

Red Hat Enterprise Linux 4

Introduzione al System Administration



Red Hat Enterprise Linux 4: Introduzione al System Administration

Copyright © 2005 Red Hat, Inc.



Red Hat, Inc.

1801 Varsity Drive Raleigh NC 27606-2072 USA Telefono: +1 919 754 3700 Telefono: 888 733 4281 Fax: +1 919 754 3701 PO Box 13588 Research Triangle Park NC 27709 Stati Uniti

rhel-isa(1T)-4-Print-RHI (2004-08-25T17:11)

Copyright © 2005 by Red Hat, Inc. Questo materiale può essere distribuito solo secondo i termini e le condizioni della Open Publication License, V1.0 o successiva (l'ultima versione è disponibile all'indirizzo <http://www.opencontent.org/openpub/>).

La distribuzione di versioni modificate di questo documento è proibita senza esplicita autorizzazione del detentore del copyright.

La distribuzione per scopi commerciali del libro o di una parte di esso sotto forma di opera stampata è proibita se non autorizzata dal detentore del copyright.

Red Hat ed il logo "Shadow Man" di Red Hat, sono dei marchi registrati di Red Hat, Inc. negli Stati Uniti e nelle altre Nazioni.

Tutti gli altri marchi elencati sono di proprietà dei rispettivi proprietari.

Il codice GPG della chiave security@redhat.com è:

CA 20 86 86 2B D6 9D FC 65 F6 EC C4 21 91 80 CD DB 42 A6 0E

Sommario

Introduzione	i
1. Informazioni specifiche sull'architettura	i
2. Convenzioni del documento.....	i
3. Attivate la vostra sottoscrizione	iv
3.1. Fornite un login di Red Hat	v
3.2. Fornite il numero della vostra sottoscrizione.....	v
3.3. Collegate il vostro sistema	v
4. Ed altro ancora.....	v
4.1. Inviateci suggerimenti.....	v
1. La filosofia di un amministratore di sistema	1
1.1. Automatizzare ogni cosa.....	1
1.2. Documentare ogni cosa.....	2
1.3. Comunicare il più possibile	3
1.3.1. Comunicate ai vostri utenti cosa farete.....	3
1.3.2. Comunicate ai vostri utenti cosa state facendo	4
1.3.3. Comunicate ai vostri utenti cosa avete fatto	5
1.4. Conoscere le vostre risorse	5
1.5. Conoscere i vostri utenti	6
1.6. Conoscere la vostra compagnia	6
1.7. La sicurezza non può essere considerata come un fattore secondario	6
1.7.1. I rischi dell'ingegneria sociale	7
1.8. Pianificare a priori.....	8
1.9. Aspettatevi un imprevisto	8
1.10. Informazioni specifiche su Red Hat Enterprise Linux	9
1.10.1. Automazione	9
1.10.2. Documentazione e comunicazione	9
1.10.3. Sicurezza	10
1.11. Risorse aggiuntive.....	11
1.11.1. Documentazione installata	11
1.11.2. Siti web utili.....	11
1.11.3. Libri correlati	12
2. Controllo delle risorse.....	13
2.1. Concetti di base.....	13
2.2. Controllo della prestazione del sistema	13
2.3. Controllo della capacità del sistema	14
2.4. Cosa controllare?	14
2.4.1. Controllo della potenza della CPU	15
2.4.2. Controllo della larghezza della banda.....	16
2.4.3. Controllo della memoria	17
2.4.4. Controllo dello storage.....	17
2.5. Informazioni specifiche su Red Hat Enterprise Linux	18
2.5.1. free.....	18
2.5.2. top.....	19
2.5.3. vmstat.....	21
2.5.4. La suite di Sysstat per i tool di controllo delle risorse.....	22
2.5.5. OProfile.....	25
2.6. Risorse aggiuntive.....	28
2.6.1. Documentazione installata	28
2.6.2. Siti web utili.....	29
2.6.3. Libri correlati	29

3. Larghezza di banda e potenza di processazione.....	31
3.1. Larghezza di banda	31
3.1.1. Bus	31
3.1.2. Datapath	32
3.1.3. Potenziali problemi relativi alla larghezza di banda	32
3.1.4. Soluzioni potenziali relative alla larghezza di banda	33
3.1.5. In generale.	33
3.2. Potenza di processazione	34
3.2.1. Dati relativi alla potenza di processazione.....	34
3.2.2. Utenze interessate all'uso della potenza di processazione.....	34
3.2.3. Migliorare la carenza di CPU	35
3.3. Informazioni specifiche su Red Hat Enterprise Linux	38
3.3.1. Controllo della larghezza di banda su Red Hat Enterprise Linux	38
3.3.2. Controllo dell'utilizzo della CPU su Red Hat Enterprise Linux.....	40
3.4. Risorse aggiuntive.....	44
3.4.1. Documentazione installata	44
3.4.2. Siti web utili	44
3.4.3. Libri correlati	44
4. Memoria virtuale e fisica.....	47
4.1. Percorsi per l'accesso allo storage	47
4.2. Gamma dello Storage.....	47
4.2.1. Registri CPU	48
4.2.2. Memoria Cache	48
4.2.3. Memoria principale — RAM.....	49
4.2.4. Dischi fissi.....	50
4.2.5. Storage di backup off-line.....	51
4.3. Concetti di base della memoria virtuale.....	51
4.3.1. Spiegazione semplice della memoria virtuale.....	51
4.3.2. Backing Store — Il principio centrale della memoria virtuale.....	52
4.4. Memoria virtuale: Dettagli.....	53
4.4.1. Page Fault.....	53
4.4.2. Il working set	54
4.4.3. Swapping.....	54
4.5. Implicazioni sulle prestazioni della memoria virtuale	55
4.5.1. Come ottenere una prestazione pessima	55
4.5.2. Come ottenere la prestazione migliore	55
4.6. Red Hat Enterprise Linux-Informazioni specifiche	56
4.7. Risorse aggiuntive.....	58
4.7.1. Documentazione installata	59
4.7.2. Siti web utili	59
4.7.3. Libri correlati	59
5. Gestione dello storage.....	61
5.1. Panoramica dell'hardware dello storage.....	61
5.1.1. Disk Platter.....	61
5.1.2. Dispositivo di lettura/scrittura dei dati.....	61
5.1.3. Braccia d'accesso.....	62
5.2. Concetti relativi allo storage	63
5.2.1. Addressing basato sulla geometria.....	63
5.2.2. Addressing basato sul blocco.....	65
5.3. Interfacce del dispositivo mass storage.....	65
5.3.1. Background storico	65
5.3.2. Interfacce standard al giorno d'oggi	66
5.4. Caratteristiche delle prestazioni dei dischi fissi	69
5.4.1. Limitazioni elettriche/meccaniche	69
5.4.2. Prestazione e carico I/O	70

5.5. Rendere utilizzabile lo storage.....	72
5.5.1. Partizioni/Sezioni.....	72
5.5.2. File System	73
5.5.3. Struttura della directory	76
5.5.4. Abilitare l'accesso allo storage	76
5.6. Tecnologie avanzate per lo storage	77
5.6.1. Storage accessibile tramite la rete.....	77
5.6.2. Storage basato su RAID.....	78
5.6.3. Logical Volume Management	83
5.7. Gestione giornaliera dello storage	84
5.7.1. Controllo dello spazio disponibile	85
5.7.2. Problematiche riguardanti il disk quota	87
5.7.3. Problematiche relative al file.....	88
5.7.4. Aggiunta/rimozione dello storage.....	89
5.8. Una parola sui Backup. . .	95
5.9. Informazioni specifiche su Red Hat Enterprise Linux.....	95
5.9.1. Convenzioni del device naming	96
5.9.2. Informazioni di base sui file system.....	98
5.9.3. Montaggio dei file system.....	100
5.9.4. Storage accessibile dalla rete con Red Hat Enterprise Linux	103
5.9.5. Montaggio automatico del file system con <code>/etc/fstab</code>	104
5.9.6. Aggiunta/rimozione dello storage.....	105
5.9.7. Implementazione dei Disk Quota.....	109
5.9.8. Creazione dei RAID array	113
5.9.9. Gestione giornaliera dei RAID Array	114
5.9.10. Logical Volume Management	115
5.10. Risorse aggiuntive.....	115
5.10.1. Documentazione installata	116
5.10.2. Siti web utili	116
5.10.3. Libri correlati	116
6. Gestione degli User Account e dell'accesso alle risorse.....	119
6.1. Gestione degli User Account	119
6.1.1. Nome utente	119
6.1.2. Password	122
6.1.3. Informazioni per il controllo dell'accesso	126
6.1.4. Gestione degli account e dell'accesso alle risorse giorno per giorno	127
6.2. Gestione delle risorse dell'utente.....	129
6.2.1. Personale in grado di accedere ai dati condivisi	129
6.2.2. Da dove si possono accedere i dati condivisi.....	130
6.2.3. Limiti implementati per prevenire l'abuso delle risorse	131
6.3. Informazioni specifiche su Red Hat Enterprise Linux.....	131
6.3.1. User Account, Gruppi, e Permessi.....	131
6.3.2. File che controllano gli User Account ed i Gruppi.....	133
6.3.3. User Account e Group Application.....	137
6.4. Risorse aggiuntive.....	138
6.4.1. Documentazione installata	139
6.4.2. Siti web utili	139
6.4.3. Libri correlati	140

7. Stampanti e stampa.....	141
7.1. Tipi di stampante.....	141
7.1.1. Considerazioni sulla stampa	141
7.2. Stampanti ad impatto	142
7.2.1. Stampanti dot-Matrix	143
7.2.2. Stampanti Daisy-Wheel	143
7.2.3. Stampanti di linea	143
7.2.4. Materiali di consumo per la stampante ad impatto.....	143
7.3. Stampanti a getto 'Inkjet'.....	144
7.3.1. Materiali di consumo per le stampanti Inkjet	144
7.4. Stampanti laser.....	145
7.4.1. Stampanti laser a colori.....	145
7.4.2. Materiali di consumo per stampanti a colori.....	145
7.5. Altri tipi di stampante	146
7.6. Tecnologie e linguaggi di stampa	146
7.7. Networked contro le stampanti locali	147
7.8. Informazioni specifiche su Red Hat Enterprise Linux	147
7.9. Risorse aggiuntive.....	149
7.9.1. Documentazione installata	149
7.9.2. Siti web utili	149
7.9.3. Libri correlati	149
8. Pianificazione di un system disaster	151
8.1. Tipi di disaster.....	151
8.1.1. Problemi hardware	151
8.1.2. Problemi software	157
8.1.3. Problematiche riguardanti l'ambiente.....	159
8.1.4. Errori umani	166
8.2. Backup	170
8.2.1. Dati diversi: Diversi requisiti di backup	170
8.2.2. Software di backup: Acquistare o Creare	172
8.2.3. Tipi di backup	173
8.2.4. Media di backup.....	174
8.2.5. Conservazione dei backup	176
8.2.6. Problematiche riguardanti il processo di ripristino.....	176
8.3. Disaster Recovery	177
8.3.1. Creazione, prova ed implementazione di un piano per un disaster recovery..	178
8.3.2. Siti di backup: Cold, Warm, Hot.....	179
8.3.3. Disponibilità software e hardware	180
8.3.4. Disponibilità dei dati di backup	180
8.3.5. Connettività della rete al sito di backup.....	180
8.3.6. Assunzione di personale per il sito di backup.....	180
8.3.7. Ritorno alla normalità	181
8.4. Informazioni specifiche di Red Hat Enterprise Linux	181
8.4.1. Supporto software	181
8.4.2. Tecnologie di backup	182
8.5. Risorse aggiuntive.....	185
8.5.1. Documentazione	185
8.5.2. Siti web utili	185
8.5.3. Libri correlati	186
Indice.....	187
Colophon.....	195

Introduzione

Benvenuti alla *Red Hat Enterprise Linux Introduzione al System Administration*.

Red Hat Enterprise Linux Introduzione al System Administration contiene informazioni utili per i nuovi amministratori di sistema Red Hat Enterprise Linux. Queste informazioni *non* hanno la presunzione di insegnarvi come eseguire compiti particolari con Red Hat Enterprise Linux; esse forniscono semplicemente l'opportunità di acquisire quella conoscenza posseduta da molti amministratori esperti.

Questa guida presuppone una esperienza limitata come utenti di Linux, senza averne alcuna come amministratori. Se non avete alcuna conoscenza di Linux (ed in particolare di Red Hat Enterprise Linux), fareste meglio acquistare un manuale introduttivo su Linux.

Ogni capitolo in *Red Hat Enterprise Linux Introduzione al System Administration* ha la seguente struttura:

- Materiale introduttivo generico — Questa sezione affronta gli argomenti senza discutere i dettagli inerenti a specifici sistemi operativi, alla tecnologia o alla metodologia.
- Red Hat Enterprise Linux-materiale specifico — Questa sezione indirizza in generale gli aspetti inerenti a Linux e in particolare quelli inerenti a Red Hat Enterprise Linux.
- Risorse aggiuntive per studi più approfonditi — Questa sezione include riferimenti ad altri manuali di Red Hat Enterprise Linux, siti web utili, e libri che contengono informazioni utili per l'argomento in questione.

Adottando una struttura costante, i lettori possono affrontare *Red Hat Enterprise Linux Introduzione al System Administration* molto più facilmente. Per esempio, un amministratore esperto di sistema, con una conoscenza limitata di Red Hat Enterprise Linux può evitare le sole sezioni inerenti a Red Hat Enterprise Linux stesso, mentre un nuovo amministratore potrebbe iniziare leggendo solo le sezioni che riguardano la parte generica, e usare quelle relative a Red Hat Enterprise Linux come una introduzione ad ulteriori risorse.

Red Hat Enterprise Linux System Administration Guide rappresenta una risorsa eccellente nell'eseguire compiti specifici in un ambiente Red Hat Enterprise Linux. Gli amministratori che necessitano maggiori informazioni dovrebbero consultare *Red Hat Enterprise Linux Reference Guide*.

Le versioni HTML, PDF, e RPM del manuale sono disponibili sul CD di documentazione di Red Hat Enterprise Linux e online su <http://www.redhat.com/docs/>.



Nota Bene

Anche se questo manuale riporta le informazioni più aggiornate, vi consigliamo di leggere le *Release Note di Red Hat Enterprise Linux*, per informazioni che potrebbero non essere state incluse prima della finalizzazione di questa documentazione. Tali informazioni possono essere trovate sul CD #1 di Red Hat Enterprise Linux e online su <http://www.redhat.com/docs/>.

1. Informazioni specifiche sull'architettura

Se non diversamente riportato, tutte le informazioni contenute in questo manuale valgono solo per i processori x86 e per i processori che contengono le tecnologie Intel® Extended Memory 64 Technology (Intel® EM64T) e AMD64. Per informazioni specifiche all'architettura, fate riferimento a *Red Hat Enterprise Linux Installation Guide*.

2. Convenzioni del documento

Consultando il presente manuale, vedrete alcune parole stampate con caratteri, dimensioni e stili differenti. Si tratta di un metodo sistematico per mettere in evidenza determinate parole; lo stesso stile grafico indica l'appartenenza a una specifica categoria. Le parole rappresentate in questo modo includono:

comando

I comandi di Linux (e altri comandi di sistemi operativi, quando usati) vengono rappresentati in questo modo. Questo stile indica che potete digitare la parola o la frase nella linea di comando e premere [Invio] per invocare il comando. A volte un comando contiene parole che vengono rappresentate con uno stile diverso (come i file name). In questi casi, tali parole vengono considerate come parte integrante del comando e, dunque, l'intera frase viene visualizzata come un comando. Per esempio:

Utilizzate il comando `cat testfile` per visualizzare il contenuto di un file chiamato `testfile`, nella directory corrente.

file name

I file name, i nomi delle directory, i percorsi e i nomi del pacchetto RPM vengono rappresentati con questo stile grafico. Ciò significa che un file o una directory particolari, sono rappresentati sul vostro sistema da questo nome. Per esempio:

Il file `.bashrc` nella vostra home directory contiene le definizioni e gli alias della shell bash per uso personale.

Il file `/etc/fstab` contiene le informazioni relative ai diversi dispositivi del sistema e file system.

Installate il pacchetto RPM `webalizer` per utilizzare un programma di analisi per il file di log del server Web.

applicazione

Questo stile grafico indica che il programma citato è un'applicazione per l'utente finale "end user" (contrariamente al software del sistema). Per esempio:

Utilizzate **Mozilla** per navigare sul Web.

[tasto]

I pulsanti della tastiera sono rappresentati in questo modo. Per esempio:

Per utilizzare la funzionalità [Tab], inserite una lettera e poi premete il tasto [Tab]. Il vostro terminale mostra l'elenco dei file che iniziano con quella lettera.

[tasto]-[combinazione]

Una combinazione di tasti viene rappresentata in questo modo. Per esempio:

La combinazione dei tasti [Ctrl]-[Alt]-[Backspace] esce dalla vostra sessione grafica e vi riporta alla schermata grafica di login o nella console.

testo presente in un'interfaccia grafica

Un titolo, una parola o una frase trovata su di una schermata dell'interfaccia GUI o una finestra, verrà mostrata con questo stile: Il testo mostrato in questo stile, viene usato per identificare una particolare schermata GUI o un elemento della schermata GUI, (per esempio il testo associato a una casella o a un campo). Esempio:

Selezionate la casella di controllo, **Richiedi la password**, se desiderate che lo screen saver richieda una password prima di scomparire.

livello superiore di un menu o di una finestra dell'interfaccia grafica

Quando vedete una parola scritta con questo stile grafico, si tratta di una parola posta per prima in un menu a tendina. Facendo clic sulla parola nella schermata GUI, dovrebbe comparire il resto del menu. Per esempio:

In corrispondenza di **File** in un terminale di GNOME, è presente l'opzione **Nuova tabella** che vi consente di aprire più prompt della shell nella stessa finestra.

Se dovete digitare una sequenza di comandi da un menu GUI, essi verranno visualizzati con uno stile simile al seguente esempio:

Per avviare l'editor di testo **Emacs**, fate clic sul **pulsante del menu principale** (sul pannello) => **Applicazioni => Emacs**.

pulsante di una schermata o una finestra dell'interfaccia grafica

Questo stile indica che il testo si trova su di un pulsante in una schermata GUI. Per esempio:

Fate clic sul pulsante **Indietro** per tornare all'ultima pagina Web visualizzata.

output del computer

Il testo in questo stile, indica il testo visualizzato ad un prompt della shell come ad esempio, messaggi di errore e risposte ai comandi. Per esempio:

Il comando `ls` visualizza i contenuti di una directory. Per esempio:

Desktop	about.html	logs	paulwesterberg.png
Mail	backupfiles	mail	reports

L'output restituito dal computer in risposta al comando (in questo caso, il contenuto della directory) viene mostrato con questo stile grafico.

prompt

Un prompt, ovvero uno dei modi utilizzati dal computer per indicare che è pronto per ricevere un vostro input. Ecco qualche esempio:

```
$
#
[stephen@maturin stephen]$
leopard login:
```

input dell'utente

Il testo che l'utente deve digitare sulla linea di comando o in un'area di testo di una schermata di un'interfaccia grafica, è visualizzato con questo stile come nell'esempio riportato:

Per avviare il sistema in modalità di testo, dovete digitare il comando **text** al prompt `boot:.`

replaceable

Il testo usato per gli esempi, il quale deve essere sostituito con i dati forniti dall'utente, è visualizzato con questo stile. Nel seguente esempio, il *<numero della versione>* viene mostrato in questo stile:

La directory per la fonte del kernel è `/usr/src/<version-number>/`, dove *<version-number>* è la versione del kernel installato su questo sistema.

Inoltre, noi adottiamo diverse strategie per attirare la vostra attenzione su alcune informazioni particolari. In base all'importanza che tali informazioni hanno per il vostro sistema, questi elementi verranno definiti nota bene, suggerimento, importante, attenzione o avvertenza. Per esempio:

**Nota Bene**

Ricordate che Linux distingue le minuscole dalle maiuscole. In altre parole, una rosa non è una ROSA né una rOsA.

**Suggerimento**

La directory `/usr/share/doc` contiene una documentazione aggiuntiva per i pacchetti installati sul vostro sistema.

**Importante**

Se modificate i file di configurazione DHCP, le modifiche non avranno effetto se non si riavvia il demone DHCP.

**Attenzione**

Non effettuate operazioni standard come utente root. Si consiglia di utilizzare sempre un account utente normale, a meno che non dobbiate amministrare il sistema.

**Avvertenza**

Fate attenzione a rimuovere solo le partizioni necessarie per Red Hat Enterprise Linux. Rimuovere altre partizioni può comportare una perdita dei dati oppure una corruzione dell'ambiente del sistema.

3. Attivate la vostra sottoscrizione

Prima di poter accedere alle informazioni sulla gestione del software e sul servizio, e prima di poter consultare la documentazione di supporto inclusa con la vostra registrazione, è necessario attivare la vostra sottoscrizione registrandovi con Red Hat. Il processo di registrazione include le seguenti fasi:

- Fornire un login di Red Hat
- Fornire il numero della vostra sottoscrizione
- Collegare il vostro sistema

La prima volta che eseguirete un avvio dell'installazione di Red Hat Enterprise Linux, vi verrà richiesto di registrarvi con Red Hat usando l'**Agent Setup**. Seguendo le diverse prompt dell'**Agent Setup**, sarete in grado di completare le fasi necessarie alla registrazione, attivando così la vostra sottoscrizione.

Se non siete in grado di completare la registrazione durante l'**Agent Setup** (il quale necessita di un accesso di rete), allora provate il processo di registrazione online di Red Hat disponibile su <http://www.redhat.com/register/>.

3.1. Fornite un login di Red Hat

Se non siete in possesso di un login di Red Hat, allora sarete in grado di crearne uno quando vi verrà richiesto durante l'**Agent Setup**, oppure online su:

<https://www.redhat.com/apps/activate/newlogin.html>

Un login di Red Hat vi permette di accedere a:

- Aggiornamenti software, errata e gestione tramite Red Hat Network
- Risorse per il supporto tecnico di Red Hat, documentazione, e Knowledgebase

Se avete dimenticato il vostro login di Red Hat, potete trovarlo online su:

https://rhn.redhat.com/help/forgot_password.pxt

3.2. Fornite il numero della vostra sottoscrizione

Il numero della vostra sottoscrizione si trova all'interno del pacchetto da voi ordinato. Se tale pacchetto non include il suddetto numero, allora la vostra sottoscrizione è già stata attivata, per questo motivo potete saltare questa fase.

Potete fornire il numero della vostra sottoscrizione quando richiesto durante l'**Agent Setup**, oppure visitando <http://www.redhat.com/register/>.

3.3. Collegare il vostro sistema

Il Red Hat Network Registration Client sarà utile per il collegamento del vostro sistema in modo tale da permettervi di ottenere gli aggiornamenti, e la gestione dei diversi sistemi. Sono disponibili tre diversi modi per eseguire tale collegamento:

1. Durante l'**Agent Setup** — Controllare le opzioni **Invia informazioni hardware** e **Invia un elenco del pacchetto del sistema** quando richiesto.
2. Dopo il completamento dell'**Agent Setup** — Dal **Menu Principale**, andate su **Tool di sistema**, per poi selezionare **Red Hat Network**.
3. Dopo il completamento dell'**Agent Setup** — Inserire il seguente comando dalla linea di comando come utente root:
 - `/usr/bin/up2date --register`

4. Ed altro ancora

Red Hat Enterprise Linux Introduzione al System Administration fa parte dell'impegno sempre crescente di Red Hat, nel fornire un supporto utile per gli utenti di Red Hat Enterprise Linux. Ad ogni nuova release di Red Hat Enterprise Linux, cercheremo di includere la documentazione nuova e quella aggiornata.

4.1. Inviateci suggerimenti

Se individuate degli errori nella *Red Hat Enterprise Linux Introduzione al System Administration*, o pensate di poter contribuire al miglioramento di questa guida, contattateci subito! Inviare i vostri suggerimenti. Vi preghiamo di inviare un report in Bugzilla (<http://bugzilla.redhat.com/bugzilla>) sul componente `rhel-isa`.

Assicuratevi di indicare l'identificatore del manuale:

```
rhel-isa(IT)-4-Print-RHI (2004-08-25T17:11)
```

Se indicate il suddetto identificatore, saremo in grado di sapere con esattezza di quale versione siete in possesso.

Se avete dei suggerimenti per migliorare la documentazione, cercate di essere il più specifici possibile. Se avete trovato un errore, vi preghiamo di includere il numero della sezione, e alcune righe di testo, in modo da agevolare le ricerca dell'errore stesso.

Capitolo 1.

La filosofia di un amministratore di sistema

Anche se le situazioni che si possono presentare ad un amministratore di sistema possono variare a seconda della piattaforma, vi sono alcuni temi che restano pressochè invariati. I suddetti temi costituiscono la filosofia dell'amministratore di sistema.

Questi i punti:

- Automatizzare ogni cosa
- Documentare tutto
- Comunicare il più possibile
- Conoscere le vostre risorse
- Conoscere i vostri utenti
- Conoscere la vostra compagnia
- La sicurezza non può essere considerata come un fattore secondario
- Pianificare a priori
- Prevedere qualsiasi imprevisto

Le seguenti sezioni affrontano i suddetti punti in modo dettagliato.

1.1. Automatizzare ogni cosa

Molti amministratori di sistema sono in numero — inferiori al numero di utenti, a quello dei sistemi o ad entrambi. In molti casi, l'automatizzazione è l'unico modo per venire a capo. In generale, ogni cosa o azione che si ripeta può essere automatizzata.

Ecco alcuni compiti che possono essere automatizzati:

- Controllo e riporto dello spazio disponibile su disco
- Backup
- Raccolta dei dati inerenti le prestazioni del sistema
- Gestione dello User account (creazione, cancellazione, ecc.)
- Funzioni specifiche dell'azienda (immissione di nuovi dati in un Web server, esecuzione di rapporti mensili/ trimestrali/annuali, ecc.)

Questo elenco però potrebbe risultare incompleto; le funzioni automatizzate dagli amministratori di rete, sono limitate solo dalla volontà dell'amministratore di scrivere gli script necessari. In questo caso, lasciate fare al computer la maggior parte del lavoro, questo potrebbe risultare il miglior approccio.

L'automatizzazione conferisce agli utenti una maggiore prevedibilità e consistenza del servizio.



Suggerimento

Se avete un compito che può essere automatizzato ricordatevi che voi non siete il primo amministratore al quale si presenta questa necessità. Ecco dove i benefici di un software open source si fanno più espliciti — in questa circostanza sareste in grado di automatizzare una procedura in modo da dedicare così la vostra attenzione ed il vostro tempo ad altri compiti. Per questo motivo assicuratevi

di eseguire una ricerca Web, prima di scrivere un qualcosa che risulti essere più complesso di uno script Perl.

1.2. Documentare ogni cosa

Se vi viene data la possibilità di installare un server nuovo o scrivere un documento sulle procedure su come eseguire i backup del sistema, l'amministratore di sistema medio sceglierà sempre di installare un server nuovo. È fortemente consigliato di documentare *sempre* quello che fate. Molti amministratori non documentano quello che fanno a causa di svariati motivi:

"Lo farò più tardi."

Sfortunatamente, questo non accade quasi mai. Infatti la natura dei compiti che un amministratore deve affrontare, è così caotica che in molti casi è impossibile "rinvviare". Ancora peggio, più si aspetta, e più facilmente si dimentica, creando così una documentazione non molto dettagliata (e quindi quasi inutile).

"Perché scrivere tutto? È più semplice ricordare."

A meno che non siate uno di quei pochissimi individui con una memoria fotografica, non sarete in grado di ricordare tutto. O peggio ancora, ne ricorderete solo la metà, senza sapere di aver perso l'intera storia. Ciò ne scatuisce una perdita di tempo nel cercare di ricordare quello che avete dimenticato o di correggere quello che avete scritto.

"Se lo terrò nella mia mente, non possono licenziarmi — sicurezza del posto di lavoro!"

Anche se questa tattica potrà essere valida per un pò di tempo — essa non risulta — essere molto sicura. Cercate di pensare per un momento cosa possa succedere durante una emergenza. Voi potreste non essere reperibili; la vostra documentazione potrebbe risultare importante permettendo ad un'altra persona di risolvere il problema durante la vostra assenza. In questi casi, è meglio avere la vostra documentazione in modo da risolvere il problema, che notare la vostra assenza e annoverarla come una delle cause del problema.

