to testing area ready for review.		