In aggiunta, se fate parte di una piccola organizzazione in continua crescita, sarà possibile che la stessa abbia bisogno di un nuovo amministratore. Come potete aspettarvi che questo nuovo amministratore possa sostituirvi se mantenete tutte le informazioni nella vostra testa? Ancora peggio, se non documentate niente, potrà essere richiesta sempre la vostra presenza, 'sareste indispensabili', in modo da intaccare la vostra carriera e quindi le vostre promozioni a compiti dirigenziali. Potreste trovarvi a svolgere il lavoro della persona assunta per assistervi nel vostro incarico.

Si spera che adesso abbiate capito l'importanza di una documentazione. Ciò ci porta a formulare un'altra domanda: Cosa documentare? Ecco un breve elenco:

Policy

Le policy vengono scritte per ufficializzare e chiarire la relazione che intercorre tra voi e la community dei vostri utenti. Esse chiariscono ai vostri utenti come le richieste sulle risorse vengano assistite e gestite. La natura, lo stile, ed il metodo di diffusione delle policy per la vostra community varia da organizzazione a organizzazione.

Procedure

Le procedure non sono altro che una sequenza di azioni passo dopo passo da intraprendere per poter completare un compito. Le procedure da documentare possono includere le procedure di

backup, quelle di gestione dello user account, inerenti le problematiche nel riporto delle procedure, e così via. Come l'automazione, se una procedura viene seguita più di una volta, è consigliabile documentarla.

Modifiche

Gran parte della carriera di un amministratore di sistema ruota intorno alle modifiche — alla configurazione dei sistemi per massimizzarne le prestazioni, alle modifiche dei file di configurazione, agli script e così via. Tutte queste modifiche devono essere documentate in qualche modo. In caso contrario, potreste trovarvi spiazzati a causa di una modifica eseguita in passato.

Alcune organizzazioni utilizzano metodi più complessi per registrare le suddette modifiche, ma in molti casi potrebbe essere solo necessario una semplice nota all'inizio del file modificato. Come minimo, ogni entry presente nella cronologia dovrebbe contenere:

- Il nome o le iniziali della persona che ha eseguito la modifica
- La data nella quale è stata eseguita la modifica
- La ragione per la quale è stata apportata la modifica

Il risultato, e le entry più utili:

ECB, 12-Giugno-2002 — Entry aggiornata per nuove stampanti di contabilità (per supportare l'abilità di stampa duplex della stampante)

1.3. Comunicare il più possibile

Per quanto riguarda i vostri utenti, la comunicazione non è mai abbastanza. Fate attenzione che anche le più piccole modifiche che per voi potrebbero essere irrilevanti, possono essere motivo di confusione per la gestione del personale (Human Resources).

Il metodo attraverso il quale comunicate con i vostri utenti può variare a seconda della vostra organizzazione. Alcune organizzazioni usano le email, altre utilizzano un sito web interno, altre ancora possono affidarsi alle comunicazioni di rete o IRC. Potrebbe risultare sufficiente in questo contesto, l'utilizzo di un pezzo di carta da affiggere nell'area di ricreazione. In ogni caso, utilizzate qualsiasi metodo che ritenete più opportuno.

In generale, è consigliato seguire uno di questo consigli:

1. Comunicate ai vostri utenti cosa farete
2. Comunicate ai vostri utenti cosa state per fare
3. Comunicate ai vostri utenti cosa avete fatto

Le seguenti sezioni affrontano queste fasi in maniera più dettagliata.

1.3.1. Comunicate ai vostri utenti cosa farete

Assicuratevi di avvisare i vostri utenti prima di fare qualcosa. Il tempo di preavviso per i vostri utenti varia a seconda del tipo di modifica (l'aggiornamento di un sistema operativo richiede un tempo di preavviso maggiore rispetto a quello necessario per la modifica del colore di default della schermata di login del sistema), e della natura della community dei vostri utenti (gli utenti più esperti possono essere in grado di gestire le modifiche molto più facilmente rispetto ad utenti meno esperti.)

Come minimo, è necessario descrivere:

- La natura della modifica
- Quando verrà messa in atto

- Perché è stata eseguita tale modifica
- Il tempo necessario per tale modifica
- L'impatto (se causato) che gli utenti possono aspettarsi a causa di tale modifica
- Informazioni su chi contattare in caso di necessità

Ecco una ipotetica situazione. La sezione responsabile alle finanze presenta alcuni problemi dovuti alla lentezza del proprio database. Voi disattiverete il server, aggiornerete il modulo CPU utilizzando così un modello più veloce, ed eseguirete successivamente un riavvio. Una volta aver seguito la suddetta procedura, appronterete il vostro database con uno storage basato sul RAID più veloce. Ecco cosa potreste fare per informare i vostri utenti:

Downtime del sistema previsto per venerdì notte

Inizio questo venerdì ore 18:00, tutte le applicazioni interessate non saranno disponibili per un periodo di circa quattro ore.

Durante il quale saranno apportate le modifiche all'hardware e al software del server del database in questione. Queste modifiche avranno lo scopo di diminuire il tempo necessario per eseguire le applicazioni del tipo Accounts Payable 'per la fatturazione passiva' e Accounts Receivable, ed i rapporti sull'attività settimanale.

Oltre alla modifica nel runtime, la maggior parte delle persone non dovrebbero notare altri cambiamenti. Tuttavia, coloro che hanno scritto le proprie richieste SQL, devono essere a conoscenza che la struttura di alcuni indici potrebbero essere stati modificati. Ciò viene documentato nell'intranet del sito web della compagnia, nella pagina inerente le finanze dell'azienda stessa.

Se avete delle domande, commenti o problematiche, si prega di contattare l'amministratore del sistema digitando l'interno 4321.

È importante notare:

- Comunicare in modo effettivo l'inizio e la durata di ogni downtime che si può verificare a causa della modifica.
- Assicuratevi d'informare in modo efficiente *tutti* gli utenti, sull'orario della modifica.
- Utilizzate termini facili da capire. Alle persone coinvolte in questa modifica, non importa che il nuovo modulo CPU è una unità di 2GHz con una cache L2 doppia, o che il database è stato posizionato su di un volume logico RAID 5.

1.3.2. Comunicate ai vostri utenti cosa state facendo

Questa fase è generalmente un avviso del tipo 'last minute' della modifica che stà per avvenire; pertanto dovrebbe essere una breve ripetizione del primo messaggio, rendendo più visibile la natura della modifica ("L'aggiornamento del sistema verrà eseguito QUESTA SERA."). Questo è anche il momento migliore per rispondere a qualsiasi domanda ricevuta a causa del primo messaggio.

Continuando il nostro ipotetico esempio, ecco un ultimissimo avviso:

Downtime del sistema previsto per questa sera

Ricordatevi. Il downtime del sistema annunciato questo Lunedì, avverrà come previsto questa sera alle ore 18:00. Potete trovare il messaggio originale nella pagina intranet del sito web della compagnia, nella pagina amministrativa del sistema.

Diverse persone hanno richiesto di poter terminare prima questa sera in modo da eseguire un backup del loro lavoro prima del periodo di downtime. Questo non sarà necessario, in quanto il lavoro svolto questa sera non avrà alcun impatto su quello eseguito sulla vostra workstation personale.

Ricordate, coloro che hanno salvato le proprie richieste SQL, devono essere a conoscenza che la struttura di alcuni indici sarà cambiata. Ciò viene documentato nell'intranet della compagnia sulla pagina delle finanze.

I vostri utenti sono stati avvisati, adesso sarete in grado di svolgere il vostro lavoro.

1.3.3. Comunicate ai vostri utenti cosa avete fatto

Dopo aver finito le modifiche, *dovete* informare i vostri utenti su quello che avete fatto. Ancora una volta, questo dovrebbe essere un sunto dei messaggi precedenti (ci sarà sempre qualcuno che non li ha letti.)¹

Tuttavia, vi è ancora un'aggiunta importante da fare. È vitale informare i vostri utenti dello stato attuale del sistema. L'aggiornamento non è stato molto semplice da eseguire? Il server dello storage era in grado di servire solo i sistemi della sezione d'ingegneria e non quelli delle finanze? Queste problematiche devono essere affrontate qui.

Ovviamente, se lo stato corrente differisce da quello comunicato precedentemente, dovrete allora descrivere le fasi da intraprendere per arrivare alla soluzione finale.

Nel nostro esempio, il periodo di downtime presenta alcuni problemi. Il nuovo modulo CPU non ha funzionato; una telefonata al rivenditore del sistema ha rivelato che è necessaria una nuova versione. Nell'altra ipotesi, la migrazione del database sul volume RAID ha avuto successo (anche se ha richiesto un tempo maggiore di quello previsto a causa dei problemi dovuti al modulo della CPU.)

Ecco una possibile comunicazione:

Periodo di downtime del sistema terminato

Il periodo di downtime previsto per venerdì sera (consultare la pagina di amministrazione del sistema presente nell'intranet della compagnia) è terminato. Sfortunatamente, alcune problematiche riguardanti l'hardware non hanno reso possibile il completamento di uno dei compiti previsti. A causa di ciò, i restanti compiti hanno avuto bisogno di un periodo più lungo di quello previsto. Tutti i sistemi sono tornati alla normalità alla mezzanotte.

A causa delle restanti problematiche hardware, le prestazioni AP, AR e dei rapporti sui fogli di bilancio saranno leggermente migliorati ma non in modo previsto. Un secondo periodo di downtime verrà programmato e annunciato appena le problematiche sopra indicate saranno risolte.

Si prega di notare che il periodo di downtime ha effettivamente modificato gli indici del database; le persone che hanno salvato le loro richieste SQL, dovrebbero consultare la pagina inerente alle Finanze nell'intranet della compagnia.

Si prega di contattare l'amministratore all'interno 4321 per eventuali chiarimenti.

Con questo tipo di informazioni, i vostri utenti avranno elementi sufficienti per continuare il loro lavoro, e capire come i cambiamenti possano influenzare i loro compiti.

1.4. Conoscere le vostre risorse

Il lavoro dell'amministratore di sistema è basato principalmente nel bilanciamento delle risorse disponibili nei confronti delle persone e dei programmi che ne fanno uso. Per questo motivo, la vostra

1. Assicuratevi d'inviare questo messaggio subito dopo aver terminato il vostro lavoro, *prima* di tornare a casa. Una volta lasciato l'ufficio, è molto facile dimenticare, lasciando i vostri utenti nel dubbio se usare o meno il sistema.

carriera come amministratore di sistema sarà breve e piena di stress se non riuscirete a capire le risorse a voi disponibili.

Alcune delle risorse sono quelle che a voi possono sembrare più ovvie:

- Risorse del sistema, come ad esempio la potenza di processazione, la memoria e lo spazio del disco
- Larghezza di banda della rete
- Risorse finanziarie presenti nel budget IT

Ma alcune potrebbero non essere così ovvie:

- I servizi del personale operativo, altri amministratori, oppure un assistente amministrativo
- Tempo (spesso di rilevante importanza quando esso comprende il tempo necessario per eseguire i backup di un sistema)
- Conoscenza (se presente nei libri, nella documentazione del sistema, o se posseduta da una singola persona che lavora per la compagnia da venti anni)

È importante tener presente che è di fondamentale importanza registrare le risorse a voi disponibili in modo da poter mantenerle *sempre aggiornate* — una carenza di "conoscenza delle circostanze" per quanto riguarda le risorse disponibili, potrebbe risultare più grave di una *non* conoscenza delle stesse.

1.5. Conoscere i vostri utenti

Anche se alcune persone vengono terrorizzate dal termine "utenti" (dovuto anche all'uso fatto in termini denigratori da parte di alcuni amministratori di sistema), il suddetto termine viene qui usato senza implicarne alcuna allusione. Gli utenti sono quelle persone che usano le risorse ed i sistemi per i quali voi siete i responsabili —. Per questo motivo, essi rappresentano il fulcro per il vostro successo nel gestire i vostri sistemi; se non siete in grado di capire i vostri utenti, come potete aspettarvi di capire anche le risorse a loro necessarie?

Per esempio, considerate uno sportello bancario. Il suddetto sportello usa delle applicazioni ben specifiche e richiede molto poco in termini di risorse del sistema. Un ingegnere software, d'altro canto, potrebbe utilizzare diverse applicazioni e gradire un numero sempre maggiore di risorse (per tempi di creazione più veloci). Due utenti completamente diversi con diversissime esigenze.

Assicuratevi di acquisire quanto più potete in termini di conoscenza dei vostri utenti.

1.6. Conoscere la vostra compagnia

Sia che lavoriate per una grande compagnia multinazionale o per una piccola community universitaria, è necessario conoscere sempre la natura dell'ambiente aziendale nel quale lavorate. Tutto questo può portare alla formulazione di una domanda:

Qual'è il compito dei sistemi che amministro?

La chiave per questo è rappresentata dal fatto di capire i compiti dei sistemi in un senso più globale:

- Le applicazioni che devono essere eseguite all'interno di un frangente, come ad esempio un mese, tre mesi oppure un anno
- I periodi entro i quali deve essere eseguito il mantenimento del sistema
- Nuove tecnologie che possono essere implementate per la risoluzione dei problemi aziendali

Prendendo in considerazione il lavoro della vostra compagnia, troverete che le decisioni giornaliere da voi prese risulteranno essere le più idonee sia per i vostri utenti che per voi stessi.

1.7. La sicurezza non può essere considerata come un fattore secondario

Non ha importanza cosa pensiate dell'ambiente nel quale vengono eseguiti i vostri sistemi, non potete considerare la sicurezza come un qualcosa di scontato. Anche i sistemi del tipo 'standalone' non collegati ad internet, possono risultare a rischio (anche se ovviamente il loro rischio risulta essere diverso da un sistema collegato al mondo esterno).

Per questo motivo è molto importante considerare le ripercussioni sulla sicurezza dovute alle vostre azioni. Il seguente elenco illustra le diverse problematiche da considerare:

- La natura dei possibili rischi per ogni sistema soggetto alla vostra attenzione
- La posizione, il tipo, ed il valore dei dati presenti sui sistemi
- Il tipo e la frequenza degli accessi autorizzati ai sistemi

Mentre voi state pensando alla sicurezza, non dimenticate che i potenziali intrusi non sono solo quelli che attaccano il vostro sistema dall'esterno. Molte volte la persona che perpetra una violazione risulta essere un utente interno alla compagnia. In questo modo la prossima volta che vi aggirate tra le scrivanie dell'ufficio, guardatevi intorno e domandatevi:

Cosa può accadere se *quella* persona cerca di compromettere la nostra sicurezza?



Nota Bene

Ciò *non* significa che dovete trattare i vostri colleghi come se fossero dei criminali. Significa soltanto di controllare il tipo di lavoro che una persona svolge, per capire di conseguenza il potenziale rischio che tale persona possa rappresentare nei confronti della sicurezza.

1.7.1. I rischi dell'ingegneria sociale

Mentre la prima reazione degli amministratori, per quanto riguarda le problematiche riguardanti la sicurezza, può essere quella di concentrarsi sull'aspetto tecnologico, è molto importante mantenere una certa prospettiva. Molto spesso, le violazioni che si verificano nel campo della sicurezza non presentano un fattore tecnologico, ma presentano un fattore umano.

Le persone interessate a tali violazioni, spesso usano un fattore umano per poter baipassare i controlli tecnologici dovuti all'accesso. Questo metodo è conosciuto come *ingegneria sociale*. Ecco un esempio:

L'operatore del secondo turno riceve una telefonata esterna. Colui che chiama dice di essere il CFO dell'organizzazione (il nome e le informazioni personali del CFO sono state ottenute dal sito web dell'organizzazione, nella pagina "Team Manageriale").

La persona che chiama dice di trovarsi in tutt'altra parte del mondo (forse questo fa parte di una storia inventata, oppure la vostra pagina web riporta che il CFO della vostra compagnia è presente in un determinato salone in una città di un'altra nazione).

L'impostore racconta tutte le sue disavventure; il portatile è stato rubato all'aeroporto, ed è con un cliente molto importante e necessita di accedere alla pagina intranet dell'azienda, per poter controllare così lo stato dell'account del cliente stesso. Con molta gentilezza questa persona chiede di poter accedere alle informazioni necessarie.

Sapete cosa potrebbe fare il vostro operatore in circostanze simili? Se il vostro operatore non ha una policy specifica a riguardo (in termini di procedure), voi non saprete mai con sicurezza cosa potrebbe accadere.

L'obiettivo delle policy e delle procedure è quello di fornire una guida chiara su quello che si può e non si può fare. Tuttavia, le suddette procedure funzionano se tutte le persone le osservano. Ma si presenta sempre lo stesso problema — e cioè risulterà difficile che tutte le persone rispetteranno le suddette policy. Infatti, a seconda della natura della vostra organizzazione, è possibile che voi non abbiate l'autorità sufficiente a definire tali procedure, e tanto meno a farle rispettare. E allora?

Sfortunatamente, non vi sono risposte facili a questa domanda. L'educazione degli utenti potrebbe esservi molto utile; fate ciò che potete per informare la vostra community sulle procedure inerenti la sicurezza, e metteteli a conoscenza della presenza di tale fenomeno. Organizzate dei meeting durante l'ora di pranzo. Inviare degli articoli relativi alla sicurezza tramite le mailing list dell'organizzazione. Rendetevi disponibili per rispondere ad eventuali domande che possono essere formulate da parte degli utenti.

In breve, inculcate il più possibile queste idee ai vostri utenti.

1.8. Pianificare a priori

Gli amministratori che seguono le suddette indicazioni e fanno il loro meglio per seguirne i consigli, risulteranno essere degli amministratori molto validi — ma solo per un giorno. Probabilmente l'ambiente potrebbe cambiare, ed il nostro fantastico amministratore potrebbe trovarsi con le spalle al muro. La ragione? Il nostro super amministratore non è stato in grado di pianificare il suo lavoro.

Certamente nessuno può predire il futuro con una esattezza del 100%. Tuttavia, con un pò di buon senso potrebbe risultare semplice predire diversi cambiamenti:

- Un accenno fatto durante uno dei meeting infrasettimanali, può rappresentare un segno premonitore del bisogno di supportare i vostri nuovi utenti in un prossimo futuro.
- La discussione di una possibile acquisizione può significare che voi potreste essere i responsabili di nuovi (e forse incompatibili) sistemi in uno o più luoghi remoti

Essere abili a leggere questi segnali (e rispondervi in modo efficiente), può rendere la vostra vita e quella dei vostri utenti più tranquilla.

1.9. Aspettatevi un imprevisto

Mentre la frase "aspettatevi un imprevisto" può sembrare banale, essa riflette una inconfondibile verità, infatti tutti gli amministratori devono comprendere:

Ci *saranno* momenti dove sarete presi alla sprovvista.

Una volta accettato questo fattore, cosa può fare un amministratore di sistema? La risposta sta nella flessibilità; eseguendo il vostro lavoro in modo da avere sempre sia per voi che per i vostri utenti, il maggior numero di opzioni disponibili. Prendete, per esempio, il problema dello spazio sul disco. Dato per scontato che lo spazio presente non è mai sufficiente, è ragionevole assumere che in un determinato momento avrete un *disperato* bisogno di spazio aggiuntivo.

Cosa potrebbe fare un amministratore che ha preso in considerazione il verificarsi di un imprevisto? È consigliabile avere disponibile una certa impostazione dell'unità disco in caso si verificassero alcuni problemi hardware². Un sostituto di questo tipo potrebbe essere rapidamente impiegato³ momentaneamente per risolvere il temporaneo problema di spazio del disco, dando più tempo alla

2. Ovviamente un amministratore così meticoloso usa naturalmente RAID (oppure tecnologie relative), per minimizzare l'impatto dovuto ad un problema verificatosi nell'unità disco durante il normale funzionamento.

3. Ed ancora, gli amministratori che cercano di pensare in anticipo configurano i loro sistemi in modo da poter aggiungere velocemente una unità al sistema.

risoluzione permanente del problema stesso (per esempio, seguendo la procedura standard per fornire unità disco aggiuntive).

Cercando di anticipare i problemi prima che gli stessi si possano verificare, vi troverete in una posizione dove sarete in grado di rispondere più velocemente ed in modo più efficace.

1.10. Informazioni specifiche su Red Hat Enterprise Linux

Questa sezione riporata le informazioni relative alla filosofia degli amministratori del sistema, specifiche a Red Hat Enterprise Linux.

1.10.1. Automazione

L'automazione di compiti ripetuti regolarmente con Red Hat Enterprise Linux richiede la conoscenza di diverse tecnologie. Le prime sono rappresentate da quelle tecnologie che controllano il tempo di esecuzione dei comandi o dello script. I comandi `cron` e `at` sono i più utilizzati in questi ruoli.

Incorporando un tempo di specificazione del sistema, flessibile e potente ma allo stesso tempo facile da capire, `cron` è in grado di programmare l'esecuzione di comandi e script per intervalli regolari che vanno da qualche minuto a qualche mese. Il comando `crontab` viene usato per manipolare i file che controllano il demone `cron` che programma l'esecuzione di ogni compito `cron`.

Il comando `at` (ed il comando relativo `batch`) sono più idonei per programmare l'esecuzione di script o comandi che possono essere utilizzati solo una volta. Questi comandi implementano un sottosistema di batch rudimentale che consiste in code multiple con diverse priorità di programmazione. Queste priorità sono conosciute come livelli *nice*ness (a causa del nome del comando — *nice*). Entrambi `at` e `batch` sono perfetti per compiti che devono iniziare in un determinato momento ma che non hanno un tempo limite determinato.

Successivamente abbiamo gli scripting language. Essi sono dei "linguaggi di programmazione" che l'amministratore di sistema medio utilizza per automatizzare operazioni manuali. Vi sono diversi linguaggi di programmazione (ed ogni amministratore tende ad avere quello a lui/lei più congeniale), ma i seguenti sono quelli più usati:

- La command shell `bash`
- Lo scripting language `perl`
- Lo scripting language `python`

Andando oltre le differenze più ovvie tra questi linguaggi, quella più importante è il modo in cui questi linguaggi interagiscono con altri programmi utility su di un sistema Red Hat Enterprise Linux. Gli script scritti con la shell `bash`, tendono ad avere un uso più vasto di tanti programmi utility più piccoli (per esempio, per eseguire una manipolazione di una stringa di caratteri), mentre gli script `perl` eseguono questo tipo di operazione usando le caratteristiche create all'interno del linguaggio stesso. Uno script scritto usando `python` può sfruttare pienamente le capacità orientate sull'oggetto del linguaggio, rendendo gli script più complessi, più facilmente ampliabili.

Questo significa che, per eseguire un master shell scripting, dovete avere una certa familiarità con i diversi programmi utility (come ad esempio `grep` e `sed`) che fanno parte di Red Hat Enterprise Linux. Conoscere `perl` (e `python`), d'altra parte, tende ad essere un processo più "autonomo". Tuttavia, molte impostazioni del linguaggio `perl`, sono basate sulla sintassi di vari programmi utility tradizionali di UNIX, e di conseguenza possono essere familiari agli amministratori di un sistema Red Hat Enterprise Linux con una certa esperienza di shell scripting.

1.10.2. Documentazione e comunicazione

Nell'area di documentazione e comunicazione, non vi è molto di specifico su Red Hat Enterprise Linux. Poichè le suddette aree possono consistere in commenti, file di configurazione basati sul testo, aggiornamento di una pagina web o di un invio di email, un amministratore di sistema che utilizza Red Hat Enterprise Linux, deve avere accesso agli editor di testo, editor HTML, e mail client.

Ecco un piccolo esempio dei diversi editor di testo disponibili con Red Hat Enterprise Linux:

- L'editor di testo **gedit**
- L'editor di testo **Emacs**
- L'editor di testo **Vim**

L'editor di testo **gedit** rappresenta un'applicazione strettamente grafica (in altre parole, necessita di un ambiente X Window System attivo), mentre **vim** e **Emacs** sono di natura basati sul testo.

La scelta del miglior editor di testo ha dato luogo ad un dibattito iniziato da quando esistono i computer, continuando fino ad oggi, e forse si protrarrà per moltissimi altri anni. Pertanto, il miglior approccio è quello di provare ogni editor di testo, e usare quello che ritenete più idoneo.

Per gli editor HTML, gli amministratori di sistema possono usare la funzione Composer del web browser di **Mozilla**. Naturalmente, alcuni amministratori di sistema preferiscono codificare manualmente i loro HTML, trasformando un editor di testo regolare in un tool accettabile.

Red Hat Enterprise Linux include l'email client grafico **Evolution**, **Mozilla** (anch'esso grafico), e **mutt**, il quale è basato sul testo. Poichè con gli editor di testo la scelta di una email client tende ad essere personale; il migliore approccio è quello di provare in prima persona ogni client, e usare quello a voi più idoneo.

1.10.3. Sicurezza

Come precedentemente accennato in questo capitolo, la sicurezza non può essere considerata come un fattore secondario, e la stessa sotto Red Hat Enterprise Linux risulta essere molto importante. I controlli sull'autenticazione e sull'accesso sono profondamente integrati nel sistema operativo, e sono basati su design maturati dalla lunga esperienza nella community UNIX.

Per l'autenticazione Red Hat Enterprise Linux utilizza PAM — Pluggable Authentication Modules. PAM rende possibile l'affinamento dell'autenticazione dell'utente attraverso la configurazione delle librerie condivise usate da tutte le applicazioni che supportano PAM, senza la necessità di modificare le applicazioni stesse.

Il controllo dell'accesso con Red Hat Enterprise Linux utilizza i permessi tradizionali stile-UNIX (read, write, execute), rispetto alle classificazioni user, group, e "altri". Come UNIX, Red Hat Enterprise Linux fa uso anche di bit *setuid* e *setgid*, per conferire temporaneamente il diritto d'accesso ai processi che eseguono un programma particolare, basato sull'ownership del program file. Naturalmente, tutto ciò rende necessario che tutti i programmi che vengono eseguiti con i privilegi *setuid* o *setgid*, devono essere verificati attentamente per assicurare che non esistano alcuni punti deboli.

Red Hat Enterprise Linux include anche il supporto per l'*access control lists*. Un access control list (ACL) permette un controllo minuzioso sugli utenti o gruppi che possono accedere ad un file o ad una directory. Per esempio, i permessi di un file possono negare l'accesso a tutti eccetto al proprietario del file, ed anche l'ACL del file può essere configurato per permettere solo all'utente *bob* di scrivere, e al gruppo *finance* di leggere il file.

Un altro aspetto della sicurezza è rappresentato dal fatto di poter mantenere una registrazione dell'attività del sistema. Red Hat Enterprise Linux fa un ampio uso del logging, sia ad un livello kernel che ad un livello dell'applicazione. Il logging viene controllato dal demone di logging del sistema *syslogd*, il quale è in grado di registrare in modo locale le informazioni del sistema

(normalmente conservate nella directory `/var/log/`), oppure su di un sistema remoto (il quale funge come un apposito server di registrazione per computer multipli.)

Gli Intrusion detection system (IDS) sono dei tool molto potenti utili ad ogni amministratore di sistema. Un IDS rende possibile per gli amministratori di sistema, determinare la presenza di modifiche non autorizzate su uno o più sistemi. Il design medio del sistema operativo include una funzionalità simile all'IDS.

Poichè Red Hat Enterprise Linux è installato utilizzando l'RPM Package Manager (RPM), è possibile l'utilizzo dell'RPM stesso per verificare se qualsiasi modifica fatta ai pacchetti, possa compromettere il sistema operativo. Tuttavia, poichè RPM è innanzitutto un tool di gestione del pacchetto, le sue abilità come IDS sono limitate. Nonostante tutto, rappresenta un primo passo nei confronti del controllo di modifiche non autorizzate di un sistema Red Hat Enterprise Linux.

1.11. Risorse aggiuntive

Questa sezione include diverse risorse che possono essere usate per ottenere maggiori informazioni sulla filosofia dell'amministratore del sistema e su argomenti specifici di Red Hat Enterprise Linux affrontati in questo capitolo.

1.11.1. Documentazione installata

Le seguenti risorse sono state installate nel corso di una installazione tipica di Red Hat Enterprise Linux, e possono aiutarvi ad ottenere maggiori informazioni su argomento trattati in questo capitolo.

- Pagina man di `crontab(1)` e `crontab(5)` — Imparate a programmare i comandi e gli script per una esecuzione automatica ad intervalli regolari.
- Pagina man di `at(1)` — Imparate a programmare i comandi e gli script per una esecuzione futura.
- Pagina man di `bash(1)` — Per saperne di più sulla shell di default e sulla shell script writing.
- Pagina man di `perl(1)` — Revisione degli indici delle diverse pagine man che costituiscono la documentazione online di perl.
- Pagina man di `python(1)` — Per saperne di più sulle opzioni, file, e sulle variabili dell'ambiente che controllano il Python interpreter.
- Pagina man di `gedit(1)` ed entry del menu **Aiuto** — Imparate a modificare i file di testo con questo editor di testo grafico.
- Pagina man di `emacs(1)` — Per saperne di più su questo editor di testo molto flessibile, e su come esercitarsi online.
- Pagina man di `vim(1)` — Imparate ad utilizzare questo potente editor di testo.
- Entry del menu **Contenuti d'aiuto Mozilla** — Imparate a modificare i file HTML, leggere la posta, e navigare sul web.
- Pagina man di `evolution(1)` ed entry del menu **Aiuto** — Imparate a gestire le vostre email con il vostro client email grafico.
- Pagina man di `mutt(1)` e file in `/usr/share/doc/mutt-<version>` — Per imparare a gestire le vostre email con questo email client basato sul testo.
- Pagina man di `pam(8)` e file in `/usr/share/doc/pam-<version>` — Per sapere come avviene l'autenticazione con Red Hat Enterprise Linux.

1.11.2. Siti web utili

- <http://www.kernel.org/pub/linux/libs/pam/> — La homepage del progetto Linux-PAM.
- <http://www.usenix.org/> — La homepage di USENIX. Una organizzazione professionale riunisce tutti i professionisti del computer, affidandone innovazioni e comunicazioni.
- <http://www.sage.org/> — La homepage dei System Administrators Guild. Un gruppo tecnico speciale USENIX che rappresenta una buona risorsa per tutti gli amministratori responsabili dei sistemi operativi Linux (o Linux-like).
- <http://www.python.org/> — Il sito web del linguaggio Python. Un sito eccellente per saperne di più su Python.
- <http://www.perl.org/> — Il sito web Perl Mongers. Il luogo ideale per ottenere informazioni su Perl e per collegarsi con relativa community.
- <http://www.rpm.org/> — La homepage di RPM Package Manager. Il sito web più completo per informazioni su RPM.

1.11.3. Libri correlati

La maggior parte dei libri sulla gestione del sistema, fanno davvero poco per affrontare la problematica inerente la filosofia dell'amministratore. Tuttavia, i seguenti libri presentano alcune sezioni che affrontano in modo più approfondito tale argomento:

- *Red Hat Enterprise Linux Reference Guide*; Red Hat, Inc. — Fornisce una panoramica sui luoghi dei file più importanti del sistema, sulle impostazioni dell'utente e del gruppo, e sulla configurazione PAM.
- *Red Hat Enterprise Linux Security Guide*; Red Hat, Inc. — Affronta in modo esauriente le problematiche relative alla sicurezza per gli amministratori del sistema Red Hat Enterprise Linux.
- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — Include i capitoli sulla gestione degli utenti e dei gruppi, compiti automatizzati, e gestione dei log file.
- *Linux Administration Handbook* di Evi Nemeth, Garth Snyder, and Trent R. Hein; Prentice Hall — Presenta una sezione molto utile sulle policy e sugli effetti secondari di un amministratore di sistema, incluso diverse discussioni del tipo "what-if" riguardanti l'etica.
- *Linux System Administration: A User's Guide* di Marcel Gagne; Addison Wesley Professional — Contiene un capitolo molto utile su come automatizzare i diversi compiti.
- *Solaris System Management* di John Philcox; New Riders Publishing — Anche se non specifico per Red Hat Enterprise Linux (o in generale per Linux), usando il termine "manager del sistema" invece di "amministratore del sistema", questo libro contiene 70 pagine inerenti i diversi ruoli ricoperti dall'amministratore del sistema in una normale organizzazione.

Capitolo 2.

Controllo delle risorse

Come precedentemente detto, una grossa importanza nella gestione del sistema viene data all'utilizzo efficiente delle risorse. Bilanciando le diverse risorse tra le persone ed i programmi che ne fanno uso, è possibile risparmiare ingenti risorse economiche, facilitandone l'uso delle stesse tra gli utenti. Tuttavia, tale approccio può far sorgere alcune domande:

Che cosa sono le risorse?

E:

Com'è possibile sapere quali risorse sono state utilizzate (e in quali proporzioni)?

Lo scopo di questo capitolo è quello di poter rispondere a queste domande, aiutandovi a conoscere le risorse e a scoprire il modo per poterle controllare:

2.1. Concetti di base

Prima di poter controllare le risorse, bisogna sapere quali sono quelle da controllare. Tutti i sistemi presentano le seguenti risorse:

- Potenza della CPU
- Larghezza della banda
- Memoria
- Storage

Queste risorse vengono analizzate in modo più dettagliato nei seguenti capitoli. Tuttavia, per adesso tutto quello che dovete ricordare è che le risorse hanno un impatto diretto sulle prestazioni del sistema, e conseguentemente sulla produttività e soddisfazione dei vostri utenti.

Il controllo delle risorse risulta essere una raccolta di informazioni sull'utilizzo di una o più risorse del sistema.

Tuttavia non è così facile. Innanzitutto tenete in considerazione le risorse da controllare. Successivamente è necessario esaminare ogni sistema da controllare, facendo particolare attenzione alla situazione di ogni singolo sistema.

Il sistema da controllare rientra in una delle seguenti categorie:

- Il sistema ha incontrato dei problemi e voi desiderate migliorare le sue prestazioni.
- Il sistema attualmente funziona correttamente, e voi desiderate che continui ad operare in questa direzione.

La prima categoria sta ad indicare che dovete controllare le risorse dall'aspetto delle prestazioni del sistema, mentre la seconda implica un controllo delle risorse da una prospettiva di pianificazione della capacità.

A causa dei diversi requisiti unici ad ogni prospettiva, le seguenti sezioni affrontano ogni categoria in modo più approfondito.

2.2. Controllo della prestazione del sistema

Come accennato precedentemente, il controllo delle prestazioni del sistema viene eseguito in risposta ad un problema che si è verificato nelle prestazioni stesse. Si può verificare che il sistema è troppo lento, o che i programmi (e talvolta l'intero sistema) non sono in grado di funzionare. In entrambi i casi, il controllo delle prestazioni viene normalmente fatto seguendo le suddette fasi:

1. Eseguire un controllo in modo da identificare la natura e la portata della carenza delle risorse che causano tali problemi
2. I dati prodotti vengono analizzati, facendo seguire successivamente delle azioni (generalmente eseguendo una regolazione delle prestazioni e/o ottenere hardware aggiuntivo) per la risoluzione del problema
3. Eseguire un controllo per accertarsi che tale problema è stato risolto

Per questo motivo il controllo delle prestazioni tende ad essere generalmente molto breve ma molto dettagliato.



Nota Bene

Il controllo delle prestazioni del sistema è un processo interattivo il quale necessita di una ripetizione di queste fasi, per arrivare ad ottenere le migliori prestazioni possibili del sistema stesso. Per questo le risorse del sistema ed il loro utilizzo tendono ad essere in relazione tra loro, ciò significa che l'eliminazione di un determinato problema ne potrebbe portare alla luce uno nuovo.

2.3. Controllo della capacità del sistema

Il controllo della capacità del sistema viene eseguito come parte di un programma di pianificazione della capacità stessa. Tale pianificazione utilizza un controllo delle risorse per determinare la frequenza dei cambiamenti durante l'utilizzo delle risorse stesse. Una volta conosciuta tale frequenza, è possibile eseguire una pianificazione molto più accurata sull'acquisizione di risorse aggiuntive.

Il controllo eseguito a scopo di pianificazione della capacità, differisce dal controllo delle prestazioni:

- Il controllo viene eseguito in base più o meno continua
- Il controllo non è generalmente molto dettagliato

La ragione di queste differenze sta negli obiettivi del programma di pianificazione della capacità. Tale pianificazione richiede un punto di vista molto "più ampio"; l'utilizzo anomalo o breve di una risorsa non rappresenta un grosso problema. I dati invece, vengono raccolti attraverso un determinato periodo di tempo, rendendo possibile la classificazione dell'utilizzo delle risorse in termini di cambiamenti del carico di lavoro. In ambienti più specifici, (per esempio dove viene eseguita solo un'applicazione), è possibile modellare l'impatto di una applicazione sulle risorse del sistema. Questo può essere eseguito con sufficiente accuratezza in modo da determinare, per esempio, l'impatto dovuto alla presenza di cinque utenti che eseguono una determinata applicazione durante il periodo più movimentato della giornata.

2.4. Cosa controllare?

Come affrontato precedentemente, le risorse presenti in un sistema sono la potenza della CPU, la larghezza della banda, la memoria, e lo storage. A prima vista sembrerebbe che il controllo consista solo nel monitorare questi quattro elementi.

Sfortunatamente, non è così semplice. Per esempio, considerate una unità disco. Cosa desiderate sapere sulle sue prestazioni?

- Quanto spazio è disponibile?
- Quante operazioni I/O in generale, esegue ogni secondo?
- Quanto tempo è necessario per completare una operazione I/O?
- Quante operazioni I/O vengono lette? E quante ne vengono scritte?
- A quanto ammontano i dati letti/scritti con ogni operazione I/O?

Vi sono svariati modi per studiare le prestazioni dell'unità disco; questi sono solo alcune delle domande più comuni. Il concetto fondamentale da ricordare è che vi sono diversi dati per ogni risorsa.

Le seguenti sezioni riportano le diverse informazioni inerenti il loro impiego, utili per le risorse più importanti.

2.4.1. Controllo della potenza della CPU

Nella sua forma base, il controllo della potenza della CPU risulta semplice tanto quanto il determinare se l'utilizzo della CPU raggiunge il 100%. Se tale utilizzo rimane al di sotto del 100%, non ha importanza cosa stia facendo il sistema, la potenza di processazione aggiuntiva risulta essere sempre disponibile.

Tuttavia, è molto raro che il sistema non riesca a raggiungere il 100% dell'utilizzo della CPU. A questo punto risulta essere utile esaminare i dati più importanti inerenti l'utilizzo della CPU stessa. Così facendo, diventa possibile determinare dove viene consumata la maggior parte della potenza di processazione. Ecco alcune delle statistiche più comuni sull'utilizzo della CPU:

Utente contro sistema

La percentuale di tempo impiegato per eseguire una processazione del tipo user-level contro una processazione del tipo system-level, esso è in grado di indicare che il caricamento di un sistema è dovuto soprattutto all'esecuzione delle applicazioni o all'overhead del sistema operativo. Alte percentuali user-level tendono ad essere molto positive (assumendo che gli utenti abbiano delle prestazioni soddisfacenti), mentre alte percentuali di system-level, tendono ad evidenziare problemi che necessitano un maggiore controllo.

Cambio di contesto

Un cambio di contesto si può verificare quando la CPU termina l'esecuzione di un processo e inizia l'esecuzione di un altro. A causa della necessità del sistema operativo di assumere il controllo della CPU quando si verifica un cambio di contesto, generalmente si possono verificare contemporaneamente sia diverse variazioni di contesto, che un livello molto alto del consumo della CPU da parte del system-level.

Interruzioni

Come implica il nome, le interruzioni sono situazioni dove la processazione eseguita dalla CPU viene bruscamente modificata. Le suddette interruzioni si verificano generalmente a causa di attività hardware (come ad esempio le interruzioni software che controllano la processazione delle applicazioni). Poiché tali interruzioni vengono affrontate in un system level, una frequenza di interruzione troppo elevata, favorisce un utilizzo maggiore della CPU del system-level.

Processi eseguibili

Un processo può essere presente con un diverso stato. Per esempio, si potrebbe verificare:

- L'attesa che una operazione I/O venga terminata

- Attesa per un sottosistema di gestione della memoria in grado di gestire un page fault

In questi casi, il processo non ha bisogno della CPU.

Tuttavia lo stato del processo potrebbe cambiare, diventando così un processo eseguibile. Come indicato dal nome, un processo eseguibile è quel processo che è in grado di terminare il proprio lavoro quando viene programmata la ricezione della CPU time. Se è possibile eseguire più di un processo in un determinato momento, tutti eccetto uno¹ dei processi eseguibili, devono attendere il proprio turno nella CPU. Controllando il numero di processi eseguibili, è possibile determinare il livello di legame del vostro sistema alla CPU.

Altre metriche che possono avere un impatto sull'utilizzo della CPU, tendono ad includere diversi servizi che il sistema operativo fornisce ai processi. Essi possono includere le statistiche sulla gestione della memoria, la processazione I/O, e così via. Queste statistiche accertano che, durante il controllo delle prestazioni del sistema, non vi sia alcuna relazione tra di esse. In altre parole, le statistiche sull'utilizzo della CPU possono indicare un problema nel sottosistema I/O, oppure quelle inerenti l'utilizzo della memoria potrebbero rivelare un difetto nel design dell'applicazione.

Per questo motivo quando si controllano le prestazioni del sistema, non è possibile esaminare solo una singola statistica; solo esaminando il quadro generale è possibile ottenere informazioni utili.

2.4.2. Controllo della larghezza della banda

Il controllo della larghezza della banda risulta essere molto più difficoltoso di altre risorse finora descritte. Il motivo è dato dal fatto che le statistiche riguardanti le prestazioni, si basano generalmente sul dispositivo mentre molti dei luoghi dove la larghezza della banda è importante, le stesse statistiche risultano essere basate su bus che collegano i dispositivi. In questi esempi, dove più di un dispositivo condivide un bus, potreste essere in grado di notare le statistiche inerenti ogni dispositivo, ma in compenso il carico che ne deriva con i dispositivi posizionato sui bus, sarà maggiore.

Un'altra considerazione che bisogna tener presente durante il controllo della larghezza della banda, è la possibilità che si possano verificare circostanze talidove le statistiche per i dispositivi veri e propri, possono non essere disponibili. Questo può essere particolarmente vero per i bus di espansione del sistema e per i datapath². Tuttavia, anche tali statistiche non possono essere sempre accurate al 100%, ci saranno sempre informazioni sufficienti per avere un elemento di analisi.

Alcune delle statistiche comuni relative alla larghezza della banda sono:

Byte ricevuti/inviati

Le statistiche sull'interfaccia della rete forniscono una indicazione sull'utilizzo della larghezza della banda di uno dei bus più visibili — la rete.

Conteggi dell'interfaccia e velocità

Le statistiche relative alla rete possono fornire una indicazione di eccessive collisioni, trasmettere e ricevere errori e molto altro. Attraverso l'uso di queste statistiche (in modo particolare se queste statistiche sono disponibili per più di un sistema presente sulla rete), è possibile eseguire un minimo di troubleshooting della rete ancor prima di utilizzare i tool di diagnosi.

Trasferimenti al secondo

Raccolti normalmente per i dispositivi I/O a blocco, come ad esempio dischi e tape drive ad elevata prestazione, questa statistica rappresenta un ottimo metodo per determinare se è stato raggiunto il limite di una larghezza di banda di un dispositivo. A causa della loro natura elettromeccanica, il disco ed i tape drive possono solo eseguire un certo numero di operazioni I/O al secondo, la loro prestazione peggiora rapidamente man mano che si raggiunge questo limite.

1. Prendendo in considerazione un sistema con processore singolo.
2. Maggiori informazioni su bus, datapath e sulla larghezza di banda, sono disponibili sul Capitolo 3.

2.4.3. Controllo della memoria

Se vi è una zona dove è possibile trovare delle statistiche riguardanti la prestazione, la suddetta è rappresentata dall'area di controllo sull'utilizzo della memoria. A causa della complessità dei sistemi operativi con memoria virtuale 'demand-paged', le statistiche inerenti l'utilizzo della memoria sono molteplici e varie. È qui che si verifica la maggior parte del lavoro degli amministratori di sistema per la gestione delle risorse.

Le seguenti statistiche rappresentano una panoramica di statistiche comuni per la gestione della memoria:

Page In/Page Out

Queste statistiche rendono possibile valutare il flusso di pagine dalla memoria del sistema ad un dispositivo mass storage collegato (generalmente una unità disco). Un numero elevato per entrambe queste statistiche potrebbe significare che il sistema ha una carenza di memoria fisica, e sta eseguendo un *thrashing*, oppure che utilizza più risorse per muovere le pagine in entrata e in uscita dalla memoria, che nell'eseguire le applicazioni.

Pagine attive/inattive

Queste statistiche mostrano come vengono usate le pagine che risiedono nella memoria. Una carenza di pagine inattive può indicare una carenza di memoria fisica.

Pagine disponibili, condivise, 'buffered' e 'cached'

Queste statistiche sono in grado di fornire migliori informazioni rispetto a quelle riguardanti la pagina inattiva/attiva. Usando le suddette statistiche, è possibile determinare le informazioni generali sull'utilizzo della memoria.

Swap In/Swap Out

Queste statistiche mostrano il comportamento di swapping generale del sistema. Una velocità molto elevata può indicare una carenza di memoria fisica.

Un controllo efficace sull'utilizzo della memoria richiede una buona conoscenza del funzionamento di sistemi operativi con memoria virtuale 'demand-paged'. Anche se con il suddetto argomento da solo si potrebbe scrivere un libro intero, i concetti di base vengono affrontati nel Capitolo 4. Questo capitolo insieme con il tempo necessario a controllare un sistema, vi dà la possibilità di ottenere maggiori informazioni su questo argomento.

2.4.4. Controllo dello storage

Il controllo dello storage viene normalmente eseguito tramite due fasi:

- Il controllo di spazio sufficiente
- Controllo dei problemi sulla prestazione relativi allo storage

È possibile avere problemi in un'area mentre l'altra ne è esente. Per esempio, è possibile avere una unità disco senza spazio ma non avere alcun problema nelle prestazioni. Al contrario è possibile avere 99% di spazio disponibile, ma avere problemi nella prestazione.

Tuttavia, è molto probabile che un sistema normale possa avere problemi sia nell'una che nell'altra area. Per questo motivo, è possibile anche che — in alcuni casi — i problemi che si verificano in un'area possano influenzare anche l'altra. In molti casi i suddetti problemi si presentano sotto forma di prestazioni I/O molto povere, con uno spazio disponibile sull'unità disco vicino allo 0%, anche se in casi di carico I/O estremo, è possibile rallentare I/O fino ad un certo livello dove le applicazioni non possono più essere eseguite correttamente.

In qualunque caso, le seguenti statistiche sono molto utili per il controllo dello storage:

Spazio disponibile

Lo spazio disponibile rappresenta la risorsa più seguita da parte degli amministratori di sistema, è molto raro incontrare un amministratore che non controlla tale risorsa (in alcuni casi tali amministratori hanno un modo automatizzato per eseguire questo controllo).

Statistiche relative al File System

Queste statistiche (come ad esempio il numero dei file/directory, misura media dei file, ecc.), forniscono informazioni aggiuntive su di una singola percentuale di spazio disponibile. I suddetti dati rendono possibile da parte dell'amministratore, una configurazione del sistema in modo da ottenere la migliore prestazione, in quanto il carico I/O imposto da un file system pieno di piccoli file, non è lo stesso di quello imposto da un file system riempito da un singolo grande file.

Trasferimenti al secondo

Questa statistica è il modo idoneo per determinare se i limiti della larghezza di banda di un dispositivo particolare sono stati raggiunti.

Scrittura/lettura al secondo

Leggermente più dettagliate dei trasferimenti al secondo, queste statistiche permettono all'amministratore di sistema di capire maggiormente la natura dei carichi I/O presenti nel dispositivo storage. Questo può essere critico, in quanto alcune tecnologie storage possiedono diverse caratteristiche in prestazione per operazioni di lettura rispetto a quelle di scrittura.

2.5. Informazioni specifiche su Red Hat Enterprise Linux

Red Hat Enterprise Linux presenta diversi tool per il controllo delle risorse. Anche se i suddetti tool sono più numerosi di quelli qui riportati, essi sono i più rappresentativi in termini di funzionalità. I tool sono:

- `free`
- `top` (e **GNOME System Monitor**, una versione più grafica di `top`)
- `vmstat`
- La suite di Sysstat per i tool di controllo delle risorse
- Il profiler generale del sistema di OProfile

Esaminiamoli in dettaglio.

2.5.1. `free`

Il comando `free` mostra l'utilizzo della memoria del sistema. Ecco un esempio del suo output:

	total	used	free	shared	buffers	cached
Mem:	255508	240268	15240	0	7592	86188
-/+ buffers/cache:		146488	109020			
Swap:	530136	26268	503868			

La riga `Mem:` mostra l'utilizzo della memoria fisica, mentre `Swap:` mostra l'utilizzo dello spazio swap del sistema, e `-/+ buffers/cache:` mostra la quantità di memoria fisica riservata ai buffer del sistema.

Poichè `free` per default mostra le informazioni sull'utilizzo della memoria solo una volta, esso è utile solo per un controllo breve, oppure per determinare velocemente se esiste ancora il problema relativo alla memoria stessa. Anche se `free` possiede l'abilità di visualizzare ripetutamente le informazioni

relative all'uso della memoria tramite la sua opzione `-s`, l'output continua a scorrere rendendo difficile il rilevamento di modifiche nell'utilizzo della stessa memoria.



Suggerimento

Una soluzione migliore nell'uso di `free -s` sarebbe quella di eseguire `free` usando il comando `watch`. Per esempio, per visualizzare l'utilizzo della memoria ogni due secondi (il valore di default per `watch`), usare questo comando:

```
watch free
```

Il comando `watch` emette il comando `free` ogni due secondi, eseguendo l'aggiornamento, ripulendo la schermata e scrivendo il nuovo output nella stessa posizione. Questo rende il tutto molto più semplice nel determinare i cambiamenti sull'utilizzo della memoria attraverso un certo periodo di tempo, in quanto `watch` crea una panoramica singola sull'aggiornamento senza poter eseguire uno scorrimento dei dati. È possibile controllare il ritardo tra gli aggiornamenti usando l'opzione `-n`, mentre è possibile evidenziare i cambiamenti utilizzando l'opzione `-d` come mostrato nel seguente comando:

```
watch -n 1 -d free
```

Per maggiori informazioni, consultate la pagina man di `watch`.

Il comando `watch` viene eseguito fino a quando non viene interrotto con [Ctrl]-[C]. Il comando `watch` è da ricordare in quanto potrà esservi molto utile in diverse situazioni.

2.5.2. top

Mentre `free` visualizza solo le informazioni relative alla memoria, il comando `top` fa un pò di tutto. Uso della CPU, statistiche sul processo, utilizzo della memoria — `top` controlla più o meno tutto. Diversamente dal comando `free`, il comportamento di default di `top` è quello di esecuzione continua; non vi è il bisogno di usare il comando `watch`. Eccone un esempio:

```
14:06:32 up 4 days, 21:20, 4 users, load average: 0.00, 0.00, 0.00
77 processes: 76 sleeping, 1 running, 0 zombie, 0 stopped
CPU states:  cpu    user    nice  system    irq  softirq  iowait    idle
              total  19.6%    0.0%    0.0%    0.0%    0.0%    0.0%  180.2%
              cpu00   0.0%    0.0%    0.0%    0.0%    0.0%    0.0%    100.0%
              cpu01   19.6%    0.0%    0.0%    0.0%    0.0%    0.0%    80.3%
Mem:  1028548k av,    716604k used,    311944k free,          0k shrd,    131056k buff
      324996k actv,    108692k in_d,    13988k in_c
Swap: 1020116k av,     5276k used,   1014840k free                382228k cached
```

PID	USER	PRI	NI	SIZE	RSS	SHARE	STAT	%CPU	%MEM	TIME	CPU	COMMAND
17578	root	15	0	13456	13M	9020	S	18.5	1.3	26:35	1	rhnp-applet-gu
19154	root	20	0	1176	1176	892	R	0.9	0.1	0:00	1	top
1	root	15	0	168	160	108	S	0.0	0.0	0:09	0	init
2	root	RT	0	0	0	0	SW	0.0	0.0	0:00	0	migration/0
3	root	RT	0	0	0	0	SW	0.0	0.0	0:00	1	migration/1
4	root	15	0	0	0	0	SW	0.0	0.0	0:00	0	keventd
5	root	34	19	0	0	0	SWN	0.0	0.0	0:00	0	ksoftirqd/0
6	root	35	19	0	0	0	SWN	0.0	0.0	0:00	1	ksoftirqd/1
9	root	15	0	0	0	0	SW	0.0	0.0	0:07	1	bdflush
7	root	15	0	0	0	0	SW	0.0	0.0	1:19	0	kswapd
8	root	15	0	0	0	0	SW	0.0	0.0	0:14	1	kscand
10	root	15	0	0	0	0	SW	0.0	0.0	0:03	1	kupdated
11	root	25	0	0	0	0	SW	0.0	0.0	0:00	0	mdrecoveryd

La schermata viene divisa in due sezioni. Quella superiore contiene le informazioni relative allo stato generale del sistema — l'uptime, il carico medio, i conteggi del processo, lo stato della CPU, e le statistiche sull'utilizzo sia per la memoria che per lo spazio di swap. La sezione inferiore visualizza le statistiche process-level. È possibile cambiare ciò che viene mostrato mentre `top` è in esecuzione. Per esempio, `top` per default visualizza solo i processi idle e non-idle. Per visualizzare i processi non-idle, pigiare [i]; pigiandolo una seconda volta si ritorna nella modalità display di default.



Avvertimento

Anche se `top` appare come un programma di sola visualizzazione, in realtà non lo è. Questo perché `top` utilizza dei comandi a caratteri singoli per eseguire diverse operazioni. Per esempio, se siete registrati come root è possibile modificare la priorità oppure eseguire il kill di qualsiasi processo presente nel sistema. Per questo motivo, fino a quando non avete revisionato la schermata d'aiuto di `top` (digitare [?] per visualizzarla), è più sicuro digitare [q] (il quale abbandona `top`).

2.5.2.1. GNOME System Monitor — Un `top` grafico

Se vi sentite più a vostro agio con le graphical user interface, **GNOME System Monitor** potrebbe fare al caso vostro. Come `top`, **GNOME System Monitor** visualizza le informazioni relative allo stato generale del sistema, ai conteggi del processo, all'utilizzo dello spazio di swap e della memoria e alle statistiche del livello del processo.

Tuttavia, **GNOME System Monitor** si spinge leggermente oltre, includendo anche alcune rappresentazioni grafiche della CPU, della memoria, e sull'utilizzo dello swap, insieme alla voce sull'utilizzo dello spazio tabellare del disco. Un esempio della schermata **Process Listing** di **GNOME System Monitor** appare in Figura 2-1.

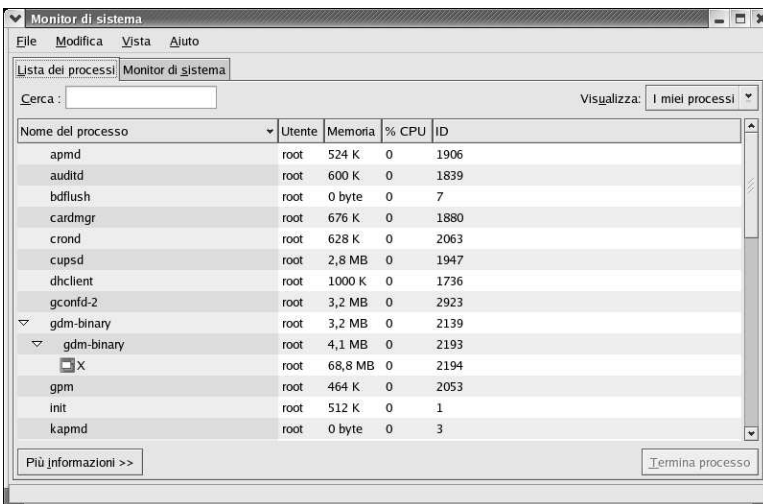


Figura 2-1. Schermata Process Listing di GNOME System Monitor

Si possono visualizzare informazioni aggiuntive per un processo specifico facendo clic sul processo desiderato e poi sul tasto **More Info**.

Per visualizzare la CPU, la memoria e le statistiche sull'utilizzo del disco, fate clic sul pannello **Monitor del sistema**.

2.5.3. vmstat

Per una conoscenza più approfondita sulle prestazioni del sistema, provate `vmstat`. Con `vmstat` è possibile ottenere una panoramica del processo, della memoria, di swap, I/O, del sistema, e sull'attività della CPU attraverso una riga composta da numeri:

procs		memory			swap		io		system			cpu			
r	b	swpd	free	buff	cache	si	so	bi	bo	in	cs	us	sy	id	wa
0	0	5276	315000	130744	380184	1	1	2	24	14	50	1	1	47	0

La prima riga divide i campi in sei categorie, include il processo, la memoria, swap, I/O, il sistema e le statistiche relative alla CPU. La seconda riga identifica maggiormente i contenuti di ogni campo, facilitando un controllo rapido dei dati per specifiche statistiche.

I campi relativi al processo sono:

- `r` — Il numero di processi eseguibili in attesa per l'accesso alla CPU
- `b` — Il numero di processi che sono in uno stato di 'sleep'

I campi relativi alla memoria sono:

- `swpd` — La quantità utilizzata di memoria virtuale
- `free` — La quantità di memoria disponibile
- `buff` — La quantità di memoria usata dai buffer
- `cache` — La quantità di memoria utilizzata come page cache

I campi relativi allo swap sono:

- `si` — La quantità di memoria scambiata in entrata da un disco
- `so` — La quantità di memoria scambiata in uscita dal disco

I campi relativi a I/O sono:

- `bi` — Blocchi inviati ad un dispositivo a blocco
- `bo` — Blocchi ricevuti da un dispositivo a blocco

I campi relativi al sistema sono:

- `in` — Il numero di interruzioni al secondo
- `cs` — Il numero di modifiche del contesto al secondo

I campi relativi alla CPU sono:

- `us` — La percentuale di tempo attraverso il quale la CPU ha eseguito il codice user-level
- `sy` — La percentuale di tempo attraverso il quale la CPU ha eseguito il codice system-level
- `id` — La percentuale di tempo nel quale la CPU è rimasta in posizione di idle
- `wa` — attesa I/O

Quando si esegue `vmstat` senza opzioni viene visualizzata solo una riga. Questa riga contiene informazioni calcolate dall'ultimo avvio del sistema.

Tuttavia molti amministratori di sistema non fanno affidamento ai dati contenuti in questa riga, in quanto i dati vengono raccolti in momenti diversi. Molti amministratori invece usano l'abilità di `vmstat` di visualizzare ripetutamente i dati sull'utilizzo delle risorse a determinati intervalli. Per esempio, il comando `vmstat 1` visualizza una nuova riga ogni secondo, mentre il comando `vmstat 1 10` visualizza una nuova riga al secondo per dieci secondi.

Se usato da un amministratore esperto, `vmstat` può essere utilizzato per determinare velocemente l'uso delle risorse e le problematiche inerenti le prestazioni. Ma per poter ottenere dettagli più specifici su questi argomenti, è necessario un tool diverso — un tool capace di raccogliere dati in modo più approfondito e di condurre analisi.

2.5.4. La suite di Sysstat per i tool di controllo delle risorse

Mentre i tool precedenti possono essere utili per ottenere informazioni più dettagliate sulle prestazioni del sistema in tempi molto brevi, gli stessi sono poco utili nel fornire un quadro corretto sull'utilizzo delle risorse del sistema. In aggiunta, ci sono alcuni aspetti inerenti le prestazioni del sistema che non possono essere controllati usando i suddetti tool.

Per questo motivo è necessario un tool più sofisticato. Sysstat rappresenta tale tool.

Sysstat contiene i seguenti tool relativi alla raccolta di I/O e delle statistiche della CPU:

`iostat`

Visualizza una panoramica sull'utilizzo della CPU, insieme alle statistiche I/O per una o più unità disco.

`mpstat`

Visualizza le statistiche della CPU in modo più dettagliato.

Sysstat contiene anche i tool in grado di raccogliere i dati riguardanti l'uso delle risorse del sistema, e creare dei report giornalieri basati sui dati stessi. Questi tool sono:

`sadc`

Conosciuto come raccoglitore di dati sull'attività del sistema, `sadc` raccoglie le informazioni sull'utilizzo delle risorse del sistema e le scrive su di un file.

`sar`

Eseguendo dei rapporti dai file creati da `sadc`, i rapporti di `sar` possono essere generati in modo interattivo o scritti su di un file per analisi più approfondite.

Le seguenti sezioni analizzano i tool in modo più approfondito.

2.5.4.1. Comando `iostat`

Il comando `iostat` fornisce una panoramica della CPU e delle statistiche I/O su disco:

```
Linux 2.4.20-1.1931.2.231.2.10.ent (pigdog.example.com) 07/11/2003

avg-cpu:  %user   %nice    %sys    %idle
           6.11    2.56    2.15   89.18

Device:            tps   Blk_read/s   Blk_wrtn/s   Blk_read   Blk_wrtn
dev3-0              1.68         15.69         22.42    31175836   44543290
```

Sotto la prima riga (la quale contiene la versione del kernel del sistema e l'hostname, insieme con la data corrente), `iostat` visualizza una panoramica sull'utilizzo medio della CPU del sistema dall'ultimo riavvio. Il rapporto sull'uso della CPU include quanto segue:

- La percentuale di tempo trascorso in modalità utente (esecuzioni delle applicazioni, ecc)
- Percentuale di tempo trascorso in modalità utente (per processi che hanno modificato la loro priorità usando `nice(2)`)
- Percentuale di tempo trascorso in modalità kernel
- Percentuale di tempo trascorso in idle

Sotto il rapporto sull'utilizzo della CPU, si trova il rapporto sull'utilizzo del dispositivo. Il suddetto rapporto contiene una riga per ogni dispositivo del disco attivo sul sistema, e include le seguenti informazioni:

- La specificazione del dispositivo, visualizzato come `dev<major-number>-sequence-number`, dove `<major-number>` è il numero maggiore del dispositivo³, e `<sequence-number>` è una sequenza numerica che parte da zero.
- Il numero di trasferimenti (o di operazioni I/O) al secondo.
- Il numero di blocchi da 512-byte letti al secondo.
- Il numero di blocchi da 512-byte scritti al secondo.
- Il numero totale di blocchi da 512-byte letti.
- Il numero totale di blocchi da 512-byte scritti.

Questo è solo un esempio delle informazioni che si possono ottenere usando `iostat`. Per maggiori informazioni consultare la pagina man di `iostat(1)`.

2.5.4.2. Comando `mpstat`

Il comando `mpstat` potrebbe apparire simile al rapporto sull'utilizzo della CPU prodotto da `iostat`:

```
Linux 2.4.20-1.1931.2.231.2.10.ent (pigdog.example.com) 07/11/2003

07:09:26 PM CPU %user %nice %system %idle intr/s
07:09:26 PM all 6.40 5.84 3.29 84.47 542.47
```

Con l'eccezione di una colonna aggiuntiva che mostra le interruzioni al secondo gestite dalla CPU. Tuttavia, la situazione cambia se viene utilizzata l'opzione `-P ALL` di `mpstat`:

```
Linux 2.4.20-1.1931.2.231.2.10.ent (pigdog.example.com) 07/11/2003

07:13:03 PM CPU %user %nice %system %idle intr/s
07:13:03 PM all 6.40 5.84 3.29 84.47 542.47
07:13:03 PM 0 6.36 5.80 3.29 84.54 542.47
07:13:03 PM 1 6.43 5.87 3.29 84.40 542.47
```

Su sistemi multiprocessori, `mpstat` presenta la possibilità per ogni CPU di essere visualizzata individualmente, permettendo di sapere come è stata usata ogni CPU.

3. I numeri maggiori del dispositivo si possono trovare utilizzando `ls -l`, in modo da visualizzare il file del dispositivo desiderato in `/dev/`. Il numero maggiore in questo esempio appare dopo la specificazione del gruppo del dispositivo.

2.5.4.3. Comando `sadc`

Come precedentemente osservato, il comando `sadc` raccoglie i dati che riguardano l'uso del sistema, e li scrive su di un file per ulteriori analisi. Per default, i dati vengono scritti sui file nella directory `/var/log/sa/`. I file vengono chiamati `sa<dd>`, dove `<dd>` è la data corrente espressa con due caratteri.

`sadc` viene eseguito normalmente dallo script `sa1`. Questo script viene invocato generalmente da `cron` tramite il file `sysstat`, che si trova in `/etc/cron.d/`. Lo script `sa1` richiama `sadc` per un intervallo singolo di un secondo. Per default, `cron` esegue `sa1` ogni 10 minuti, aggiungendo i dati raccolti durante ogni intervallo al file `/var/log/sa/sa<dd>` corrente.

2.5.4.4. Comando `sar`

Il comando `sar` fornisce i rapporti sull'utilizzo del sistema in base ai dati raccolti da `sadc`. Come configurato in Red Hat Enterprise Linux, `sar` viene eseguito automaticamente per processare i file raccolti automaticamente da `sadc`. I file report sono scritti su `/var/log/sa/` e sono chiamati `sar<dd>`, dove `<dd>` è la rappresentazione a due caratteri della data precedente sempre a due caratteri.

`sar` viene eseguito normalmente dallo script `sa2`. Questo script viene invocato periodicamente da `cron` tramite il file `sysstat`, il quale si trova in `/etc/cron.d/`. Per default, `cron` esegue `sa2` una volta al giorno alle 23:53, permettendogli di fornire un report per tutti i dati giornalieri.

2.5.4.4.1. Leggere i rapporti di `sar`

Il formato di un rapporto `sar` fornito dalla configurazione di default di Red Hat Enterprise Linux, consiste in sezioni multiple, con ogni sezione contenente un tipo specifico di dati. Poichè `sadc` è configurato per eseguire un intervallo di un secondo ogni dieci minuti, i report `sar` di default contengono dati che incrementano ogni dieci minuti, da 00:00 a 23:50⁴.

Ogni sezione del rapporto inizia con una testata che descrive i dati contenuti nella sezione stessa. La testata viene ripetuta ad intervalli regolari attraverso la sezione, rendendola più semplice per l'interpretazione dei dati mentre si esegue un 'paging' attraverso il rapporto. Ogni sezione termina con una riga contenente informazioni generali sui dati riportati in quella sezione.

Ecco un esempio di un rapporto della sezione `sar`, con i dati da 00:30 a 23:40 rimossi per ottenere spazio:

00:00:01	CPU	%user	%nice	%system	%idle
00:10:00	all	6.39	1.96	0.66	90.98
00:20:01	all	1.61	3.16	1.09	94.14
...					
23:50:01	all	44.07	0.02	0.77	55.14
Average:	all	5.80	4.99	2.87	86.34

In questa sezione vengono mostrate le informazioni sull'utilizzo della CPU. Questo è molto simile ai dati visualizzati da `iostat`.

Altre sezioni possono avere più di una riga per volta contenente dati, come mostrato in questa sezione generata da dati sull'utilizzo della CPU raccolti su di un sistema 'dual-processor':

00:00:01	CPU	%user	%nice	%system	%idle
00:10:00	0	4.19	1.75	0.70	93.37
00:10:00	1	8.59	2.18	0.63	88.60

4. .A causa dei cambiamenti dei carichi del sistema, il periodo entro il quale i dati sono stati raccolti può variare di uno o due secondi.

00:20:01	0	1.87	3.21	1.14	93.78
00:20:01	1	1.35	3.12	1.04	94.49
...					
23:50:01	0	42.84	0.03	0.80	56.33
23:50:01	1	45.29	0.01	0.74	53.95
Average:	0	6.00	5.01	2.74	86.25
Average:	1	5.61	4.97	2.99	86.43

Ci sono un totale di diciassette sezioni presenti nei rapporti generati dalla configurazione `sar` di default di Red Hat Enterprise Linux; alcuni di essi vengono affrontati nei capitoli successivi. Per maggiori informazioni sui dati contenuti in ogni sezione, consultate la pagina `man` di `sar(1)`.

2.5.5. OProfile

Il profiler generale del sistema di OProfile è un tool di controllo a basso overhead. OProfile fa uso di hardware di controllo delle prestazioni del sistema⁵ per determinare la natura dei problemi relativi alle prestazioni.

L'hardware di controllo delle prestazioni fa parte dello stesso processo. Prende la forma di un contatore speciale, incrementato ogni volta che si verifica un certo evento (come ad esempio se un processo non è in posizione idle oppure i dati richiesti non sono presenti nella cache). Alcuni processori presentano più di un contatore, e permettono di eseguire una selezione di diversi eventi per ogni contatore.

I contatori possono essere caricati con un valore iniziale e produrre una interruzione ogni volta che lo stesso contatore presenta un overflow. Caricando un contatore con diversi valori iniziali, è possibile variare la velocità alla quale vengono prodotte le interruzioni. In questo modo è possibile controllare la velocità campione, e così, il livello dei dettagli ottenuti dai dati raccolti.

Da un lato, impostando il contatore in modo da generare un overflow interrupt con ogni evento, si possono fornire dati sulla prestazione molto dettagliati (ma con un overhead elevatissimo). Al contrario impostare il contatore in modo da generare il minor numero di interruzioni possibile, si fornisce una panoramica generale delle prestazioni del sistema (con nessun overhead). Il segreto di un controllo efficace viene rappresentato dalla selezione di una velocità campione sufficientemente alta in modo da catturare i dati necessari, ma non così alta da far verificare un sovraccarico del sistema con un overhead del controllo delle prestazioni.



Avvertimento

Potete configurare OProfile in modo da produrre un overhead sufficiente per rendere il sistema non più utilizzabile. In questo senso è necessario esercitare molta attenzione nella selezione dei valori del contatore. Per questa ragione, il comando `opcontrol` supporta l'opzione `--list-events`, la quale visualizza il tipo di evento disponibile per il processore attualmente installato, insieme con i valori del contatore minimi suggeriti.

È importante tener sempre presente il rapporto tra la velocità campione e l'overhead quando si utilizza OProfile.

5. OProfile è in grado di usare anche un meccanismo di fallback (conosciuto anche come `TIMER_INT`), per quelle architetture del sistema che presentano delle carenze per quanto riguarda l'hardware di controllo delle prestazioni.

2.5.5.1. Componenti di OProfile

Oprofile consiste dei seguenti componenti:

- Software di raccolta dati
- Software di analisi dei dati
- Software per l'interfaccia amministrativa

Il software per la raccolta dati consiste del modulo del kernel di `oprofile.o`, e del demone `oprofiled`.

Il software per l'analisi dei dati include i seguenti programmi:

`op_time`

Visualizza il numero e le relative percentuali di esempi presi per ogni file eseguibile

`oprofp`

Visualizza il numero e le relative percentuali di esempi presi per la funzione, per le istruzioni individuali, o per output in stile `gprof`

`op_to_source`

Visualizza i codici sorgente annotati e/o dell'assembly listings

`op_visualise`

Visualizza graficamente i dati raccolti

Questi programmi rendono possibile visualizzare i dati raccolti in diversi modi.

Il software per l'interfaccia amministrativa controlla tutti gli aspetti di raccolta dei dati, dalla specificazione dell'evento da controllare, all'avvio e arresto della raccolta dei dati stessi. Tutto questo viene fatto usando il comando `opcontrol`.

2.5.5.2. Esempio della sessione di OProfile

Questa sezione mostra un controllo OProfile e una sessione per l'analisi dei dati da una configurazione iniziale ad un'analisi finale dei dati. Essa rappresenta solo una panoramica introduttiva, per maggiori informazioni consultare *Red Hat Enterprise Linux System Administration Guide*.

Utilizzate `opcontrol` per configurare il tipo di dati da raccogliere con il seguente comando:

```
opcontrol \
--vmlinux=/boot/vmlinux-`uname -r` \
--ctr0-event=CPU_CLK_UNHALTED \
--ctr0-count=6000
```

Le opzioni qui usate dirgono `opcontrol` su:

- Dirigono OProfile su di una copia del kernel corrente in esecuzione (`--vmlinux=/boot/vmlinux-`uname -r``)
- Specificano l'utilizzo del contatore 0 del processore e che l'evento da controllare è il periodo quando la CPU esegue le istruzioni (`--ctr0-event=CPU_CLK_UNHALTED`)
- Specificano che OProfile deve raccogliere gli esempi ogni 6000 volte che si verifica l'evento specificato (`--ctr0-count=6000`)

Successivamente, controllate che il modulo del kernel di `oprofile` venga caricato usando il comando `lsmod`:

```
Module                Size Used by    Not tainted
oprofile              75616   1
...
```

Confermate che il file system di OProfile (posizionato in `/dev/oprofile/`) attraverso il comando `ls /dev/oprofile/`:

```
0 buffer      buffer_watershed  cpu_type  enable      stats
1 buffer_size cpu_buffer_size  dump      kernel_only
```

(Il numero esatto di file varia a seconda del tipo del processo.)

A questo punto, il file `/root/.oprofile/daemonrc` contiene le impostazioni necessarie dal software di raccolta dei dati:

```
CTR_EVENT[0]=CPU_CLK_UNHALTED
CTR_COUNT[0]=6000
CTR_KERNEL[0]=1
CTR_USER[0]=1
CTR_UM[0]=0
CTR_EVENT_VAL[0]=121
CTR_EVENT[1]=
CTR_COUNT[1]=
CTR_KERNEL[1]=1
CTR_USER[1]=1
CTR_UM[1]=0
CTR_EVENT_VAL[1]=
one_enabled=1
SEPARATE_LIB_SAMPLES=0
SEPARATE_KERNEL_SAMPLES=0
VMLINUX=/boot/vmlinux-2.4.21-1.1931.2.349.2.2.entsmp
```

Successivamente, utilizzare `opcontrol` per iniziare la raccolta dei dati con il comando `opcontrol --start`:

```
Using log file /var/lib/oprofile/oprofiled.log
Daemon started.
Profiler running.
```

Verificate che il demone `oprofiled` viene eseguito con il comando `ps x | grep -i oprofiled`:

```
32019 ?      S      0:00 /usr/bin/oprofiled --separate-lib-samples=0 ...
32021 pts/0   S      0:00 grep -i oprofiled
```

(La linea di comando di `oprofiled` visualizzata da `ps` è più lunga; tuttavia, essa è stata troncata per motivi di formattazione.)

Il sistema viene adesso monitorato, con i dati raccolti per tutti gli eseguibili presenti sul sistema. I dati vengono conservati nella directory `/var/lib/oprofile/samples/`. I file in questa directory seguono una convenzione di nomi molto inusuale. Ecco un esempio:

```
}usr}bin}less#0
```

Questa convenzione utilizza il percorso assoluto di ogni file contenente il codice degli eseguibili, con la barra (/) sostituita dalle parentesi graffe (}), e finendo con il simbolo cancelletto (#) seguito da un numero (in questo caso, 0.) Per questo motivo, il file usato in questo esempio rappresenta i dati raccolti mentre `/usr/bin/less` era in esecuzione.

Una volta raccolti i dati utilizzare uno dei tool di analisi per visualizzarli. Un lato positivo di OProfile è rappresentato dal fatto che non è necessario arrestare la raccolta dei dati prima di eseguire un'analisi. Tuttavia, è necessario aspettare la scrittura sul disco di un set di esempi, o usare il comando `opcontrol --dump` per forzare tale procedura.

Nel seguente esempio, `op_time` viene usato per visualizzare (con un ordine inverso — da un numero maggiore di esempi a quello più basso), gli esempi raccolti:

```
3321080 48.8021 0.0000 /boot/vmlinuz-2.4.21-1.1931.2.349.2.2.entsm
761776 11.1940 0.0000 /usr/bin/oprofiled
368933 5.4213 0.0000 /lib/tls/libc-2.3.2.so
293570 4.3139 0.0000 /usr/lib/libgobject-2.0.so.0.200.2
205231 3.0158 0.0000 /usr/lib/libgdk-x11-2.0.so.0.200.2
167575 2.4625 0.0000 /usr/lib/libglib-2.0.so.0.200.2
123095 1.8088 0.0000 /lib/libcrypto.so.0.9.7a
105677 1.5529 0.0000 /usr/X11R6/bin/XFree86
...
```

È buona idea quando si produce un rapporto in modo interattivo, usare `less`, in quanto i suddetti rapporti possono essere lunghi centinaia di righe. L'esempio riportato è stato troncato per questa ragione.

Il formato per questo rapporto in particolare è quello di una riga prodotta per ogni file eseguibile dai quali sono stati presi come esempio. Ogni riga segue questo formato:

```
<sample-count> <sample-percent> <unused-field> <executable-name>
```

Dove:

- `<sample-count>` rappresenta il numero di esempio raccolti
- `<sample-percent>` rappresenta la percentuale di tutti gli esempi raccolti per questo specifico eseguibile
- `<unused-field>` è un campo non in uso
- `<executable-name>` rappresenta il nome del file che contiene il codice dell'eseguibile per il quale sono stati raccolti gli esempi.

Questo rapporto (fornito sulla maggior parte dei sistemi in idle) mostra che quasi la metà di tutti gli esempi sono stati presi in considerazione mentre la CPU esegue il codice all'interno dello stesso kernel. Successivamente nella riga vi è il demone per la raccolta dei dati di OProfile, seguito da diverse librerie e dal server del sistema X Window, XFree86. Vale la pena notare che per il sistema che esegue questa sessione, il valore di 6000 usato dal contatore, rappresenta il valore minimo consigliato da `opcontrol --list-events`. Questo significa che — almeno per questo sistema particolare — l'overhead di OProfile nel peggiore delle ipotesi consuma approssimativamente l'11% della CPU.

2.6. Risorse aggiuntive

Questa sezione include diverse risorse che possono essere usate per ottenere maggiori informazioni sulle risorse usate per il controllo e su specifici argomenti riguardanti Red Hat Enterprise Linux, affrontati in questo capitolo.

2.6.1. Documentazione installata

Le seguenti risorse sono state installate nel corso di una installazione tipica di Red Hat Enterprise Linux

- Pagina man di `free(1)` — Imparate a visualizzare le statistiche riguardanti la memoria utilizzata e quella disponibile.
- Pagina man di `top(1)` — Imparate a visualizzare l'utilizzo della CPU e le statistiche del process-level.
- Pagina man di `watch(1)` — Imparate ad eseguire periodicamente il programma specificato, visualizzando un output completo.
- Entry del menù **Aiuto** di **GNOME System Monitor** — Imparate a visualizzare graficamente il processo, la memoria della CPU, e le statistiche sull'utilizzo dello spazio del disco.
- Pagina man di `vmstat(8)` — Imparate a visualizzare una breve panoramica del processo, della memoria, di swap, I/O, del sistema e dell'utilizzo della CPU.
- Pagina man di `iostat(1)` — Imparate a visualizzare le statistiche della CPU e I/O.
- Pagina man di `mpstat(1)` — Imparate a visualizzare le statistiche individuali della CPU su sistemi multiprocessori.
- Pagina man di `sadc(8)` — Imparate a raccogliere i dati inerenti l'utilizzo del sistema.
- Pagina man di `sal(8)` — Per saperne di più sullo script che esegue periodicamente `sadc`.
- Pagina man di `sar(1)` — Imparate ad eseguire dei riporti sull'utilizzo della risorsa del sistema.
- Pagina man di `sa2(8)` — Imparate a forinisce i file report giornalieri sull'utilizzo della risorsa del sistema.
- Pagina man di `nice(1)` — Imparate a modificare la priorità dello scheduling del processo.
- Pagina man di `oprofile(1)` — Imparate a visualizzare la prestazione del sistema.
- Pagina man di `op_visualise(1)` — Imparate a visualizzare graficamente i dati di OProfile.

2.6.2. Siti web utili

- http://people.redhat.com/alikins/system_tuning.html — Informazioni per il tuning del sistema per i server di Linux. Un avvicinamento al tuning delle prestazioni e al controllo delle risorse per i server.
- <http://www.linuxjournal.com/article.php?sid=2396> — Tool di controllo sulle prestazioni di Linux. Questa pagina si propone più per gli amministratori interessati alla scrittura di una soluzione grafica personalizzata delle prestazioni. Scritto diversi anni fa, è possibile che alcune informazioni non possono essere più applicate, ma il concetto generale è ancora valido.
- <http://oprofile.sourceforge.net/> — Sito web del progetto di OProfile. Include risorse valide per OProfile, incluso gli indicatori alle mailing list ed il canale #oprofile IRC.

2.6.3. Libri correlati

I seguenti libri affrontano le diverse problematiche reletive al controllo delle risorse e rappresentano gli strumenti ideali per gli amministratori del sistema Red Hat Enterprise Linux

- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — include informazioni su diversi tool utilizzati per il controllo delle risorse, incluso OProfile.

- *Linux Performance Tuning and Capacity Planning* di Jason R. Fink and Matthew D. Sherer; Sams — Fornisce alcune panoramiche più dettagliate dei tool di controllo per le risorse finora affrontate, e ne include altri che possono risultare utili per requisiti più specifici.
- *Red Hat Linux Security and Optimization* di Mohammed J. Kabir; Red Hat Press — Approssimativamente composto da 150 pagine, questo libro affronta le problematiche relative alla prestazione. Include alcuni capitoli dedicati alle problematiche riguardanti le prestazioni specifiche della rete, Web, email, e file servers.
- *Linux Administration Handbook* di Evi Nemeth, Garth Snyder, and Trent R. Hein; Prentice Hall — Fornisce un breve capitolo simile in contenuto a questo libro, ma include anche una sezione molto interessante sulla diagnosi di un sistema che ha rallentato improvvisamente le sue prestazioni.
- *Linux System Administration: A User's Guide* di Marcel Gagne; Addison Wesley Professional — Contiene un breve capitolo sul controllo delle prestazioni e sul tuning.

Capitolo 3.

Larghezza di banda e potenza di processazione

Delle due risorse affrontate in questo capitolo, una (la larghezza di banda), risulta spesso essere per gli amministratori di sistema difficile da capire, mentre l'altra (potenza di processazione) risulta essere generalmente un concetto più semplice.

In aggiunta, potrebbe sembrare che queste due risorse non siano così in relazione tra loro — perchè dunque metterle insieme?

La ragione è che le suddette risorse sono basate su hardware che dipende direttamente dalla capacità di un computer di muovere e processare i dati. Per questo motivo essi possono essere correlati tra loro.

3.1. Larghezza di banda

Come definizione di base, la larghezza di banda rappresenta la capacità di trasferire i dati — in altre parole, la quantità di dati che è possibile trasferire da un punto ad un altro in un determinato arco di tempo. Avere una comunicazione dei dati da punto a punto, implica due cose:

- Un insieme di conduttori elettrici usati per creare una comunicazione a basso livello
- Un protocollo per facilitare una efficiente ed affidabile comunicazione dati

Vi sono due tipi di componenti del sistema idonei a questo scopo:

- Bus
- Datapath

La seguente sezione ne affronta in dettaglio le caratteristiche.

3.1.1. Bus

Come precedentemente osservato, i bus abilitano una comunicazione point-to-point e utilizzano una sorta di protocollo, per assicurare che tutte le comunicazioni siano eseguite in modo controllato. Tuttavia, i bus possiedono delle caratteristiche ben definite:

- Caratteristiche elettriche standardizzate (come il numero di conduttori, i livelli di voltaggio, le velocità di segnalazione ecc.)
- Caratteristiche meccaniche standardizzate (come il tipo di connettore, la misura della scheda, la struttura fisica ecc.)
- Protocollo standardizzato

La parola "standardizzato" è importante perchè i bus rappresentano il modo primario con il quale i diversi componenti del sistema sono collegati tra loro.

In molti casi, i bus permettono una interconnessione di hardware prodotto da diversi rivenditori; senza standardizzazione ciò non potrebbe essere possibile. Tuttavia, anche in situazioni dove un bus appartiene ad un produttore, la standardizzazione risulta essere molto importante perchè permette allo stesso produttore una più facile implementazione di componenti diversi tra loro, utilizzando una interfaccia comune — lo stesso bus.

3.1.1.1. Esempi di bus

Non ha importanza dove cerciate, ci saranno sempre dei bus. Ecco i più comuni:

- Bus di mass storage (ATA e SCSI)
- Reti ¹ (Ethernet e Token Ring)
- Bus di memoria (PC133 e Rambus®)
- Bus di espansione (PCI, ISA, USB)

3.1.2. Datapath

I datapath possono essere difficili da identificare, ma come i bus, essi sono dovunque. Anche loro possono permettono una comunicazione point-to-point. Tuttavia, a differenza dei bus, i datapath:

- Utilizzano un protocollo più semplice (se presente)
- Possiedono (se presente) una piccola standardizzazione meccanica

Il motivo per queste differenze è che i datapath sono generalmente interni ad alcuni componenti del sistema, e non vengono utilizzati per facilitare un collegamento ideale tra diversi componenti. Per questo i datapath sono ottimizzati per una situazione particolare, dove sono preferiti velocità e basso costo alla flessibilità per compiti generali più lenta e costosa.

3.1.2.1. Esempi di datapath

Ecco alcuni datapath tipici:

- CPU per un datapath cache sul-chip
- Processori grafici per il datapath per la memoria video

3.1.3. Potenziali problemi relativi alla larghezza di banda

Ci sono due modi attraverso i quali si possono verificare i problemi relativi alla larghezza di banda (sia per i bus che per i datapath):

1. Il bus o il datapath possono rappresentare una risorsa condivisa. In questa situazione, un livello elevato di contesa per il bus, riduce la disponibilità effettiva della larghezza di banda per tutti i dispositivi presenti sul bus.

Un bus SCSI con diverse unità disco attive, può rappresentare un buon esempio. Tali unità possono saturare il bus SCSI, lasciando disponibile una larghezza di banda molto piccola per qualsiasi altro dispositivo presente sullo stesso bus. Il risultato finale è che tutti gli I/O per qualsiasi dispositivo presente su questo bus sono lenti, anche se ogni dispositivo non è molto attivo.

2. Bus o datapath potrebbero rappresentare una risorsa con un numero fisso di dispositivi collegati ad essa. In questo caso, le caratteristiche elettriche del bus (e in alcuni casi la natura del protocollo usato), limitano la larghezza di banda disponibile. Questo generalmente interessa più i datapath che i bus. Ecco perchè gli adattatori grafici tendono a operare più lentamente quando utilizzano risoluzioni più elevate e/o profondità del colore — per ogni refresh della schermata,

1. Invece di un bus del tipo intra-system, le reti possono essere concepite come un bus *inter-system*.

vi è un numero maggiore di dati da indirizzare lungo il datapath che collega la memoria video ed i processori grafici.

3.1.4. Soluzioni potenziali relative alla larghezza di banda

Fortunatamente, i problemi relativi alla larghezza di banda possono essere risolti. Infatti, vi sono diversi approcci da seguire:

- Distribuire il carico
- Ridurre il carico
- Aumentare la capacità

Le seguenti sezioni affrontano tale approccio in modo più dettagliato.

3.1.4.1. Distribuire il carico

Il primo approccio è quello di distribuire in modo omogeneo l'attività del bus. In altre parole, se un bus è sovraccarico e un altro è in una posizione di idle, la situazione potrebbe migliorare se si sposta parte del carico su quest'ultimo bus.

Come amministratori di un sistema, questo rappresenta il primo approccio da considerare, in quanto spesso accade che sono presenti bus aggiuntivi nel vostro sistema. Per esempio, molti PC includono almeno due canali ATA (il quale rappresenta un altro nome per identificare un bus). Se avete due unità disco ATA e due canali ATA, perché entrambe le unità devono essere sullo stesso canale?

Anche se la configurazione del vostro sistema non include bus aggiuntivi, la distribuzione del carico potrebbe sempre rappresentare un approccio molto indicato. Le spese inerenti l'hardware per affrontare tale approccio, dovrebbero essere minori rispetto la sostituzione di un bus esistente con un hardware con una capacità superiore.

3.1.4.2. Ridurre il carico

Come primo impatto, la riduzione del carico e la distribuzione dello stesso possono sembrare parte di due procedure diverse ma simili tra loro. Dopo tutto, con la distribuzione del carico, si ha anche una riduzione dello stesso (almeno sul bus sovraccarico), non vi sembra?

Anche se questo punto di vista è corretto, esso non ha gli stessi effetti di una riduzione del carico in modo *globale*. La chiave è rappresentata dal fatto di poter determinare se vi sono alcuni aspetti inerenti il carico del sistema, che causano il sovraccarico di un bus. Per esempio, se una rete è sovraccarica a causa della presenza di attività necessarie. Forse un piccolo file provvisorio è il destinatario di numerosi segnali I/O lettura/scrittura. Se il suddetto file risiede su di un file server di tipo 'networked', è possibile eliminare una gran parte del traffico di rete, operando in modo locale con il file.

3.1.4.3. Aumento della capacità

La soluzione più ovvia per una larghezza di banda insufficiente è quella di aumentarla in qualche modo. Tuttavia, tale procedura potrebbe risultare costosa. Per aumentare la propria larghezza di banda, il controller SCSI (e probabilmente tutti i dispositivi ad esso collegati), dovrebbe essere sostituito con un hardware più veloce. Se il controller SCSI è una scheda separata, esso rappresenterà un processo piuttosto semplice, ma se il suddetto controller è parte di una scheda madre di un sistema, allora potrebbe essere più difficile giustificare una spesa così onerosa.

3.1.5. In generale...

Tutti gli amministratori di sistema dovrebbero essere a conoscenza della larghezza di banda, e come la configurazione e l'utilizzo del sistema possa influenzarla. Sfortunatamente, non è sempre facile determinare quale sia il problema relativo alla larghezza di banda. Talvolta, il problema non viene rappresentato dal bus di per sé, ma può essere relativo ad un componente ad esso collegato.

Per esempio, considerate un adattatore SCSI collegato ad un bus PCI. Se sono presenti alcuni problemi inerenti la prestazione con l'I/O del disco SCSI, essi potrebbero essere causati da un adattatore SCSI con prestazioni scadenti, anche se i bus PCI e SCSI non sono vicini alle rispettive capacità per la larghezza di banda.

3.2. Potenza di processazione

Spesso conosciuta come potenza della CPU, cicli CPU, e diversi altri nomi, la potenza di processazione rappresenta l'abilità di un computer di manipolare i dati. Tale potenza varia a seconda dell'architettura (e della velocità dell'orologio) della CPU — generalmente le CPU con una velocità di orologio più elevata e quelle che supportano un numero di dati più elevati, possiedono una potenza di processazione più elevata rispetto alle CPU più lente che supportano un numero minore di dati.

3.2.1. Dati relativi alla potenza di processazione

Ecco due dei dati principali inerenti la potenza di processazione che dovrete tener presente:

- La potenza di processazione è fissa
- La potenza di processazione non può essere conservata

La potenza di processazione è fissa, e cioè la CPU non può superare un certo limite. Per esempio, se avete bisogno di aggiungere due numeri insieme (una operazione che necessita di una sola istruzione sulla maggior parte delle architetture), una CPU particolare può eseguire questa operazione ad una sola velocità. Con qualche eccezione, non è possibile *rallentare* la velocità con la quale una CPU processa le istruzioni, tanto meno aumentarla.

La potenza di processazione è anche limitata. Cioè, ci sono limiti dovuti al tipo di CPU da collegare ad ogni computer. Alcuni sistemi sono capaci di supportare un'ampia gamma di CPU con diverse velocità, mentre altri potrebbero non essere aggiornabili².

La potenza di processazione non può essere conservata per un uso futuro. In altre parole, se una CPU è in grado di processare 100 milioni d'istruzioni al secondo, un secondo in posizione idle, risulta essere uguale a 100 milioni d'istruzioni che non sono state processate.

Se teniamo presenti questi fattori e li esaminiamo attraverso una prospettiva diversa, una CPU "produce" una serie di istruzioni eseguite ad una velocità fissa. E se la CPU "produce" alcune istruzioni precedentemente eseguite, ciò significa che vi è qualcosa che le "consuma". La sezione successiva affronta tale aspetto.

3.2.2. Utente interessato all'uso della potenza di processazione

Sono presenti due utenze principali per la potenza di processazione:

- Le applicazioni

2. Questa situazione porta ad una situazione conosciuta come *forklift upgrade*, e cioè ad una sostituzione completa del computer.

- Lo stesso sistema operativo

3.2.2.1. Le applicazioni

Le normali utenze della potenza di processazione sono le applicazioni ed i programmi che desiderate far eseguire dal computer. Da uno spreadsheet ad un database, le applicazioni sono la ragione per la quale avete un computer.

Un sistema con una sola CPU è in grado di eseguire un compito per volta. Per questo, se la vostra applicazione è in esecuzione, non è possibile eseguire altro. È anche vero che — se state eseguendo qualcosa di diverso da un'applicazione, allora la stessa non può essere eseguita.

Come mai è possibile eseguire diverse applicazioni contemporaneamente su sistemi operativi moderni? La risposta è che essi sono sistemi operativi del tipo 'multitasking'. In altre parole, essi sono in grado di creare l'illusione che è possibile eseguire diverse cose contemporaneamente, quando in effetti ciò non può essere fatto. Il segreto è quello di conferire ad ogni processo una frazione di secondo dove il processo stesso è in grado di essere eseguito sulla CPU, prima di rendere disponibile quest'ultima ad un altro processo nella frazione di secondo successiva. Se questi *cambi di contesto* avvengono molto frequentemente, è possibile raggiungere l'illusione di eseguire applicazioni multiple contemporaneamente.

Naturalmente, le applicazioni eseguono diversi altri compiti oltre a quello di manipolare i dati usando la CPU. Esse possono attendere l'input di un utente o eseguire l'I/O per dispositivi come le unità disco ed i display grafici. Quando si verificano questi eventi, l'applicazione non ha più bisogno della CPU. A questo punto la CPU stessa può essere utilizzata per altri processi in grado di eseguire altre applicazioni senza rallentare l'applicazione in attesa.

In aggiunta, la CPU può essere utilizzata da un'altra utenza della potenza di processazione: il sistema operativo.

3.2.2.2. Il sistema operativo

È difficile determinare il consumo della potenza di processazione da parte del sistema operativo. Il motivo è rappresentato dal fatto che i sistemi operativi utilizzano un mix di codici system-level e process-level, per svolgere il loro lavoro. Mentre per esempio, è più semplice utilizzare un monitor per visualizzare l'attività del processo che esegue un *demone* o un *servizio*, non risulta essere altrettanto semplice determinare la quantità di potenza di processazione consumata dalla processazione relativa all'I/O del system-level (la quale viene determinata normalmente all'interno del contesto del processo che richiede l'I/O.)

In generale, è possibile dividere questo tipo di overhead del sistema operativo in due tipi:

- Gestione del sistema operativo
- Attività relative al processo

La gestione del sistema operativo include l'attività di programmazione del processo e della gestione della memoria, mentre le attività relative al processo includono qualsiasi processo che supporta il sistema operativo, ad esempio un processo che gestisce l'event logging dell'intero sistema oppure il cache flushing I/O.

3.2.3. Migliorare la carenza di CPU

Là dove non è disponibile una potenza di processazione sufficiente per il lavoro da svolgere è possibile:

- Ridurre il carico

- Aumentare la capacità

3.2.3.1. Riduzione del carico

La riduzione del carico della CPU è un approccio che si può intraprendere senza alcuna spesa. Il segreto è quello di identificare gli aspetti del carico del sistema sotto il vostro controllo, che possono essere eliminati. A questo scopo vi suggeriamo di fare particolare attenzione a tre aree principali:

- Riduzione dell'overhead del sistema operativo
- Riduzione dell'overhead dell'applicazione
- Eliminazione completa delle applicazioni

3.2.3.1.1. Riduzione dell'overhead del sistema operativo

Per ridurre l'overhead del sistema operativo, è necessario esaminare l'attuale carico del sistema e determinare gli aspetti che danno luogo ad un numero elevato di overhead. Queste aree possono includere:

- Riduzione della necessità di una programmazione frequente del processo
- Riduzione della quantità I/O eseguita

Non aspettatevi miracoli, in un sistema configurato in modo corretto, è improbabile notare una migliore prestazione riducendo l'overhead del sistema operativo stesso. Questo è dovuto dal fatto che il suddetto sistema, per definizione, presenta una minima quantità di overhead. Tuttavia, se il vostro sistema è in esecuzione con poca RAM, sareste in grado di ridurre l'overhead alleviando la carenza della stessa RAM.

3.2.3.1.2. Riduzione dell'overhead dell'applicazione

La riduzione dell'overhead dell'applicazione significa accertarsi che l'applicazione stessa abbia tutto il necessario per essere eseguita in modo corretto. Alcune applicazioni hanno mostrato diversi comportamenti sotto diversi ambienti — un'applicazione può richiedere molte risorse quando processa alcuni tipi di dati, ma comportarsi in modo diverso quando ne processa altri.

È quindi necessario comprendere tutte le applicazioni eseguite sul vostro sistema per poterle eseguire più efficientemente. Spesso questo procedimento comporta lavorare a stretto contatto con i vostri utenti e/o con gli sviluppatori dell'organizzazione, in modo da scoprire i diversi modi attraverso i quali le applicazioni possono essere meglio eseguite.

3.2.3.1.3. Eliminazione completa delle applicazioni

A seconda dell'organizzazione, questo approccio potrebbe non essere a voi disponibile in quanto spesso non è responsabilità dell'amministratore del sistema decidere quali applicazioni devono essere eseguite. Tuttavia, se siete in grado di identificare qualsiasi applicazione conosciuta come "CPU hogs", allora sareste in grado di influenzare il loro ritiro.

Potreste non essere i soli ad essere interessati seguendo questa procedura. Anche gli utenti coinvolti dovrebbero far parte di questo processo; in molti casi essi potrebbero avere la conoscenza ed il potere per eseguire i cambiamenti necessari all'applicazione.

**Suggerimento**

Ricordatevi che non è necessario eliminare un'applicazione da ogni sistema presente nella vostra organizzazione. Infatti potreste essere in grado di muovere una particolare applicazione particolarmente 'affamata' di CPU, da un sistema sovraccarico ad un altro sistema che è quasi in una posizione di idle.

3.2.3.2. Aumento della capacità

Ovviamente, se non risulta essere possibile ridurre la domanda di potenza di processazione, allora dovrete cercare il modo di aumentare la potenza di processazione disponibile. Per fare ciò è necessario l'uso di risorse economiche.

3.2.3.2.1. Migliorare la CPU

L'approccio più semplice è quello di determinare se la CPU del vostro sistema può essere migliorata. La prima fase è quella di determinare se l'attuale CPU può essere rimossa. Alcuni sistemi (in modo particolare i portatili) possiedono CPU in grado di essere integrate, rendendo così impossibile l'esecuzione di un miglioramento. Il resto dei sistemi possiedono delle CPU 'socketed' che rendono possibili, almeno in teoria, l'esecuzione di tale miglioramento.

Successivamente, dovete eseguire alcune ricerche per determinare se esiste una CPU più veloce per la configurazione del vostro sistema. Per esempio, se avete una CPU di 1GHz, ed esiste una unità di 2GHz dello stesso tipo, allora è possibile eseguire un miglioramento.

Per finire, dovete determinare la velocità massima dell'orologio che può essere supportata dal vostro sistema. Per continuare l'esempio sopra indicato, se il vostro sistema supporta solo processori che vengono eseguiti con un valore di 1GHz o inferiore, anche se esiste una tipologia corretta di CPU di 2GHz, non è possibile eseguire una semplice sostituzione della stessa.

Se non riuscite a installare nel vostro sistema una CPU più veloce, allora la vostra opzione sarà limitata alla sostituzione della scheda madre oppure all'adozione del 'forklift upgrade', cioè della sostituzione dell'intero computer, ma questa procedura risulta essere molto più costosa.

Tuttavia, alcune configurazioni rendono possibile un approccio leggermente diverso. Invece di sostituire la CPU corrente, perchè non aggiungerne un'altra?

3.2.3.2.2. La Multiprocessazione Simmetrica è l'opzione migliore per voi?

La multiprocessazione simmetrica (conosciuta anche come SMP), rende possibile per un sistema computerizzato, avere più di una CPU in grado di condividere tutte le risorse del sistema. Ciò significa che, a differenza di un sistema uniprocessor, un sistema SMP potrebbe presentare più di un processo in esecuzione contemporaneamente.

Questo potrebbe sembrare come il sogno principale di un amministratore di sistema. Diciamo innanzitutto che SMP rende possibile aumentare la potenza della CPU di un sistema, anche se non sono disponibili CPU con una maggiore velocità dell'orologio, — aggiungendo soltanto un'altra CPU. Tuttavia questa flessibilità presenta delle avvertenze.

La prima è che tutti i sistemi sono in grado di supportare operazioni SMP. Il vostro sistema deve presentare una scheda madre creata per supportare processori multipli. Se non li supporta, è necessario (almeno) eseguire un miglioramento della stessa scheda madre.

La seconda è che SMP aumenta l'overhead del sistema. Con più lavoro da programmare per le CPU, il sistema operativo richiede un maggior numero di cicli CPU per l'overhead. Un altro aspetto è quello che con CPU multiple, è possibile il verificarsi di maggiori dispute per le risorse del sistema. A causa

di questi fattori, il miglioramento di un sistema dual-processor in una unità quad-processor, non corrisponde ad un aumento del 100% in potenza disponibile della CPU. Infatti, a seconda dell'hardware corrente, del carico di lavoro e dell'architettura del processore, è possibile raggiungere un punto dove l'aggiunta di un altro processore potrebbe *ridurre* le prestazioni del sistema.

Un altro punto da considerare è che SMP non aiuta i carichi di lavoro che consistono in un'applicazione monolitica con un singolo livello di esecuzione. In altre parole, se un programma di simulazione compute-bound molto grande viene eseguito come un processo senza thread, esso non verrà eseguito più velocemente su di un sistema SMP rispetto ad una macchina con un processore singolo. Infatti, potrebbe essere anche più lento, a causa di maggiori overhead dovuti all'SMP. Per queste ragioni molti amministratori di sistema convengono nel fatto che quando si tratta di potenza CPU, la potenza di processazione con un singolo livello rappresenta la scelta migliore. Tale scelta fornisce una maggiore potenza CPU con minore restrizione.

Anche se gli argomenti affrontati sembrerebbero portare all'idea che SMP non è mai la scelta giusta, vi sono alcune circostanze nelle quali SMP può risultare la scelta più idonea. Per esempio, ambienti in grado di eseguire molteplici applicazioni compute-bound rappresentano il caso ideale per SMP. La ragione per questo è data dal fatto che le applicazioni che non fanno altro che eseguire calcoli per lunghi periodi, mantengono le contese presenti tra i processi attivi (e quindi, l'overhead del sistema operativo) ad un livello minimo, mentre gli stessi processi mantengono ogni CPU occupata.

Un'altra cosa da ricordare riguardante l'SMP, è che la prestazione di un sistema SMP tende a 'peggiore' in modo più omogeneo con l'aumentare del carico di lavoro del sistema. Questa caratteristica rende i sistemi SMP molto più diffusi negli ambienti server e multi-utente, in quanto i diversi processi possono avere un minor impatto sul carico generale del sistema su di una macchina a processori multipli.

3.3. Informazioni specifiche su Red Hat Enterprise Linux

Il controllo della larghezza di banda e l'impiego della CPU sotto Red Hat Enterprise Linux, comporta l'uso dei tool affrontati nel Capitolo 2; quindi, se non avete ancora letto il capitolo in questione, fatelo prima di continuare.

3.3.1. Controllo della larghezza di banda su Red Hat Enterprise Linux

Come accennato nella Sezione 2.4.2, è difficile controllare direttamente l'utilizzo della larghezza di banda. Tuttavia, esaminando le statistiche inerenti il livello del dispositivo, è possibile in modo approssimativo sapere se l'insufficienza della larghezza di banda rappresenti un problema per il vostro sistema.

Con `vmstat` è possibile determinare se l'attività generale del dispositivo è eccessiva, controllando i campi `bi` e `bo`, in aggiunta, prendendo nota dei campi `si` e `so`, è possibile ottenere maggiori informazioni sull'attività del disco causata dallo swap I/O:

procs			memory			swap		io		system			cpu		
r	b	w	swpd	free	buff	cache	si	so	bi	bo	in	cs	us	sy	id
1	0	0	0	248088	158636	480804	0	0	2	6	120	120	10	3	87

In questo esempio, il campo `bi` mostra due blocchi/secondo scritti sui dispositivi a blocco (principalmente sulle unità disco), mentre il campo `bo` mostra sei blocchi/secondo letti dai dispositivi a blocco. Possiamo determinare che nessuna di queste attività è stata causata dal processo di swap, in quanto i campi `si` e `so` mostrano entrambi una velocità relativa allo swap I/O pari a 0 kilobytes/secondo.

Utilizzando `iostat`, è possibile ottenere maggiori informazioni sull'attività relativa al disco:

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (raptor.example.com) 07/21/2003

avg-cpu:  %user   %nice    %sys    %idle
           5.34     4.60     2.83    87.24

Device:            tps   Blk_read/s   Blk_wrtn/s   Blk_read   Blk_wrtn
dev8-0              1.10         6.21        25.08     961342    3881610
dev8-1              0.00         0.00         0.00         16         0
```

Questo output mostra che il dispositivo con il numero 8 (in questo caso `/dev/sda`, il primo disco SCSI) ha una media leggermente maggiore ad una operazione I/O al secondo (il campo `tps`). La maggior parte dell'attività I/O per questo dispositivo è stata scritta (il campo `Blk_wrtn/s`) con una velocità di poco superiore a 25 blocchi al secondo (il campo `Blk_wrtn/s`).

Se avete bisogno di maggiori informazioni, utilizzate l'opzione `-x` di `iostat`:

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (raptor.example.com) 07/21/2003

avg-cpu:  %user   %nice    %sys    %idle
           5.37     4.54     2.81    87.27

Device:            rrqm/s   wrqm/s     r/s     w/s   rsec/s   wsec/s   rkB/s   kB/s   avgrq-sz
/dev/sda          13.57     2.86     0.36   0.77   32.20    29.05    16.10   14.53   54.52
/dev/sda1          0.17     0.00     0.00   0.00     0.34     0.00     0.17     0.00  133.40
/dev/sda2          0.00     0.00     0.00   0.00     0.00     0.00     0.00     0.00   11.56
/dev/sda3          0.31     2.11     0.29   0.62     4.74    21.80     2.37    10.90   29.42
/dev/sda4          0.09     0.75     0.04   0.15     1.06     7.24     0.53     3.62   43.01
```

La prima cosa da ricordare è che l'output `iostat` mostra ora delle statistiche suddivise in livelli. Utilizzando `df` per associare i mount point ai nomi del dispositivo, è possibile usare questo rapporto per determinare se, per esempio, la partizione che contiene `/home/` presenta un carico di lavoro eccessivo.

Ogni output proveniente da `iostat -x` è lungo e contiene più informazioni rispetto a questo; ecco il remainder di ogni riga (con l'aggiunta di una colonna per facilitarne la lettura):

```
Device:      avgqu-sz   await   svctm   %util
/dev/sda      0.24     20.86    3.80    0.43
/dev/sda1     0.00    141.18   122.73    0.03
/dev/sda2     0.00      6.00     6.00    0.00
/dev/sda3     0.12    12.84     2.68    0.24
/dev/sda4     0.11    57.47     8.94    0.17
```

In questo esempio, è interessante notare che `/dev/sda2` è la partizione di swap del sistema; è evidente dai molteplici campi che presentano `0.00` per questa partizione, che lo swapping non rappresenta un problema su questo sistema.

Un altro punto molto interessante da notare è `/dev/sda1`. Le statistiche per questa partizione sono molto inusuali; l'attività generale sembrerebbe bassa, ma per quale motivo la misura media della richiesta I/O (il campo `avgrq-sz`), il tempo medio d'attesa (il campo `await`), ed il tempo medio di servizio (il campo `svctm`) sono più grandi delle altre partizioni? La risposta è che questa partizione contiene la directory `/boot/`, la quale rappresenta il luogo dove viene conservata la ramdisk iniziale ed il kernel. Quando il sistema esegue l'avvio, gli I/O di lettura (da notare che solo i campi `rsec/s` e `rkB/s` non presentano il valore zero; qui non viene eseguita alcuna procedura di scrittura) usati durante il processo d'avvio sono per un alto numero di blocchi, risultando così in un'attesa relativamente lunga ed in visualizzazioni `iostat` dei tempi di servizio.

È possibile utilizzare `sar` per una panoramica più duratura delle statistiche I/O; per esempio, `sar -b` visualizza un rapporto generale di I/O:

Linux 2.4.21-1.1931.2.349.2.2.entsmp (raptor.example.com) 07/21/2003

12:00:00 AM	tps	rtps	wtps	bread/s	bwrtn/s
12:10:00 AM	0.51	0.01	0.50	0.25	14.32
12:20:01 AM	0.48	0.00	0.48	0.00	13.32
...					
06:00:02 PM	1.24	0.00	1.24	0.01	36.23
Average:	1.11	0.31	0.80	68.14	34.79

Le statistiche, come il display iniziale di `iostat`, sono raggruppate per tutti i dispositivi a blocco.

Un altro rapporto relativo a I/O viene riprodotto utilizzando `sar -d`:

Linux 2.4.21-1.1931.2.349.2.2.entsmp (raptor.example.com) 07/21/2003

12:00:00 AM	DEV	tps	sect/s
12:10:00 AM	dev8-0	0.51	14.57
12:10:00 AM	dev8-1	0.00	0.00
12:20:01 AM	dev8-0	0.48	13.32
12:20:01 AM	dev8-1	0.00	0.00
...			
06:00:02 PM	dev8-0	1.24	36.25
06:00:02 PM	dev8-1	0.00	0.00
Average:	dev8-0	1.11	102.93
Average:	dev8-1	0.00	0.00

Questo rapporto fornisce informazioni per ogni dispositivo, ma con pochi dettagli.

Anche se non vi sono delle statistiche precise in grado di mostrare l'utilizzo della larghezza di banda o del datapath, noi possiamo determinare almeno il comportamento del dispositivo e usare la loro attività per determinare indirettamente il carico del bus.

3.3.2. Controllo dell'utilizzo della CPU su Red Hat Enterprise Linux

A differenza della larghezza di banda, il controllo dell'utilizzo della CPU è molto più semplice. Da una percentuale singola di utilizzo della CPU in **GNOME System Monitor**, alle statistiche più dettagliate riportate da `sar`, è possibile determinare in modo più accurato la quantità di potenza della CPU consumata e da chi.

Andando oltre **GNOME System Monitor**, `top` rappresenta il primo tool di controllo delle risorse nel Capitolo 2 in grado di fornire una rappresentazione più dettagliata sull'uso della CPU. Ecco un rapporto `top` di una workstation dual-processor:

9:44pm up 2 days, 2 min, 1 user, load average: 0.14, 0.12, 0.09										
90 processes: 82 sleeping, 1 running, 7 zombie, 0 stopped										
CPU0 states: 0.4% user, 1.1% system, 0.0% nice, 97.4% idle										
CPU1 states: 0.5% user, 1.3% system, 0.0% nice, 97.1% idle										
Mem: 1288720K av, 1056260K used, 232460K free, 0K shrd, 145644K buff										
Swap: 522104K av, 0K used, 522104K free 469764K cached										
PID	USER	PRI	NI	SIZE	RSS	SHARE	STAT	%CPU	%MEM	TIME COMMAND
30997	ed	16	0	1100	1100	840	R	1.7	0.0	0:00 top
1120	root	5	-10	249M	174M	71508	S <	0.9	13.8	254:59 X
1260	ed	15	0	54408	53M	6864	S	0.7	4.2	12:09 gnome-terminal
888	root	15	0	2428	2428	1796	S	0.1	0.1	0:06 sendmail
1264	ed	15	0	16336	15M	9480	S	0.1	1.2	1:58 rhn-applet-gui
1	root	15	0	476	476	424	S	0.0	0.0	0:05 init
2	root	0K	0	0	0	0	SW	0.0	0.0	0:00 migration_CPU0
3	root	0K	0	0	0	0	SW	0.0	0.0	0:00 migration_CPU1

```

4 root      15   0   0   0   0 SW    0.0  0.0  0:01 keventd
5 root      34  19   0   0   0 SWN   0.0  0.0  0:00 ksoftirqd_CPU0
6 root      34  19   0   0   0 SWN   0.0  0.0  0:00 ksoftirqd_CPU1
7 root      15   0   0   0   0 SW    0.0  0.0  0:05 kswapd
8 root      15   0   0   0   0 SW    0.0  0.0  0:00 bdflush
9 root      15   0   0   0   0 SW    0.0  0.0  0:01 kupdated
10 root     25   0   0   0   0 SW    0.0  0.0  0:00 mdrecoveryd

```

Le prime informazioni riguardanti la CPU vengono visualizzate sulla prima riga: il carico medio. Il carico medio è un numero che corrisponde al numero medio di processi eseguibili sul sistema. Il carico medio viene spesso elencato attraverso un insieme di tre numeri (come fatto da `top`), i quali rappresentano il carico medio nei precedenti 1, 5 e 15 minuti, indicando che il sistema in questo esempio, non era molto occupato.

La riga successiva, anche se non strettamente legata all'utilizzo della CPU, ha una relazione indiretta, in quanto mostra il numero di processi eseguibili (qui, solo uno -- ricordate questo numero, in quanto rappresenta qualcosa di speciale in questo esempio). Il numero di processi eseguibili è un buon indicatore di come potrebbe essere il legame tra la CPU ed il sistema.

Successivamente si possono notare due righe che mostrano l'utilizzo corrente dei due CPU nel sistema. Le statistiche inerenti l'utilizzo mostrano il consumo dei cicli CPU, e se lo stesso sia dovuto alla processazione user-level o system-level; viene altresì inclusa una statistica che mostra la quantità di tempo CPU consumato dai processi che presentano uno scheduling di priorità modificato. Per finire, vi è anche una statistica riguardante il tempo in posizione idle.

Procedendo nella sezione relativa al processo, possiamo trovare che il processo che utilizza maggiormente la potenza CPU è proprio `top`; in altre parole, l'unico processo eseguibile era `top` ed era intento ad eseguire una "foto" di se stesso.



Suggerimento

È importante ricordare che l'atto di esecuzione di un monitor del sistema, influenza le statistiche sull'utilizzo della risorsa che ricevete. Tutti i monitor basati sul software, in alcuni casi, non sono esenti da quanto sopra indicato.

Per ottenere una conoscenza più dettagliata sull'utilizzo della CPU, è necessario cambiare i tool. Se esaminiamo l'output di `vmstat`, possiamo avere una concezione leggermente diversa sul nostro sistema:

```

      procs          memory      swap      io      system      cpu
r  b  w  swpd  free  buff  cache  si  so  bi  bo  in  cs  us  sy  id
1  0  0      0 233276 146636 469808  0  0   7   7  14  27  10  3  87
0  0  0      0 233276 146636 469808  0  0   0   0 523 138  3  0  96
0  0  0      0 233276 146636 469808  0  0   0   0 557 385  2  1  97
0  0  0      0 233276 146636 469808  0  0   0   0 544 343  2  0  97
0  0  0      0 233276 146636 469808  0  0   0   0 517  89  2  0  98
0  0  0      0 233276 146636 469808  0  0   0  32 518 102  2  0  98
0  0  0      0 233276 146636 469808  0  0   0   0 516  91  2  1  98
0  0  0      0 233276 146636 469808  0  0   0   0 516  72  2  0  98
0  0  0      0 233276 146636 469808  0  0   0   0 516  88  2  0  97
0  0  0      0 233276 146636 469808  0  0   0   0 516  81  2  0  97

```

In questo caso abbiamo usato il comando `vmstat 1 10` per controllare il sistema ogni secondo per dieci volte. Le statistiche relative alla CPU (i campi `us`, `sy`, e `id`) sembrano simili a quelle mostrate da `top`, e addirittura possono apparire anche meno dettagliate. Tuttavia, a differenza di `top`, possiamo anche ottenere alcune informazioni su come è stata usata la CPU.

Se esaminiamo i campi `system` possiamo notare che la CPU in media, gestisce 500 interruzioni al secondo e si smista tra i processi da 80 a 400 volte al secondo. Se vi pare che essa sia un'attività intensa, ricredetevi, in quanto la processazione user-level (il campo `us`) si comporta solo al 2%, mentre la processazione system-level (il campo `sy`) è generalmente sotto l'1%. Ancora, questo rappresenta un sistema idle.

Ricontrollando ciò che offrono i tool Sysstat, possiamo dire che `iostat` e `mpstat` forniscono poche informazioni aggiuntive a confronto di ciò che è stato ottenuto con `top` e `vmstat`. Tuttavia, `sar` produce un numero di rapporti che possono essere utili durante il controllo della CPU.

Il primo rapporto viene ottenuto dal comando `sar -q`, il quale visualizza la lunghezza della coda di esecuzione, il numero totale dei processi, ed il carico medio inerenti ad uno e cinque minuti. Ecco un esempio:

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (falcon.example.com)      07/21/2003

12:00:01 AM      runq-sz    plist-sz    ldavg-1    ldavg-5
12:10:00 AM           3         122        0.07       0.28
12:20:01 AM           5         123        0.00       0.03
...
09:50:00 AM           5         124        0.67       0.65
Average:           4         123        0.26       0.26
```

In questo esempio, il sistema è sempre occupato (se consideriamo che è possibile eseguire più di un processo in ogni determinato momento), ma non è sovraccarico (grazie al fatto che questo sistema possiede più di un processore).

Il prossimo rapporto `sar` relativo alla CPU viene fornito dal comando `sar -u`:

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (falcon.example.com)      07/21/2003

12:00:01 AM      CPU      %user      %nice      %system      %idle
12:10:00 AM      all       3.69       20.10       1.06       75.15
12:20:01 AM      all       1.73       0.22       0.80       97.25
...
10:00:00 AM      all      35.17       0.83       1.06       62.93
Average:      all       7.47       4.85       3.87       83.81
```

Le statistiche contenute in questo rapporto non variano rispetto a quelle fornite da molti altri tool. Il beneficio più grande è che `sar` è in grado di rendere i dati continuamente disponibili, e quindi, si mostra più utile per ottenere delle medie valide più a lungo, oppure potrebbe risultare utile per la produzione di grafici sull'utilizzo della CPU.

Su sistemi con processori multipli, il comando `sar -U` può fornire statistiche per un unico o per tutti i processori. Ecco un esempio di output di `sar -U ALL`:

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (falcon.example.com)      07/21/2003

12:00:01 AM      CPU      %user      %nice      %system      %idle
12:10:00 AM      0        3.46       21.47       1.09       73.98
12:10:00 AM      1        3.91       18.73       1.03       76.33
12:20:01 AM      0        1.63       0.25       0.78       97.34
12:20:01 AM      1        1.82       0.20       0.81       97.17
...
10:00:00 AM      0      39.12       0.75       1.04       59.09
10:00:00 AM      1      31.22       0.92       1.09       66.77
Average:      0        7.61       4.91       3.86       83.61
Average:      1        7.33       4.78       3.88       84.02
```

Il comando `sar -w` riporta il numero di cambiamenti del contesto al secondo, rendendo possibile ottenere informazioni aggiuntive su dove siano stati usati i cicli CPU:

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (falcon.example.com)      07/21/2003

12:00:01 AM      cswch/s
12:10:00 AM      537.97
12:20:01 AM      339.43
...
10:10:00 AM      319.42
Average:         1158.25
```

È anche possibile fornire due diversi rapporti `sar` sull'attività d'interruzione. Il primo, (riprodotto usando il comando `sar -I SUM`) visualizza una singola statistica inerente "le interruzioni al secondo":

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (falcon.example.com)      07/21/2003

12:00:01 AM      INTR      intr/s
12:10:00 AM      sum      539.15
12:20:01 AM      sum      539.49
...
10:40:01 AM      sum      539.10
Average:         sum      541.00
```

Utilizzando il comando `sar -I PROC`, è possibile suddividere l'attività d'interruzione a seconda del processore (su sistemi con processori multipli) e tramite il livello d'interruzione (da 0 a 15):

```
Linux 2.4.21-1.1931.2.349.2.2.entsmp (pigdog.example.com)      07/21/2003

12:00:00 AM      CPU      i000/s  i001/s  i002/s  i008/s  i009/s  i011/s  i012/s
12:10:01 AM      0      512.01    0.00    0.00    0.00    3.44    0.00    0.00

12:10:01 AM      CPU      i000/s  i001/s  i002/s  i008/s  i009/s  i011/s  i012/s
12:20:01 AM      0      512.00    0.00    0.00    0.00    3.73    0.00    0.00
...
10:30:01 AM      CPU      i000/s  i001/s  i002/s  i003/s  i008/s  i009/s  i010/s
10:40:02 AM      0      512.00    1.67    0.00    0.00    0.00    15.08    0.00
Average:         0      512.00    0.42    0.00    N/A    0.00    6.03    N/A
```

Questo rapporto (il quale è stato troncato in modo da rientrare in una unica pagina), include una colonna per ogni livello d'interruzione (per esempio, il campo `i002/s` riporta la velocità per il livello d'interruzione 2). Se si fosse trattato di un sistema a processore multiplo, dovrebbe essere riportata una riga per periodo per ogni CPU.

Un altro punto molto importantnte da ricordare su questo rapporto, è che `sar` è in grado di aggiungere o rimuovere dei campi d'interruzione, se nessun dato viene raccolto per un campo specifico. Il rapporto sopra riportato fornisce un esempio di quanto detto, la fine dello stesso include i livelli d'interruzione (3 e 10), non presenti all'inizio del periodo.

**Nota Bene**

Vi sono altri due rapporti `sar` relativi all'interruzione — `sar -I ALL` e `sar -I XALL`. Tuttavia, la configurazione di default per la utility per la raccolta dei dati `sadc`, non raccoglie le informazioni necessarie per questi rapporti. Tutto questo può essere cambiato modificando il file `/etc/cron.d/sysstat`, e modificando questa riga:

```
*/* * * * * root /usr/lib/sa/sa1 1 1
```

in questo modo:

```
*/* * * * * root /usr/lib/sa/sa1 -I 1 1
```

Ricordatevi che questa modifica causa una raccolta aggiuntiva di informazioni da parte di `sadc`, risultando in misure più grandi del data file. Per questo motivo, assicuratevi che la configurazione del vostro sistema sia in grado di supportare un consumo aggiuntivo dello spazio.

3.4. Risorse aggiuntive

Questa sezione include le risorse aggiuntive che possono essere usate per ottenere maggiori informazioni su argomenti specifici di Red Hat Enterprise Linux affrontati in questo capitolo.

3.4.1. Documentazione installata

Le seguenti risorse sono state installate nel corso di una installazione tipica di Red Hat Enterprise Linux, e possono esservi d'aiuto per ottenere maggiori informazioni su argomenti affrontati in questo capitolo.

- Pagina man di `vmstat(8)` — Imparate a visualizzare una breve panoramica sul processo, sulla memoria, sullo swap, I/O, sul sistema e sull'utilizzo della CPU.
- Pagina man di `iostat(1)` — Imparate a visualizzare le statistiche della CPU e I/O.
- Pagina man di `sar(1)` — Imparate a fornire i rapporti sull'utilizzo delle risorse del sistema.
- Pagina man di `sadc(8)` — Imparate a raccogliere i dati sull'utilizzo del sistema.
- Pagina man di `sa1(8)` — Per saperne di più sullo script che esegue periodicamente `sadc`.
- Pagina man `top(1)` — Imparate a visualizzare l'utilizzo della CPU e le statistiche del process-level.

3.4.2. Siti web utili

- http://people.redhat.com/alikins/system_tuning.html — Informazioni sulla regolazione del sistema. Una serie di approcci per la regolazione della prestazione e per il controllo della risorsa per i server.
- <http://www.linuxjournal.com/article.php?sid=2396> — Tool di controllo della prestazione per Linux. Questo Linux Journal page si rivolge principalmente agli amministratori interessati alla scrittura di una soluzione grafica personalizzata della prestazione. Scritto diversi anni fa, alcune informazioni potrebbero non essere più utili, ma il concetto di base è ancora valido.

3.4.3. Libri correlati

I seguenti libri affrontano diverse problematiche relative al controllo delle risorse, e rappresentano delle risorse utili per gli amministratori di sistema Red Hat Enterprise Linux:

- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — Include un capitolo che riporta la maggior parte dei tool di controllo delle risorse fin qui descritti.
- *Linux Performance Tuning and Capacity Planning* di Jason R. Fink e Matthew D. Sherer; Sams — Fornisce alcune panoramiche più dettagliate sui tool di controllo delle risorse fin qui descritti, e include altri tool che potrebbero essere più appropriati per requisiti più specifici.
- *Linux Administration Handbook* di Evi Nemeth, Garth Snyder, e Trent R. Hein; Prentice Hall — Fornisce un breve capitolo simile al contenuto di questo libro, e include una sezione molto interessante sulla diagnosi di un sistema che improvvisamente ha subito un rallentamento delle proprie prestazioni.
- *Linux System Administration: A User's Guide* di Marcel Gagne; Addison Wesley Professional — Contiene un breve capitolo sul controllo e sulla regolazione della prestazione.

Capitolo 4.

Memoria virtuale e fisica

Al giorno d'oggi i computer preposti ad un uso generale, vengono identificati come *stored program computers*. Come indica il nome, gli stored program computer caricano le istruzioni (la creazione di blocchi di programmi), all'interno di storage interni dove in seguito possono eseguire le suddette istruzioni.

Gli stored program computer, utilizzano per i propri dati lo stesso storage. Questo è in contrasto ai computer che utilizzano la loro configurazione hardware per controllare la loro funzionalità (come ad esempio i vecchi computer basati su di un pannello di connessione).

Il luogo dove i programmi sono stati conservati sui primi stored program computer, poteva presentare diversi nomi e utilizzare molteplici tecnologie. Dai puntini su di un tubo catodico, alla pressione di impulsi nelle colonne di mercurio. Fortunatamente, i computer moderni utilizzano delle tecnologie che rendono possibile una capacità di conservazione migliore con una misura minore.

4.1. Percorsi per l'accesso allo storage

Una cosa da ricordare durante questo capitolo, è che i computer cercano di accedere allo storage in determinati modi. Di fatti la maggior parte degli accessi per lo storage, tendono a mostrare uno (o entrambi) dei seguenti attributi:

- L'accesso tende ad essere sequenziale
- L'accesso tende ad essere localizzato

Accesso sequenziale significa che se la CPU accede l'indirizzo N , è probabile che l'indirizzo $N+1$ verrà accesso successivamente. Ciò ha un certo senso, in quanto molti programmi consistono in sezioni larghe di istruzioni, eseguite — in ordine — uno dopo l'altra.

Accesso localizzato significa che se si accede all'indirizzo X , è probabile che altri indirizzi presenti intorno a X saranno accessibili in futuro.

Questi attributi sono cruciali perchè permettono di avere uno storage più veloce e più piccolo invece di uno più grande e più lento. Questo rappresenta la base per implementare una memoria virtuale. Ma prima di discutere di memoria virtuale, dobbiamo esaminare le diverse tecnologie storage attualmente in uso.

4.2. Gamma dello Storage

I computer moderni utilizzano una varietà di tecnologie storage. Ogni tecnologia si orienta verso una funzione specifica, con relative velocità e capacità.

Queste tecnologie sono:

- Registri CPU
- Memoria Cache
- RAM
- Dischi fissi
- Storage di back-up off line (nastro, disco ottico, ecc)

In termini di capacità e costi, queste tecnologie formano una vasta gamma. Per esempio, i registri CPU sono:

- Molto veloci (tempo di accesso in pochi nanosecondi)
- Capacità bassa (generalmente meno di 200 bytes)
- Capacità di espansione molto limitata (sarà necessario un cambiamento nell'architettura della CPU)
- Costoso (più di un dollaro per byte)

Tuttavia, lo storage di backup off-line è:

- Molto lento (il tempo di accesso si potrebbe tradurre in giorni se il media di backup deve ricoprire lunghe distanze)
- Capacità molto elevata (10 - 100 gigabytes)
- Capacità essenziali di espansione molto limitate (limitato solo dallo spazio necessario per ospitare il media di backup)
- Non è costoso (centesimi per byte)

Utilizzando tecnologie diverse con capacità diverse è possibile migliorare il design del sistema in modo da aumentare le prestazioni diminuendo il costo. Le sezioni seguenti affrontano ogni tecnologia all'interno della gamma dello storage.

4.2.1. Registri CPU

Al giorno d'oggi ogni design della CPU include diversi registri a seconda dei compiti, dalla conservazione dell'indirizzo d'istruzione, a compiti generali di data storage e di manipolazione. I registri CPU funzionano alla stessa velocità del resto della CPU; se così non fosse, essi potrebbero avere un impatto negativo nei confronti delle prestazioni generali del sistema. Il motivo per tale impatto nelle prestazioni, è rappresentato dal fatto che tutte le operazioni eseguite dalla CPU necessitano l'utilizzo dei registri.

Il numero dei registri CPU (ed il loro utilizzo) dipendono strettamente dal design della CPU stesso. Non vi è alcun modo per modificare il numero di registri CPU a meno che non migrate ad un CPU con diversa architettura. Per queste ragioni, il numero di registri CPU può essere considerato costante, in quanto modificarli potrebbe essere molto laborioso e costoso.

4.2.2. Memoria Cache

Il compito della memoria cache è quello di comportarsi come un buffer tra i registri della CPU e la memoria del sistema molto più lenta e grande — generalmente indicata come RAM¹. La Memoria Cache possiede una velocità di funzionamento simile alla CPU stessa, quando la stessa CPU accede i dati nella cache, la stessa non viene messa in attesa.

La Memoria Cache viene configurata in modo che se i dati devono essere letti da una RAM, l'hardware del sistema esegue un controllo per determinare se i dati desiderati sono presenti nella cache. Se i suddetti dati sono presenti nella cache, essi vengono ripresi velocemente e usati dalla CPU. Tuttavia, se i dati non sono presenti nella cache, vengono letti dalla RAM e durante il loro trasferimento alla CPU, vengono posti nella cache (in caso gli stessi dati servissero più in avanti). Dal punto di vista della CPU, tutto questo viene eseguito in modo pultito, così l'unica differenza tra

1. "RAM è l'acronimo di "Random Access Memory," e rappresenta un termine che può essere usato per ogni tecnologia storage che permette un accesso non-sequenziale dei dati conservati, quando gli amministratori usano il termine RAM, essi indicano la memoria principale del sistema .

l'accedere i dati in una cache e i dati in una RAM, è solo la quantità di tempo necessario per ritornare i dati stessi.

In termini di capacità dello storage, la cache è molto più piccola della RAM. Per questo motivo ogni byte presente nella RAM, non può avere una propria posizione unica nella cache. In questo caso è necessario dividere la cache in sezioni, capaci di essere usate per conservare diverse aree della RAM in tempi diversi. Anche se vi è una differenza in misura tra la cache e la RAM, data la sequenza e la natura dell'accesso allo storage, una piccola quantità di cache è in grado di velocizzare l'accesso ad una quantità maggiore di RAM.

Quando si scrivono i dati direttamente dalla CPU, le cose si complicano leggermente. Si possono usare due diversi approcci. In entrambi i casi, i dati vengono prima scritti su di una cache. Tuttavia, poiché il compito di una cache è quello di funzionare come una copia veloce dei contenuti di determinate porzioni della RAM, ogni volta che una parte dei dati varia il proprio valore, il nuovo valore deve essere scritto sia sulla memoria cache che sulla RAM. Diversamente, i dati presenti sia sulla cache che sulla RAM non corrisponderanno più.

I due approcci differiscono nel modo di eseguire tale procedura. Un approccio, conosciuto come caching *write-through*, scrive immediatamente i dati modificati su di una RAM. Il caching *Write-back* tuttavia, ritarda la scrittura sulla RAM di dati modificati. Il motivo è quello di ridurre la riscrittura dei dati modificati sulla RAM, ogni qualvolta i suddetti dati vengono modificati.

La cache del tipo 'Write-through' è leggermente più semplice da implementare, per questa ragione è la più comune. La cache 'Write-back' è leggermente più complessa; in aggiunta a conservare i dati, è necessario mantenere un meccanismo in grado di evidenziare i dati conservati nella cache come dati puliti (cioè i dati presenti nella cache sono gli stessi presenti nella RAM), o sporchi (cioè i dati nella cache sono stati modificati, e quindi i dati nella RAM non sono più aggiornati). È necessario altresì implementare un modo per ripulire periodicamente le entry sporche presenti nella cache sulla RAM.

4.2.2.1. Livelli della cache

I sottosistemi odierni della cache possono avere diversi livelli; e cioè, ci potrebbe essere più di un set di cache tra la CPU e la memoria principale. I livelli sono spesso numerati, con i numeri più bassi accostati alla CPU. Molti sistemi possiedono due livelli cache:

- L1 cache viene spesso posizionato direttamente sul chip della CPU, e funziona alla stessa velocità della CPU
- L2 cache è spesso parte del modulo della CPU, funziona alla stessa velocità della CPU (o molto simile) ed è generalmente un po' più grande e lento di L1 cache

Alcuni sistemi (normalmente server ad alte prestazioni) possiedono anche l'L3 cache, il quale è generalmente parte della scheda madre del sistema. Probabilmente, L3 cache è più grande (e forse più lento) di L2 cache.

L'obiettivo di tutti i sottosistemi cache —, siano essi singoli o multi-level —, è quello di ridurre il tempo di accesso medio ad una RAM.

4.2.3. Memoria principale — RAM

RAM rappresenta la parte principale dello storage elettronico nei sistemi moderni. Viene usata come storage sia per i dati che per i programmi mentre gli stessi sono in uso. La velocità della RAM in molti sistemi moderni si identifica tra la velocità della cache memory e quella dei dischi fissi.

La funzionalità di base della RAM è molto semplice. Al suo livello di base ci sono i chip della RAM — circuiti integrati che eseguono l'operazione di "remembering." Questi chip possiedono quattro tipi di collegamenti:

- Collegamenti di alimentazione (per alimentare il circuito all'interno del chip)
- Collegamenti dati (per abilitare il trasferimento dei dati in entrata o in uscita dal chip)
- Collegamenti lettura/scrittura (per controllare i dati da conservare o da riprendere dal chip)
- Collegamenti dell'indirizzo (per determinare dove bisogna leggere/scrivere i dati sul chip)

Ecco riportate le fasi necessarie per conservare i dati nella RAM:

1. I dati da conservare vengono presentati ai collegamenti.
2. L'indirizzo tramite il quale i dati devono essere conservati viene presentato ai collegamenti dell'indirizzo stesso.
3. Il collegamento lettura/scrittura è impostato in modalità di scrittura.

La ripresa dei dati è molto semplice:

1. L'indirizzo dei dati desiderati viene presentato ai collegamenti dell'indirizzo stesso.
2. Il collegamento lettura/scrittura viene impostato in modalità lettura.
3. I dati desiderati vengono letti dai collegamenti dei dati stessi.

Anche se queste fasi possono sembrare molto semplici, esse vengono eseguite ad una velocità molto elevata, con una durata per ogni fase che si aggira intorno ai nanosecondi.

Quasi tutti i chip della RAM creati al giorno d'oggi vengono venduti come *moduli*. Ogni modulo consiste in un numero di chip RAM individuali collegati ad una piccola scheda del circuito. La struttura elettrica e meccanica del modulo, si attiene ai diversi standard del settore rendendo possibile così l'acquisto della memoria da diversi rivenditori.



Nota Bene

Il beneficio più importante per un sistema che utilizza i moduli RAM standard, è quello di mantenere un costo basso della RAM, grazie al fatto che è possibile acquistare i moduli da più rivenditori.

Anche se molti computer usano moduli RAM standard, ci sono sempre delle eccezioni. Le più comuni sono i computer portatili (ed anche qui si inizia a notare una certa standardizzazione), e server high-end. Tuttavia, anche qui è probabile che possano essere disponibili moduli RAM di terzi, presupponendo che il sistema sia relativamente diffuso e non un ultimissimo modello.

4.2.4. Dischi fissi

Tutte le tecnologie fino adesso affrontate sono in natura *volatili*. In altre parole, i dati contenuti in uno storage volatile, vengono persi quando si disabilita l'alimentazione.

I dischi fissi, d'altro canto, sono *non-volatili*, — i dati contenuti restano presenti anche se si disabilita l'alimentazione. Per questo motivo, i dischi fissi occupano un posto particolare nella gamma dello storage. La suddetta natura li rende ideali per la conservazione dei programmi e dei dati utili a lungo termine. Un altro aspetto unico dei dischi fissi è quello che a differenza della memoria RAM e cache, non è possibile eseguire programmi in modo diretto quando questi sono conservati in dischi fissi; è necessario invece scriverli prima all'interno della RAM.

Diverso tra cache e RAM è la velocità di conservazione dei dati e del loro recupero; i dischi fissi sono più lenti delle tecnologie elettroniche usate per la memoria cache e RAM. Tale differenza è dovuta principalmente alla loro natura elettromeccanica. Per il trasferimento da e per il disco fisso vi sono

quattro fasi distinte. Il seguente elenco mostra le suddette fasi, insieme con il tempo relativo ad ognuna di esse ad essere completata:

- Movimento del braccio d'accesso (5.5 millisecondi)
- Rotazione del disco (.1 millisecondi)
- Testine di lettura/scrittura dei dati (.00014 millisecondi)
- Trasferimento dati per/da le elettroniche dell'unità (.003 Millisecondi)

Di questi, solo l'ultima fase non dipende da alcuna operazione meccanica.



Nota Bene

Anche se vi è ancora molto da imparare, le tecnologie per il disk storage vengono affrontate più in dettaglio nel Capitolo 5. È necessario per adesso ricordarsi dell'enorme differenza in velocità che vi è tra la RAM e le tecnologie basate sul disco, e che la loro capacità di conservazione generalmente è maggiore di quella della RAM, almeno di un minimo di 10 volte fino ad un massimo di oltre 100.

4.2.5. Storage di backup off-line

Lo storage di backup off-line rappresenta un passo in avanti in termini di capacità nei confronti del disco fisso (maggiore) e in velocità (più lento). In questo caso le capacità sono effettivamente limitate solo dalla vostra abilità di procurare e conservare i media in grado di essere rimossi.

Le attuali tecnologie usate in questi dispositivi varia moltissimo. Ecco quelle più diffuse:

- Nastro magnetico
- Disco ottico

Avere dei media in grado di essere rimossi significa anche che il tempo d'accesso diventa ancora più lungo, in particolare quando i dati desiderati sono su di un media non presente sul dispositivo storage. Tale situazione viene alleviata dall'uso di dispositivi specifici in grado di caricare e scaricare automaticamente i media, ma la capacità di conservazione di questi dispositivi è piuttosto limitata. Anche nelle migliori delle ipotesi, i tempi di accesso vengono misurati in secondi, rappresentando comunque un tempo di accesso molto lungo anche rispetto alle prestazioni dei dischi fissi ad alte prestazioni, generalmente misurati in multi-millisecondi e considerati relativamente lenti.

Ora che abbiamo affrontato le diverse tecnologie per lo storage, esploriamo adesso i concetti di base per la memoria virtuale.

4.3. Concetti di base della memoria virtuale

Anche se la tecnologia usata al giorno d'oggi per la creazione delle varie tecnologie storage è notevole, un normale amministratore di sistema non ha necessità di conoscere tutti i vari dettagli. Infatti vi è solo un elemento che gli amministratori devono tener presente:

Non vi è mai abbastanza RAM.

Anche se questa affermazione potrebbe sembrare spiritosa, molti designer di sistemi operativi hanno lavorato moltissimo per ridurre l'impatto causato da questa carenza. Per far fronte a tale problematica hanno deciso d'implementare una *memoria virtuale* — un modo per combinare la RAM con una unità storage più lente, in modo da conferire al sistema più RAM di quella installata originariamente.

4.3.1. Spiegazione semplice della memoria virtuale

Iniziamo con una ipotetica applicazione. Il codice della macchina che crea la suddetta applicazione è di 10000 byte di misura. Sono necessari altresì altri 5000 byte per la conservazione dei dati e dei buffer I/O. Ciò significa che, per eseguire questa applicazione sono necessari 15000 byte di RAM; non un byte in meno, altrimenti l'applicazione non sarà in grado di essere eseguita.

Questi 15000 byte sono conosciuti come lo *spazio dell'indirizzo* dell'applicazione. Rappresenta il numero di indirizzi unici necessari per contenere sia l'applicazione che i dati. Nei primi computer, la quantità di RAM disponibile doveva essere maggiore dello spazio dell'indirizzo dell'applicazione più grande da eseguire; al contrario, l'applicazione compariva con un errore del tipo "out of memory".

L'approccio seguente conosciuto come *overlaying* ha cercato di trovare una soluzione a tale problema, permettendo ai programmatori di decidere in quale parte dell'applicazione potesse risiedere la memoria. In questo modo, il codice richiesto solo una volta per l'inizializzazione, poteva essere sovrascritto (overlayed) con quello da usare più in avanti. Mentre tale operazione è riuscita a facilitare in qualche modo la soluzione per la carenza di memoria, essa rappresentava un processo complesso e propenso ad errori. Tale approccio non era in grado altresì di trovare una soluzione per la carenza di memoria dell'intero sistema durante l'esecuzione dello stesso. In altre parole, un programma di tipo 'overlayed' potrebbe richiedere minor memoria per essere eseguito rispetto ad un programma non overlayed, ma se il sistema non ha memoria sufficiente per il programma 'overlayed', si avrà lo stesso risultato finale — un errore del tipo out of memory.

Con la memoria virtuale si pone il concetto dello spazio dell'indirizzo di una applicazione sotto un diverso punto di vista. Invece di sapere la *quantità* di memoria necessaria ad una applicazione per essere eseguita, un sistema operativo con memoria virtuale va alla ricerca continua "della quantità minima di memoria necessaria ad una applicazione per essere eseguita".

Mentre la nostra ipotetica applicazione necessita di tutti e 15000 byte per essere eseguita, cercate di ricordare l'argomento affrontato nella Sezione 4.1 — l'accesso della memoria tende ad essere sequenziale e localizzato. Per questo motivo, la quantità di memoria necessaria ad eseguire l'applicazione in ogni istante, è minore di 15000 byte — generalmente molto meno. Considerate i diversi tipi di accesso della memoria necessari per eseguire una istruzione singola della macchina:

- L'istruzione viene letta dalla memoria.
- I dati richiesti dall'istruzione sono letti dalla memoria.
- Dopo il completamento dell'istruzione, i risultati della stessa vengono riscritti all'interno della memoria.

Il numero di byte necessari per ogni accesso della memoria varia a seconda dell'architettura della CPU, dal tipo d'istruzione e dei dati. Tuttavia, anche se una istruzione necessita di 100 byte di memoria per ogni tipo di accesso, i 300 byte necessari rappresentano una quantità molto inferiore ai 15000 byte dello spazio dell'indirizzo. Se si potesse trovare un modo per avere informazioni sui requisiti della memoria di un'applicazione durante l'esecuzione della stessa, sarebbe possibile mantenere una applicazione in esecuzione usando una quantità di memoria minore rispetto al suo spazio.

Ma tutto ciò fa sorgere una domanda:

Se solo parte dell'applicazione è presente nella memoria in ogni istante, dov'è il resto?

4.3.2. Backing Store — Il principio centrale della memoria virtuale

Una risposta a questa domanda è che il resto dell'applicazione rimane nel disco. In altre parole, il disco agisce come un *backing store* per la RAM; un mezzo per lo storage più grande e più lento, che funziona come un "backup" per un mezzo più piccolo e veloce. Questo potrebbe apparire un pò problematico in termini di prestazione — dopo tutto, le unità del disco sono molto più lente della RAM.

È possibile comunque trarre vantaggio dal comportamento sequenziale e localizzato delle applicazioni, eliminando molte delle problematiche inerenti le prestazioni nell'uso delle unità del disco come backing store per la RAM. Questo viene fatto intervenendo sul sottosistema della memoria virtuale, in modo da cercare di assicurare che le parti necessarie dell'applicazione — o che potrebbero essere utili in futuro —, siano conservate nella RAM solo per il tempo necessario.

Questo potrebbe essere simile al rapporto tra la cache e la RAM: e cioè combinando un tipo di storage piccolo e veloce, con uno più grande e più lento, comportandosi quindi come uno storage più grande e rapido.

Tenendo presente questo paragone, affrontiamo il processo in modo più dettagliato.

4.4. Memoria virtuale: Dettagli

Introduciamo prima un nuovo concetto: *Lo spazio dell'indirizzo virtuale*. Tale spazio rappresenta la quantità massima di spazio per l'indirizzo disponibile per una applicazione. Lo spazio dell'indirizzo virtuale varia a seconda dell'architettura del sistema e del sistema operativo. Esso dipende dall'architettura perché è quest'ultima che definisce il numero di bit disponibili per l'indirizzo. Dipende anche dal sistema operativo e a seconda del modo con il quale è stato implementato il sistema operativo stesso potrebbe apportare dei limiti aggiuntivi, superiori o inferiori, a quelli imposti dall'architettura.

La parola "virtuale" nello spazio dell'indirizzo virtuale, sta ad indicare il numero totale di luoghi della memoria indirizzabili in modo unico disponibili per una applicazione, e *non* la quantità di memoria fisica installata nel sistema o dedicata all'applicazione.

Nel caso del nostro esempio, lo spazio per l'indirizzo virtuale è di 15000 byte.

Per implementare la memoria virtuale è necessario, per il sistema del computer, avere un hardware speciale per la gestione della memoria. Questo hardware è spesso conosciuto come *MMU* (Memory Management Unit). Senza di esso, quando la CPU accede alla RAM, le posizioni della RAM non cambiano mai — l'indirizzo della memoria 123 rappresenta sempre la stessa posizione fisica all'interno della RAM.

Tuttavia, con una MMU, gli indirizzi della memoria attraversano una fase di traslazione prima dell'accesso alla memoria. Ciò significa che l'indirizzo 123 della memoria, potrebbe essere direzionato verso l'indirizzo fisico 82043 prima, e verso l'indirizzo privato 20468 poi, l'overhead nel seguire in modo individuale la traslazione virtuale e fisica per bilioni di byte di memoria, potrebbe essere troppo grande. Invece, l'MMU divide la RAM in *pagine* — sezioni contigue di memoria di misura predeterminata gestite dalla MMU come entità singole.

Mantenere le informazioni di queste pagine e delle traslazioni dell'indirizzo, potrebbe sembrare una fase non necessaria e confusionaria. Tuttavia risulta molto importante implementare la memoria virtuale. Per questo motivo vi consigliamo di tenere in giusta considerazione questo punto.

Prendendo in considerazione la nostra ipotetica applicazione con uno spazio dell'indirizzo virtuale di 15000 byte, considerate che la prima istruzione dell'applicazione acceda i dati conservati nell'indirizzo 12374. Tuttavia, prendete in considerazione anche che il nostro computer possiede solo 12288 byte di RAM fisica. Cosa succede se la CPU cerca di accedere l'indirizzo 12374?

Quello che succede viene chiamato *page fault*.

4.4.1. Page Fault

Un page fault rappresenta una sequenza di eventi che si presentano quando un programma cerca di accedere i dati (o codice) presenti nel proprio spazio dell'indirizzo, ma che gli stessi non si trovano nella RAM del sistema. Il sistema operativo deve gestire le page fault facendo risiedere la memoria dei dati in questione, permettendo al programma di continuare le sue funzioni come se il page fault non fosse mai accaduto.

Nel caso della nostra ipotetica applicazione, la CPU presenta prima l'indirizzo desiderato (12374) alla MMU. Tuttavia, la MMU non ha alcuna transazione per questo indirizzo. Così interrompe la CPU e causa l'esecuzione del software, conosciuto come gestore della page fault. Tale gestore determina cosa si deve fare per risolvere tale problema. E cioè:

- Trova sul disco dove risiede la pagina desiderata e la legge (questo è il caso se il page fault avviene per una pagina del codice)
- Determina se la pagina desiderata è già nella RAM (ma non è assegnata al processo corrente), e riconfigura la MMU in modo da poterlo indicare
- Indica una pagina speciale che contiene solo degli zeri, e assegna una nuova pagina solo se il processo cerca di scrivere sulla pagina speciale (tale processo viene chiamato *copy on write*, ed è spesso usato per pagine che non contengono alcun dato inizializzato)
- Ottiene la pagina desiderata da un altro luogo (tale procedura viene affrontata in dettaglio più in avanti)

Mentre le prime tre fasi sono relativamente semplici, l'ultima non lo è affatto. Per questo motivo è necessario affrontare argomenti aggiuntivi.

4.4.2. Il working set

Il gruppo di pagine della memoria fisica attualmente dedicato ad un processo specifico è chiamato *working set* per il processo stesso. Il numero di pagine nel working set può aumentare o diminuire, a seconda della disponibilità delle pagine stesse.

Il working set aumenta come un processo page fault. Al contrario esso diminuisce con la diminuzione delle pagine disponibili. Per evitare la consumazione completa della memoria, le pagine devono essere rimosse dal working set del processo e trasformate in pagine disponibili per un loro eventuale utilizzo. Il sistema operativo diminuisce i working set del processo nei seguenti modi:

- Scrivendo su pagine modificate in un'area dedicata, su di un dispositivo di tipo mass storage (generalmente conosciuti come spazio di *swapping* o *paging*)
- Contrassegnando pagine non modificate come libere (non vi è alcuna necessità di scrivere queste pagine su disco in quanto queste non sono state modificate)

Per determinare i working set appropriati per tutti i processi, il sistema operativo deve possedere tutte le informazioni sull'utilizzo per tutte le pagine. In questo modo, il sistema operativo determina le pagine usate in modo attivo (risiedendo sempre nella memoria) e quelle non utilizzate (e quindi da rimuovere dalla memoria.) In molti casi, una sorta di algoritmo non usato di recente, determina le pagine che possono essere rimosse dal working set dei processi.

4.4.3. Swapping

Anche se lo swapping (scrittura delle pagine modificate sullo spazio swap del sistema) rappresenta una operazione normale della funzione di un sistema, è possibile che si verifichi uno swapping eccessivo. Il motivo per il quale bisogna fare attenzione ad uno swapping eccessivo, è la possibilità che si possa verificare, in modo costante, la seguente situazione:

- Vengono scambiate le pagine di un processo
- Il processo diventa eseguibile e cerca di accedere una pagina precedentemente scambiata
- La pagina è erroneamente rientrata nella memoria (forzando probabilmente le pagine di un altro processo ad essere scambiata)
- Dopo qualche istante, la pagina viene scambiata nuovamente

Se la sequenza di questi eventi è molto diffusa ecco che si verifica il *thrashing*, il quale indica una insufficienza di RAM per il carico di lavoro attuale. Il Thrashing è estremamente dannoso per le prestazioni del sistema, in quanto la CPU ed i carichi I/O che possono essere generati, possono avere il sopravvento sul carico imposto dal reale funzionamento del sistema. In casi estremi, il sistema potrebbe non funzionare utilmente, spendendo tutte le sue risorse spostando le pagine da e per la memoria.

4.5. Implicazioni sulle prestazioni della memoria virtuale

Mentre la memoria virtuale rende possibile per i computer la gestione di applicazioni più complesse, è necessario, come con qualsiasi tool molto potente, anche far fronte al suo lato negativo. Tale lato negativo è rappresentato dalla prestazione — un sistema operativo con memoria virtuale ha molto più da elaborare rispetto ad un sistema operativo incapace di supportare la stessa memoria virtuale. Ciò significa che non si avrà mai una buona prestazione con la presenza di memoria virtuale quando la stessa applicazione risiede 100% nella memoria stessa.

Tuttavia, ciò non rappresenta un problema insormontabile. I benefici di una memoria virtuale sono veramente molteplici, e con un pò d'impegno è possibile ottenere una buona prestazione. È consigliabile esaminare le risorse del sistema che sono state influenzate da un uso eccessivo del sottosistema della memoria virtuale.

4.5.1. Come ottenere una prestazione pessima

Per un momento, prendete in considerazione quello che avete letto in questo capitolo, e considerate le risorse del sistema usate per una attività molto intensa di swapping e page fault:

- RAM — È normale che la RAM disponibile sia bassa (altrimenti non ci sarebbe bisogno di uno swap o page fault).
- Disco — Mentre lo spazio del disco potrebbe non essere intaccato, potrebbe esserlo invece la larghezza della banda I/O (a causa di attività swapping e paging molto pesante).
- CPU — La CPU espande i cicli eseguendo il processo necessario per supportare la gestione della memoria, e impostando le operazioni necessarie I/O per le operazioni di swapping e paging.

La natura correlata di questi carichi, rende più semplice la comprensione di come la carenza delle risorse possa causare dei problemi alla prestazione.

È solo necessario avere una RAM troppo piccola, un'attività page fault molto intensa, e un sistema che opera molto vicino ai propri limiti in termini di CPU o disco I/O. A questo punto, il sistema viene affetto da thrashing, ne consegue quindi, una prestazione molto povera.

4.5.2. Come ottenere la prestazione migliore

Nel migliore dei casi, l'overhead dovuto dal supporto della memoria virtuale presenta un carico minimo aggiuntivo per un sistema configurato in modo corretto:

- xRAM — RAM sufficiente per tutti i working set in grado di gestire qualsiasi page fault ²
- Disco — A causa della limitata attività di page fault, la larghezza della banda del disco I/O verrà influenzata in maniera molto limitata

2. Un sistema attivo incontrerà *sempre* qualche livello di page fault, a causa della loro introduzione nella memoria di applicazioni appena lanciate.

- CPU — La maggior parte dei cicli CPU sono dedicati all'esecuzione delle applicazioni invece di eseguire il codice di gestione della memoria del sistema

È da ricordare quindi che l'impatto sulle prestazioni della memoria virtuale è minimo se essa viene usata il meno possibile. Questo significa che per avere una buona prestazione del sottosistema della memoria virtuale, bisogna avere una RAM sufficiente.

La fase successiva (ma meno importante) è avere una capacità sufficiente di CPU e del disco I/O. Tuttavia, ricordate che queste risorse aiutano a preservare la prestazione del sistema nei riguardi di attività di swapping e faulting molto elevata; questi requisiti fanno poco per aiutare la prestazione del sottosistema della memoria virtuale (anche se gli stessi possono avere un certo impatto nelle prestazioni generali del sistema).

4.6. Red Hat Enterprise Linux-Informazioni specifiche

A causa della complessità per essere un sistema operativo con memoria virtuale 'demand-paged', il controllo delle risorse relative alla memoria con Red Hat Enterprise Linux può essere un pò ambiguo. Per questo, è consigliabile iniziare con tool molto più semplici.

Usando *free*, è possibile avere una panoramica (piuttosto semplicistica) della memoria e dell'utilizzo di swap. Ecco un esempio:

	total	used	free	shared	buffers	cached
Mem:	1288720	361448	927272	0	27844	187632
-/+ buffers/cache:		145972	1142748			
Swap:	522104	0	522104			

Possiamo notare che questo sistema possiede 1.2GB di RAM, e solo 350MB di essa è in uso. Come previsto con un sistema con tanta RAM disponibile, 500MB di swap non vengono utilizzati.

Paragonate quell'esempio con questo:

	total	used	free	shared	buffers	cached
Mem:	255088	246604	8484	0	6492	111320
-/+ buffers/cache:		128792	126296			
Swap:	530136	111308	418828			

Questo sistema ha circa 256MB di RAM, la maggior parte della quale è in uso, lasciando solo 8MB disponibili. Vengono utilizzati 100MB dei 512MB della partizione swap in uso. Anche se questo sistema è sicuramente molto più limitato in termini di memoria rispetto al primo sistema, per determinare se questa limitazione è in grado di causare dei problemi sulla prestazione, è necessario andare più a fondo.

Anche se più enigmatico di *free*, *vmstat* ha la capacità di visualizzare maggiori informazioni oltre alle statistiche sull'utilizzo della memoria. Ecco l'output di *vmstat 1 10*:

procs				memory			swap		io				system		cpu	
r	b	w	swpd	free	buff	cache	si	so	bi	bo	in	cs	us	sy	id	
2	0	0	111304	9728	7036	107204	0	0	6	10	120	24	10	2	89	
2	0	0	111304	9728	7036	107204	0	0	0	0	526	1653	96	4	0	
1	0	0	111304	9616	7036	107204	0	0	0	0	552	2219	94	5	1	
1	0	0	111304	9616	7036	107204	0	0	0	0	624	699	98	2	0	
2	0	0	111304	9616	7052	107204	0	0	0	48	603	1466	95	5	0	
3	0	0	111304	9620	7052	107204	0	0	0	0	768	932	90	4	6	
3	0	0	111304	9440	7076	107360	92	0	244	0	820	1230	85	9	6	
2	0	0	111304	9276	7076	107368	0	0	0	0	832	1060	87	6	7	
3	0	0	111304	9624	7092	107372	0	0	16	0	813	1655	93	5	2	

2 0 2 111304 9624 7108 107372 0 0 0 972 1189 1165 68 9 23

Durante questi 10 secondi, la quantità di memoria disponibile (il campo *free*) in qualche modo varia, ed è presente anche I/O relativo allo swap (i campi *si* e *so*), in linea di massima questo sistema funziona abbastanza bene. Tuttavia si possono verificare delle problematiche nel caso in cui si verificasse un aumento del carico di lavoro.

Nella ricerca delle problematiche riguardanti la memoria, è spesso necessario determinare come il sottosistema della memoria virtuale di Red Hat Enterprise Linux, faccia uso della memoria del sistema. Usando *sar*, è possibile esaminare l'aspetto delle prestazioni del sistema più in dettaglio.

Rivisionando il report di *sar -r*, possiamo esaminare l'utilizzo di swap e della memoria più da vicino:

```
Linux 2.4.20-1.1931.2.231.2.10.ent (pigdog.example.com) 07/22/2003

12:00:01 AM kbmemfree kbmemused %memused kbmemshrd kbbuffers kbcached
12:10:00 AM 240468 1048252 81.34 0 133724 485772
12:20:00 AM 240508 1048212 81.34 0 134172 485600
...
08:40:00 PM 934132 354588 27.51 0 26080 185364
Average: 324346 964374 74.83 0 96072 467559
```

I campi *kbmemfree* e *kbmemused* mostrano le statistiche tipiche della memoria disponibile e di quella utilizzata, con la percentuale di memoria usata, visualizzata nel campo *%memused*. I campi *kbbuffers* e *kbcached* mostrano il numero di kilobyte di memoria che viene assegnato ai buffer e alla cache dei dati dell'intero sistema.

Il campo *kbmemshrd* è sempre zero per i sistemi (come ad esempio Red Hat Enterprise Linux) che utilizzano il kernel 2.4 di Linux.

Le righe per questo esempio sono state troncate in modo da rientrare nella pagina. Ecco il remainder di ogni riga, con aggiunto il timestamp sulla sinistra per rendere la lettura più semplice:

```
12:00:01 AM kbswpfree kbswpused %swpused
12:10:00 AM 522104 0 0.00
12:20:00 AM 522104 0 0.00
...
08:40:00 PM 522104 0 0.00
Average: 522104 0 0.00
```

Per l'utilizzo di swap, i campi *kbswpfree* e *kbswpused* mostrano la quantità di spazio swap usato e quello disponibile in kilobyte, con il campo *%swpused* viene mostrato il suddetto spazio in percentuale.

Per saperne di più sull'attività di swap, usare il report *sar -W*. Ecco un esempio:

```
Linux 2.4.20-1.1931.2.231.2.10.entsmp (raptor.example.com) 07/22/2003

12:00:01 AM pswpin/s pswpout/s
12:10:01 AM 0.15 2.56
12:20:00 AM 0.00 0.00
...
03:30:01 PM 0.42 2.56
Average: 0.11 0.37
```

Qui possiamo notare che in linea di massima, vi è stata una quantità pari a tre volte in meno di pagine importate dallo swap (*pswpin/s*) rispetto a quelle esportate (*pswpout/s*).

Per capire meglio come vengono usate le pagine, consultare il report `sar -B`:

```
Linux 2.4.20-1.1931.2.231.2.10.entsmp (raptor.example.com) 07/22/2003

12:00:01 AM pgpgin/s pgpgout/s activepg inadtypg inaclnpg inatargp
12:10:00 AM 0.03 8.61 195393 20654 30352 49279
12:20:00 AM 0.01 7.51 195385 20655 30336 49275
...
08:40:00 PM 0.00 7.79 71236 1371 6760 15873
Average: 201.54 201.54 169367 18999 35146 44702
```

Qui possiamo determinare il numero di blocchi al secondo impaginati dal disco (`pgpgin/s`), oppure estromessi dallo stesso (`pgpgout/s`). Queste statistiche servono come misurazione dell'attività generale della memoria virtuale.

Tuttavia potete ottenere più informazioni esaminando gli altri campi presenti in questo rapporto. Il kernel di Red Hat Enterprise Linux segna tutte le pagine come attive o inattive. Come implica il nome, le pagine attive sono quelle che vengono impiegate in qualche modo (come pagine del buffer o per il processo), al contrario invece ci sono le pagine inattive, e cioè quelle non usate in alcun processo (il campo `activepg`) approssimativamente 660MB³.

Il remainder dei campi in questo rapporto, si concentra sull'elenco delle pagine inattive — che, per qualche motivo, non sono state usate. Il campo `inadtypg` mostra il numero di pagine inattive contrassegnate con *dirty* (modificate), le quali potrebbero essere scritte sul disco. Il campo `inaclnpg` invece, mostra il numero di pagine inattive *clean* (non modificate), le quali non hanno bisogno di essere scritte sul disco.

Il campo `inatargp` rappresenta la misura ideale dell'elenco delle pagine inattive. Questo valore viene calcolato dal kernel di Linux e la sua misura fa sì che l'elenco di pagine inattive rimanga sufficientemente grande, per agire come un insieme per la sostituzione delle pagine.

Per maggiori informazioni sullo stato delle pagine (e in modo particolare la frequenza con la quale cambiano le pagine), utilizzate il rapporto `sar -R`. Ecco un esempio:

```
Linux 2.4.20-1.1931.2.231.2.10.entsmp (raptor.example.com) 07/22/2003

12:00:01 AM frmpg/s shmpg/s bufpg/s campg/s
12:10:00 AM -0.10 0.00 0.12 -0.07
12:20:00 AM 0.02 0.00 0.19 -0.07
...
08:50:01 PM -3.19 0.00 0.46 0.81
Average: 0.01 0.00 -0.00 -0.00
```

Le statistiche nel rapporto `sar` sono uniche, cioè esse possono essere positive, negative o zero. Quando esse sono positive, il valore indica la velocità con la quale queste pagine sono in aumento. Quando negativo, il valore indica la velocità con la quale queste pagine sono in diminuzione. Con un valore pari a zero, la pagine nè aumentano e nè diminuiscono.

In questo caso, l'ultimo esempio mostra una velocità di poco superiore a tre pagine al secondo assegnate dall'elenco di pagine disponibili (il campo `frmpg/s`), e quasi 1 pagina per secondo aggiunta alla cache della pagina (campo `campg/s`). Questo elenco di pagine utilizzato come buffer (il campo `bufpg/s`), ha guadagnato approssimativamente una pagina ogni due secondi, mentre l'elenco della pagina della memoria condivisa (il campo `shmpg/s`), non ha nè guadagnato e nè perduto alcuna pagina.

3. La misura con Red Hat Enterprise Linux, usata in questo esempio su di un sistema x86, è di 4096 bytes. I sistemi basati su di un'altra architettura possono avere una misura diversa.

4.7. Risorse aggiuntive

Questa sezione include diverse risorse che possono essere usate per avere maggiori informazioni sul controllo delle risorse stesse, e di Red Hat Enterprise Linux affrontato in questo capitolo.

4.7.1. Documentazione installata

Le seguenti risorse vengono installate durante il corso di una installazione tipica di Red Hat Enterprise Linux, e vi consentono di ottenere maggiori informazioni sugli argomenti affrontati in questo capitolo

- Pagina man di `free(1)` — Imparate a visualizzare le statistiche sulla memoria usata e su quella ancora disponibile.
- Pagina man di `vmstat(8)` — Imparate a visualizzare una panoramica breve sull'uso del processo, della memoria, dello swap, I/O, del sistema e della CPU.
- Pagina man di `sar(1)` — Imparate a creare i rapporti sull'utilizzo delle risorse del sistema.
- Pagina man di `sa2(8)` — Imparate a creare i file di riporto sull'utilizzo giornaliero delle risorse del sistema.

4.7.2. Siti web utili

- http://people.redhat.com/alikins/system_tuning.html — System Tuning Info per i server di Linux. Un approccio sulla regolazione della prestazione e sul controllo delle risorse da parte dei server.
- <http://www.linuxjournal.com/article.php?sid=2396> — Tool sul controllo delle prestazioni di Linux. Questa pagina Journal di Linux si propone più come di interesse per gli amministratori desiderosi di scrivere una soluzione grafica della prestazione personalizzata. Scritta diversi anni fa, alcuni dei suggerimenti potrebbero essere non più validi, ma il concetto generale e di esecuzione sono simili.

4.7.3. Libri correlati

I seguenti libri affrontano varie problematiche relative al controllo delle risorse, e rappresentano risorse importanti per gli amministratori del sistema di Red Hat Enterprise Linux:

- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — Include un capitolo contenente diversi tool per la gestione delle risorse finora descritte.
- *Regolazione della prestazione di Linux e Pianificazione della capacità* di Jason R. Fink e Matthew D. Sherer; Sams — Fornisce una panoramica più dettagliata sui tool di controllo delle risorse, includendone altri che potrebbero esservi utili per compiti particolari.
- *Red Hat Linux Security and Optimization* di Mohammed J. Kabir; Red Hat Press — Circa le prime 150 pagine di questo libro affrontano le problematiche relative alla prestazione. Vengono inclusi capitoli dedicati alle problematiche specifiche della rete, Web, email e file server.
- *Linux Administration Handbook* di Evi Nemeth, Garth Snyder, and Trent R. Hein; Prentice Hall — Fornisce un breve capitolo simile nello scopo a questo manuale, ma include una sezione molto interessante sulla diagnosi di un sistema che ha improvvisamente rallentato le proprie prestazioni.
- *Linux System Administration: Guida per l'utente* di Marcel Gagne; Addison Wesley Professional — Contiene un piccolo capitolo sul controllo delle prestazioni e della regolazione.
- *Essential System Administration* (terza edizione) di Aeleen Frisch; O'Reilly & Associates — Il capitolo sulla gestione delle risorse del sistema contiene delle informazioni generali molto utili, incluse alcune relative a Linux.

- *System Performance Tuning* (seconda edizione) di Gian-Paolo D. Musumeci and Mike Loukides; O'Reilly & Associates — Anche se orientato pesantemente sulle implementazioni tradizionali di UNIX, sono presenti diverse referenze su Linux.

Capitolo 5.

Gestione dello storage

Se vi è una cosa che è in grado d'impegnare la maggior parte delle attenzioni di un amministratore, non vi è dubbio che essa sia rappresentata dalla gestione dello storage. A questo proposito potrebbe sembrare che lo spazio non sia mai sufficiente, e che sia sempre sovraccarico di attività I/O. Per questo motivo è molto importante avere una certa conoscenza dello storage del disco.

5.1. Panoramica dell'hardware dello storage

Prima di gestire lo storage è molto importante comprendere l'hardware sul quale vengono conservati i dati. Se non avete una buona conoscenza delle funzioni dei dispositivi mass-storage, vi potreste trovare in una situazione in cui vi si presenta un problema riguardante lo storage, ma non avete una conoscenza tale da sapere a cosa sia dovuto. Ottenendo alcune informazioni sul funzionamento hardware, sarete in grado di determinare se il sottosistema storage del vostro computer funzioni in modo corretto.

La maggior parte dei dispositivi mass-storage utilizzano media rotanti, e un accesso di tipo randomico dei dati sui media in questione. Ciò significa che i seguenti componenti sono quasi sempre presenti in quasi tutti i dispositivi di mass-storage:

- Disk platter
- Dispositivo di lettura/scrittura dei dati
- Braccia d'accesso

Le seguenti sezioni affrontano i suddetti componenti in modo più dettagliato.

5.1.1. Disk Platter

I media rotanti utilizzati dalla maggior parte dei dispositivi di mass storage hanno una somiglianza di un vero e proprio piatto circolare. Il piatto può essere composto da diversi materiali, come alluminio, vetro e policarbonato.

La superficie di ogni piatto viene trattata in modo da poter conservare i dati. La natura esatta del trattamento dipende dal tipo di tecnologia usata per la gestione dei dati. La tecnologia più comune è basata sulla proprietà del magnetismo; in questi casi i piatti vengono ricoperti di un composto con caratteristiche magnetiche elevate.

Un'altra tecnologia di conservazione dati molto comune, è basata su principi ottici; in questi casi i piatti sono ricoperti con un materiale che presenta delle proprietà ottiche in grado di essere modificate, permettendo così una conservazione dati seguendo una modalità ottica ¹.

Non ha importanza quale tecnologia utilizzate, i piatti ruotano in modo tale che la loro superficie ruoti oltre un altro componente — il dispositivo di lettura/scrittura dei dati.

5.1.2. Dispositivo di lettura/scrittura dei dati

Il dispositivo di lettura/scrittura dei dati rappresenta il componente che trasforma i bit ed i byte sui quali un sistema computerizzato opera, in elementi ottici o magnetici necessari per interagire con il materiale che ricopre la superficie dei piatti del disco.

1. alcuni dispositivi ottici — le unità CD-ROM — utilizzano approcci diversi per la conservazione dei dati, tali differenze verranno affrontate più in avanti.

Talvolta le condizioni nelle quali i suddetti dispositivi operano sono molto impegnativi. Per esempio in mass-storage magnetici il dispositivo di lettura/scrittura (conosciuto come *testine*), deve essere molto vicino alla superficie del piatto. Tuttavia se la testina e la superficie del disco si toccano, la frizione risultante potrebbe danneggiare severamente sia la testina che il piatto del disco. Per questo motivo, le superfici della testina e quella del disco, devono essere pulite accuratamente, con l'utilizzo da parte della testina di una pellicola d'aria sviluppata dai piatti in modo da sovrastare il disco "flying" ad una distanza minore a quella ricavata dallo spessore di un capello di una persona. Ecco perchè i dischi magnetici sono sensibili all'urto, a variazioni di temperatura, ed altro.

Le difficoltà che si possono incontrare con le testine ottiche sono diverse da quelle che si possono incontrare con l'utilizzo delle testine magnetiche — in questo caso la testina deve rimanere ad una distanza costante dal piatto. In caso contrario, le lenti utilizzate per la messa a fuoco, non sono in grado di produrre una immagine sufficientemente chiara.

In entrambi i casi le testine utilizzano una piccolissima quantità di superficie per la conservazione dei dati. Durante la rotazione del piatto sottostante alle testine, l'area interessata prende una forma di una riga circolare molto sottile.

Se questo è il funzionamento dei dispositivi mass storage, ciò significherebbe che oltre il 99% della superficie del piatto non verrebbe utilizzata. In questo caso si possono montare delle testine aggiuntive, ma per poter usufruire di tutta la superficie in questione, dovremmo montare più di mille testine. Quello che invece risulta essere necessario è una testina in grado di essere posizionata sui diversi punti del piatto.

5.1.3. Braccia d'accesso

Utilizzando una testina collegata ad un braccio in grado di muoversi sopra l'intera superficie del piatto, è possibile sfruttare pianamente la capacità di gestione dei dati del piatto. Tuttavia, il braccio d'accesso deve essere in grado di eseguire quanto segue:

- Essere in grado di muoversi molto velocemente
- Essere molto preciso

Il braccio d'accesso deve muoversi molto rapidamente, poichè il tempo necessario per muovere la testina da un punto ad un altro, rappresenta un tempo morto. Questo perchè durante tale operazione non è possibile leggere o scrivere alcun dato fino a quando il braccio non sia fermo².

Il braccio deve essere in grado di muoversi con molta precisione, in quanto l'area in questione è molto piccola. Per questo motivo, per utilizzare la capacità di conservazione del piatto, è necessario muovere con precisione le testine in modo tale da non sovrascrivere i dati già esistenti. Tale operazione suddivide la superficie del piatto in migliaia di "cerchi" concentrici o *tracce*. Il movimento del braccio da una traccia ad un'altra viene indicato come *ricerca*, ed il tempo relativo a tale ricerca viene chiamato *tempo di ricerca*.

Dove ci sono piatti multipli (o un piatto con entrambe le superfici utilizzate per la conservazione dei dati), le braccia per ogni superficie sono collegate insieme in modo da permettere l'accesso simultaneo della stessa traccia su ogni superficie. Se le tracce su ogni superficie possono essere visualizzate tramite l'accesso attraverso una traccia determinata, esse potrebbero apparire una sull'altra, ottenendo una forma cilindrica, e per questo motivo l'insieme di tracce accessibili in una determinata posizione è conosciuto come un *cilindro*.

2. In alcuni dispositivi ottici (come ad esempio le unità CD-ROM), il braccio d'accesso è in continuo movimento, causando una descrizione di un percorso a spirale attraverso la superficie del piatto. Questo rappresenta una differenza sostanziale sull'uso del mezzo di conservazione dei dati, e riflette le origini del CD-ROM come mezzo di conservazione di musica, dove il recupero continuo dei dati rappresenta una operazione più frequente rispetto alla ricerca specifica di un dato.

5.2. Concetti relativi allo storage

La configurazione dei piatti, delle testine e delle braccia d'accesso rende possibile il posizionamento della testina in qualsiasi punto della superficie di qualsiasi piatto nel dispositivo di mass storage. Tuttavia, ciò non è sufficiente; per utilizzare questa capacità di conservazione, è necessario avere un metodo in grado di conferire gli indirizzi in modo da uniformare le dimensioni delle parti disponibili dello storage.

Vi è un elemento finale appartenente a questo processo. Considerate tutte le tracce presenti nei numerosi cilindri di un dispositivo tipico di mass-storage. Poiché le tracce hanno diametri differenti, le loro circonferenze variano a loro volta. Quindi se lo storage è stato affrontato solo in una prospettiva 'track level', ogni traccia avrà diverse quantità di informazioni — traccia #0 (vicino al centro del piatto) potrebbe contenere 10,827 byte, mentre la traccia #1,258 (vicino al limite del piatto) potrebbe avere 15,382 byte.

La soluzione è dividere ogni traccia in *settori* o *blocchi* multipli con segmenti che hanno una misura costante (spesso la misura è di 512 byte). Il risultato è che ogni traccia contiene un insieme di numeri³ di settori.

Il lato negativo è rappresentato dal fatto che ogni traccia contiene spazio non utilizzato — e cioè lo spazio vicino ai settori. Nonostante il numero costante di settori in ogni traccia, la quantità di spazio non utilizzato varia — una quantità di spazio minore nelle tracce interne, ed uno spazio maggiore nelle tracce esterne. In entrambi i casi, tale spazio non utilizzato viene perso, in quanto i dati non possono essere conservati.

Tuttavia, il vantaggio dell'offsetting di spazio non utilizzato, è rappresentato dal fatto che è possibile affrontare in modo efficiente la conservazione su di un dispositivo mass storage. Sono disponibili due metodi — uno basato sulla geometria e uno basato sul blocco.

5.2.1. Addressing basato sulla geometria

Il termine *metodo basato sulla geometria* si riferisce al fatto che i dispositivi di mass storage conservano i dati in un punto fisico specifico dello storage. Nel caso del dispositivo da noi descritto, esso si riferisce a tre oggetti specifici che definiscono un punto specifico sui piatti del dispositivo:

- Cilindro
- Testina
- Settore

Le seguenti sezioni descrivono come un indirizzo ipotetico sia in grado di descrivere una posizione fisica specifica sullo storage medium.

5.2.1.1. Cilindro

Come precedentemente indicato, il cilindro denota una posizione specifica del braccio d'accesso (e quindi anche le testine di lettura/scrittura). Specificando un cilindro specifico, si possono eliminare i restanti cilindri, riducendo così la nostra ricerca ad una sola traccia per ogni superficie presente nel mass storage.

3. I primi dispositivi di mass storage utilizzavano lo stesso numero di settori per ogni traccia, col passare del tempo, i dispositivi più recenti hanno suddiviso la gamma dei cilindri in diverse "zone," e ogni zona possiede un numero di settori per traccia. Questo permette di avvalersi di uno spazio aggiuntivo tra i settori presenti vicino ai cilindri esterni, dove è maggiore lo spazio non utilizzato.

Cilindro	Testina	Settore
Cilindro	Testina	Settore
1014	X	X

Tabella 5-1. Addressing storage

Nella Tabella 5-1, è stata affrontata la prima parte di un indirizzo basato sulla geometria. Altri due componenti per questo indirizzo — la testina ed il settore — non sono stati ancora definiti.

5.2.1.2. Testina

Anche se selezioniamo un piatto particolare, in quanto ogni superficie possiede una testina di lettura/scrittura, è molto più semplice interagire con una testina specifica. Infatti, l’elettronica del dispositivo è in grado di selezionare una testina e — deselezionare tutto il resto — interagire con la testina selezionata per tutta la durata dell’operazione I/O. Tutte le altre tracce che compongono il cilindro sono state eliminate.

Cilindro	Testina	Settore
1014	2	X

Tabella 5-2. Addressing storage

Nella Tabella 5-2, le prime due parti di un indirizzo basato sulla geometria sono stati riempiti. Un componente finale a questo indirizzo — il settore — non viene definito.

5.2.1.3. Settore

Specificando un settore particolare, si è in grado di completare l’operazione di addressing, identificando in modo unico il blocco desiderato di dati.

Cilindro	Testina	Settore
1014	2	12

Tabella 5-3. Addressing storage

Nella Tabella 5-3, l’indirizzo completo basato sulla geometria è stato riempito. Il suddetto indirizzo identifica la posizione di un blocco specifico rispetto a tutti gli altri blocchi presenti nel dispositivo.

5.2.1.4. Problematiche inerenti l’addressing basato sulla geometria

Anche se l’addressing basato sulla geometria è semplice, vi è un’area in grado di causare alcuni problemi. L’ambiguità presente in questa area è rappresentata dalla numerazione dei cilindri, delle testine e dei settori.

È vero che ogni indirizzo basato sulla geometria identifica in modo unico un blocco specifico di dati, ma questo viene applicato solo se lo schema di numerazione dei cilindri, delle testine e dei settori non viene modificato. Se il suddetto schema varia (come ad esempio se varia l’hardware/software che interagisce con il dispositivo storage), allora la mappatura tra gli indirizzi basati sulla geometria ed i corrispondenti blocchi può variare, rendendo impossibile l’accesso ai dati desiderati.

Per questo motivo è stato sviluppato un approccio diverso. La sezione successiva affronta questo tipo di approccio in modo più dettagliato.

5.2.2. Addressing basato sul blocco

L'*addressing basato sul blocco* è molto più semplice dell'*addressing basato sulla geometria*. Con l'*addressing basato sul blocco*, si è in grado di conferire un numero unico ad ogni blocco. Il suddetto numero viene passato dal computer al dispositivo di mass storage, il quale esegue internamente la conversione in un indirizzo basato sulla geometria, necessario per il circuito di controllo del dispositivo.

Poichè la conversione in un indirizzo basato sulla geometria viene sempre eseguito dal dispositivo stesso, esso è sempre costante, eliminando così il problema relativo al dispositivo relativo all'*addressing basato sulla geometria*.

5.3. Interfacce del dispositivo mass storage

Ogni dispositivo usato in un sistema computerizzato deve avere un modo con il quale si collega ad un determinato sistema computerizzato. Questo punto di collegamento è conosciuto come *interfaccia*. I dispositivi mass storage non sono diversi — anch'essi hanno delle interfacce. È importante conoscere le suddette interfacce per due ragioni:

- Ci sono numerose interfacce (la maggior parte delle quali sono incompatibili)
- Le interfacce presentano prezzi e prestazioni diverse

Sfortunatamente non vi è alcuna interfaccia universale singola del dispositivo tanto meno una interfaccia singola del dispositivo mass storage. Per questo motivo, gli amministratori di sistema devono essere a conoscenza delle interfacce supportate dai sistemi dell'organizzazione. In caso contrario, si può verificare un rischio reale nell'acquisto di hardware non idoneo quando si desidera eseguire un miglioramento del sistema.

Interfacce diverse presentano diverse capacità, rendendo alcune interfacce più idonee per alcuni ambienti invece che per altri. Per esempio, la capacità da parte delle interfacce di supportare dispositivi molto veloci, rendono queste ultime idonee per ambienti server, mentre interfacce più lente saranno sufficienti per un uso desktop. Tali differenze si riflettono anche nel prezzo, ciò significa che — come sempre — la qualità ha un suo costo. L'*high-performance computing* è costoso.

5.3.1. Background storico

Col passare degli anni sono state create diverse interfacce idonee per i dispositivi mass storage. Alcune si sono perse per strada, altre invece sono ancora in uso. Tuttavia, viene riportato il seguente elenco in modo da fornire un'idea sullo scopo dello sviluppo dell'interfaccia attraverso gli ultimi trent'anni, e di fornire una prospettiva chiara sulle interfacce in uso oggi.

FD-400

Interfaccia creata originariamente per le unità floppy disk di 8-pollici negli anni 70. Utilizzo di un cavo di tipo 44-conductor con un connettore in grado di fornire sia alimentazione che dati.

SA-400

Un'altra interfaccia dell'unità floppy disk (questa volta creata originariamente verso il finire degli anni 70 per le unità floppy da 5.25 pollici). Utilizzo di un cavo di tipo 34-conductor con un connettore socket standard. Una versione leggermente modificata di questa interfaccia, è ancora in uso per il floppy da 5.25 pollici e le unità disco da 3.5 pollici.

IPI

Acronimo di Intelligent Peripheral Interface, questa interfaccia è stata usata sulle unità disco da 8 e 14 pollici, impiegate sui minicomputer negli anni 70.

SMD

Un successore di IPI, SMD (stà per Storage Module Device) ed è stato usato sulle unità disco da 8 e 14 pollici dei minicomputer negli anni 70 e 80.

ST506/412

Una interfaccia dell'unità disco che risale agli inizi degli anni 80. Utilizzata nella maggior parte dei personal computer a quei tempi, utilizzava due cavi — uno di tipo 34-conductor e l'altro di tipo 20-conductor.

ESDI

Stà per Enhanced Small Device Interface, questa interfaccia è stata considerata come successore di ST506/412, con una velocità di scambio più elevata e con un supporto maggiore delle dimensioni dell'unità. Risalente alla metà degli anni 80, ESDI utilizza lo stesso schema di connessione a due cavi del suo predecessore.

In quei tempi le interfacce proprietarie erano disponibili tramite i più grandi rivenditori di computer (principalmente IBM e DEC). Lo scopo dietro la creazione di queste interfacce era quello di cercare di proteggere il mercato, molto redditizio, delle periferiche per il loro computer. Tuttavia, a causa della loro natura proprietaria, i dispositivi compatibili con queste interfacce risultavano essere molto più costosi di dispositivi non-proprietari equivalenti. Per questo motivo, queste interfacce non sono riuscite ad affermarsi.

Mentre la maggior parte delle interfacce proprietarie sono scomparse, e le interfacce descritte all'inizio della sezione non hanno alcun valore di mercato, è importante conoscerle, in quanto esse non fanno altro che confermare un punto — niente nel settore informatico resta costante. Per questo motivo, siate sempre attenti alle nuove tecnologie riguardanti le interfacce; potreste trovarne una in grado di essere quella più idonea alle vostre esigenze, rispetto a quella più tradizionale da voi attualmente usata.

5.3.2. Interfacce standard al giorno d'oggi

Diversamente dalle interfacce proprietarie menzionate nella precedente sezione, alcune interfacce hanno avuto un maggiore consenso e quindi sviluppate secondo i criteri standard del settore. Due interfacce in particolare hanno seguito questa transizione, e rappresentano oggi il cuore del settore:

- IDE/ATA
- SCSI

5.3.2.1. IDE/ATA

IDE stà per Integrated Drive Electronics. Questa interfaccia originaria verso la fine degli anni 80, utilizza un connettore di tipo 40-pin.



Nota bene

Il nome corretto di questa interfaccia è "AT Attachment" (o ATA), ma l'utilizzo del termine "IDE" (il quale si riferisce ad un dispositivo mass storage compatibile -ATA) viene in alcuni casi ancora usato. Tuttavia il remainder di questa sezione utilizza il nome corretto — ATA.

ATA implementa una topologia bus, con ogni bus in grado di supportare due dispositivi mass storage. I suddetti dispositivi sono conosciuti come *master* e *slave*. Questi termini possono ingannare, in quanto possono implementare una forma di relazione tra i dispositivi stessi; ma non è il nostro caso. La

scelta del dispositivo master e di quello slave, viene fatta attraverso l'uso dei jumper block su ogni dispositivo.



Nota bene

Una recente innovazione è rappresentata dall'introduzione delle capacità di *selezione del cavo* su ATA. Questo tipo di innovazione necessita l'utilizzo di un cavo speciale, un controller ATA, e di dispositivi mass storage che supportano la selezione del cavo (normalmente attraverso una impostazione jumper della "selezione cavo"). Quando configurato correttamente, la selezione del cavo elimina la necessità di modificare i jumper quando si muovono i dispositivi; la posizione del dispositivo sul cavo ATA, denota se esso sia master oppure slave.

Una variazione di questa interfaccia mostra i diversi modi attraverso i quali le tecnologie possono essere implementate tra loro, altresì introduce la nostra interfaccia standard del settore. *ATAPI* rappresenta una variazione dell'interfaccia ATA, e sta per AT Attachment Packet Interface. Utilizzata principalmente dalle unità CD-ROM, ATAPI è conforme agli aspetti elettrici e meccanici dell'interfaccia ATA, ma utilizza il protocollo di comunicazione dell'interfaccia successiva — SCSI.

5.3.2.2. SCSI

Conosciuta come Small Computer System Interface, SCSI come viene conosciuta oggi, viene ideata nei primi anni 80, ed è stata dichiarata standard nel 1986. Come ATA, SCSI fa un uso della topologia bus. Tuttavia, essi non presentano altre similitudini.

L'utilizzo della topologia bus, significa che ogni dispositivo presente sul bus, deve essere in qualche modo identificato. Mentre ATA supporta solo due dispositivi differenti per ogni bus, dando ad ognuno di essi un nome specifico, SCSI esegue tale operazione assegnando ad ogni dispositivo presente su di un bus SCSI, un indirizzo numerico unico o *SCSI ID*. Ogni dispositivo su di un bus SCSI deve essere configurato (generalmente tramite jumper o per mezzo di interruttori⁴) per rispondere al proprio SCSI ID.

Prima di continuare nella nostra discussione, è importante notare come lo standard SCSI non rappresenta una interfaccia singola, ma una famiglia di interfacce. Vi sono diverse aree nelle quali SCSI varia:

- Larghezza del Bus
- Velocità del Bus
- Caratteristiche elettriche

Lo standard SCSI originale descritto riportava una topologia bus nel quale erano presenti otto righe per il trasferimento dei dati. Questo significa che i primi dispositivi SCSI erano in grado di trasferire i dati ad un byte per volta. Con il passare degli anni, lo standard è stato esteso in modo da permettere alcune implementazioni che permettevano l'uso di sedici righe, raddoppiando così la quantità di dati in grado di essere trasferiti. Le implementazioni SCSI originali "8-bit" venivano identificate come SCSI *narrow*, mentre quelle nuove a 16-bit erano conosciute come SCSI *wide*.

Originariamente, la velocità del bus per SCSI è stata impostata su 5MHz, permettendo di avere una velocità di trasferimento sul bus SCSI a 8-bit originale pari a 5MB/secondo. Tuttavia le successive revisioni dello standard, hanno raddoppiato la velocità fino ad ottenere 10MHz, e cioè 10MB/secondo

4. Alcuni storage hardware (generalmente quelli che rappresentano unità in grado di essere rimosse "carriers") vengono ideati in modo tale che all'atto del collegamento di un modulo, SCSI ID viene impostato automaticamente su di un valore appropriato.

per SCSI narrow e 20MB/secondo per SCSI wide. Come con la larghezza del bus, le modifiche nella velocità del bus hanno dato luogo a nuovi nomi, come ad esempio la velocità del bus di 10MHz chiamato *fast*. Successivi miglioramenti hanno dato luogo all'*ultra* (20MHz), *fast-40* (40MHz), e *fast-80*⁵. Miglioramenti nella velocità di trasferimento hanno portato alla creazione di diverse versioni, come ad esempio *ultra160* bus speed.

Combinando questi termini è possibile indicare diverse configurazioni SCSI. Per esempio, "ultra-wide SCSI" si riferisce ad un bus SCSI a 16-bit eseguito a 20MHz.

Lo standard SCSI originale utilizzava un *single-ended* signaling; costituito da una configurazione elettrica dove solo un conduttore è stato utilizzato per passare un segnale elettrico. Implementazioni successive hanno permesso l'utilizzo di un signaling *differenziato*, dove vengono usati due conduttori per passare i segnali. Lo SCSI differenziale (il quale è stato rinominato *high voltage differential* o SCSI HVD), possiede il beneficio di ridurre la sensibilità al rumore elettrico e permette una maggiore lunghezza del cavo, esso non è mai diventato comune nel campo informatico. Una successiva implementazione, conosciuta come *low voltage differential* (LVD), ha avuto molto più successo diventando uno dei requisiti per i bus speed più elevati.

La larghezza di un bus SCSI non solo indica la quantità di dati in grado di essere trasferiti ad ogni ciclo, ma determina anche il numero di dispositivi che possono essere collegati ad un bus. SCSI normali supportano 8 dispositivi indirizzati in modo unico, mentre gli SCSI wide ne supportano 16. In entrambi i casi, dovete assicurarvi che tutti i dispositivi siano impostati per poter utilizzare uno SCSI ID unico. Due dispositivi che condividono un ID singolo, possono causare dei problemi che a loro volta portano ad una corruzione dei dati.

Un'altra cosa da ricordare è che *ogni* dispositivo presente sul bus utilizza un ID. Ciò include lo SCSI controller. Spesso, purtroppo, gli amministratori di sistema dimenticano tutto questo, impostando inavvertitamente il dispositivo in modo da utilizzare lo stesso SCSI ID come controller del bus. Ciò significa che, in pratica, solo 7 (o 15, per SCSI wide) dispositivi possono essere presenti in un bus singolo, in quanto ogni bus deve riservare un ID per il controller.



Suggerimento

Molte implementazioni SCSI sono in grado di eseguire uno scanning del bus SCSI; ciò viene regolarmente usato per confermare che tutti i dispositivi siano configurati correttamente. Se la scansione del bus ritorna lo stesso dispositivo per ogni SCSI ID singolo, ciò significherebbe che quel dispositivo è stato impostato incorrettamente per lo stesso SCSI ID dello SCSI controller. Per risolvere questo problema, riconfigurare il dispositivo in modo da usare uno SCSI ID diverso (unico).

A causa dell'architettura del bus SCSI, è necessario *terminare* correttamente entrambe le estremità del bus. Tale terminazione può essere eseguita posizionando il carico dell'impedenza elettrica corretta, su ogni conduttore che costituisce il bus SCSI. Tale operazione rappresenta un requisito elettrico; senza di esso, i diversi segnali presenti sul bus, saranno riportati sulle estremità del bus, ingarbugliando così tutte le comunicazioni.

Molti (ma non tutti) dispositivi SCSI hanno dei terminatori interni che possono essere abilitati o disabilitati utilizzando jumper o interruttori. Sono anche disponibili dei terminatori esterni.

Un'ultima cosa da ricordare nei confronti di SCSI — esso non rappresenta solo una interfaccia standard per i dispositivi mass storage. Molti altri dispositivi (come gli scanner, le stampanti, ed i dispositivi di comunicazione) utilizzano SCSI. Anche se essi non sono molto numerosi dei dispositivi mass storage SCSI, è bene sapere della loro esistenza. È da tener presente che con l'avvento di USB e

5. Fast-80 tecnicamente non rappresenta una modifica in velocità; conserva il bus 40MHz, ma i dati sono stati programmati in modo tale che sia durante il picco che durante la caduta dell'impulso essi possano essere trasmessi, raddoppiando così la velocità di trasferimento dei dati.

IEEE-1394 (spesso chiamati Firewire), queste interfacce verranno utilizzate più spesso per questi tipi di dispositivi.



Suggerimento

Le interfacce USB e IEEE-1394 si stanno facendo strada nel mondo dei dispositivi mass storage; tuttavia, non esistono dispositivi nativi mass storage di tipo USB o IEEE-1394. Sono invece molto diffusi al giorno d'oggi i dispositivi ATA o SCSI con circuito di conversione esterno.

Non ha importanza quale interfaccia viene utilizzata dal dispositivo mass storage, le caratteristiche principali del dispositivo si basa sulla propria prestazione. Le seguenti sezioni affrontano i passi più importanti.

5.4. Caratteristiche delle prestazioni dei dischi fissi

Le caratteristiche della prestazione del disco fisso sono state precedentemente affrontate nella Sezione 4.2.4; questa sezione affronta tale argomento in dettaglio. La suddetta fase è molto importante per gli amministratori di sistema, in quanto senza una conoscenza di base sul funzionamento dei dischi fissi, è possibile eseguire delle modifiche nei confronti della configurazione del vostro sistema, in grado di influenzare in modo negativo le proprie prestazioni.

Il tempo necessario ad un disco fisso per rispondere e completare una richiesta I/O dipende da due cose:

- Dai limiti elettrici e meccanici del disco fisso
- Dal carico I/O imposto dal sistema

Le seguenti sezioni esplorano questi aspetti inerenti la prestazione del disco fisso in modo più dettagliato.

5.4.1. Limitazioni elettriche/meccaniche

Poichè i dischi fissi sono dispositivi elettro-meccanici, essi sono soggetti a diverse limitazioni sulla loro velocità e prestazione. Ogni richiesta I/O, per poter funzionare insieme, necessita di diversi componenti del drive, in modo da soddisfare la richiesta stessa. Poichè ogni componente presenta diverse caratteristiche, la prestazione generale del disco fisso viene determinata dalla somma della prestazione dei componenti individuali.

Tuttavia, i componenti elettrici risultano essere più veloci rispetto ai componenti meccanici. Per questo motivo, i componenti meccanici risultano avere un impatto maggiore sulle prestazioni generali del disco fisso.



Suggerimento

Il modo migliore per migliorare le prestazioni del disco fisso è quello di ridurre il più possibile l'attività meccanica dell'unità.

Il tempo di accesso medio di un disco fisso è di circa 8.5 millisecondi. Le seguenti sezioni analizzano questa figura in modo più dettagliato, mostrando come ogni componente influisca sulle prestazioni generali del disco fisso.

5.4.1.1. Tempo di processazione del comando

Tutti i dischi fissi di moderna concezione possiedono dei sistemi computerizzati embedded molto sofisticati in grado di controllare le loro funzioni. Essi sono in grado di eseguire i seguenti compiti:

- Interagire con il mondo esterno attraverso l'interfaccia del disco fisso
- Controllare le operazioni del resto dei componenti facenti parte del disco fisso, e ripristinare le normali funzioni dopo che si siano verificate alcune condizioni d'errore.
- Processare i dati letti e scritti sul media dello storage corrente

Anche se i microprocessori utilizzati nei dischi fissi sono relativamente potenti, i compiti a loro assegnati richiedono del tempo per la loro processazione. In media, il tempo necessario è di 003 millisecondi.

5.4.1.2. Testine di lettura/scrittura dei dati

Le testine di lettura/scrittura del disco fisso funzionano solo quando i piatti del disco sono in rotazione. Poichè la lettura e la scrittura dei dati dipende dal movimento del media posizionato sotto la testina, il tempo necessario al media contenente il settore desiderato, a passare completamente sotto la testina, risulta essere l'unico componente in grado di determinare il tempo di accesso totale. Esso è in generale pari a .0086 millisecondi per una unità di 10,000 RPM con 700 settori per traccia.

5.4.1.3. Tempo di rotazione

Poichè i disk platter del disco fisso sono in continuo movimento, quando arriva una richiesta I/O, risulterà improbabile che il piatto si troverà esattamente nella posizione corrispondente al settore desiderato. Per questo motivo, anche se il resto dell'unità è pronta ad accedere il settore desiderato, è necessario attendere la rotazione del piatto, portando il suddetto settore nella posizione giusta e cioè sotto la testina di lettura/scrittura.

Questo è il motivo per il quale i dischi fissi con elevate prestazioni, generalmente ruotano i loro piatti ad una velocità maggiore. Al giorno d'oggi una velocità pari a 15,000 RPM è idonea per le unità con elevate prestazioni, mentre 5,400 RPM è idonea per unità di tipo entry-level. Ciò ne consegue che la media è di circa 3 millisecondi per una unità di 10,000 RPM.

5.4.1.4. Movimento del braccio d'accesso

Se vi è un componente nel disco fisso che può essere considerato il punto debole, esso è rappresentato dal braccio d'accesso. Il motivo è che il braccio d'accesso si muove molto rapidamente e in modo preciso, su distanze relativamente molto lunghe. In aggiunta, il movimento del braccio d'accesso non è continuo — infatti deve accelerare molto rapidamente per avvicinarsi al cilindro desiderato per poi decelerare per evitarlo di superarlo. Quindi il braccio d'accesso deve essere resistente (per sopravvivere alle forze alle quali viene sottoposto a causa dei movimenti brevi e veloci), ma anche leggero (in questo modo vi è meno massa da accelerare/decelerare).

Raggiungere questi obiettivi è difficile, infatti il movimento del braccio d'accesso è relativamente maggiore rispetto al tempo utilizzato da altri componenti. Per questo, il movimento del braccio d'accesso risulta essere determinante sulla prestazione generale del disco fisso, aggirandosi intorno a 5.5 millisecondi.

5.4.2. Prestazione e carico I/O

Un altro fattore in grado di controllare la prestazione del disco fisso è il carico I/O al quale il disco fisso stesso viene sottoposto. Alcuni degli aspetti più specifici inerenti il carico I/O sono:

- La quantità di dati letti rispetto a quelli scritti
- Il numero corrente dei lettori/scrittori
- La posizione delle letture/scritture

Questo viene affrontato in modo più dettagliato nelle seguenti sezioni.

5.4.2.1. Letture contro scritture

Per l'unità disco generica che utilizza un media magnetico per la conservazione dei dati, il numero delle operazioni I/O di lettura rispetto al numero delle operazioni I/O di scrittura non risulta essere un problema, in quanto la lettura e la scrittura dei dati richiedono la stessa quantità di tempo⁶. Tuttavia, altre tecnologie mass storage necessitano di una quantità di tempo diverso per processare la lettura e la scrittura⁷.

L'impatto che ne consegue è che i dispositivi che necessitano di maggior tempo per processare le operazioni I/O di scrittura (per esempio), sono in grado di gestire un minor numero I/O di scrittura rispetto agli I/O di lettura. Se ci poniamo un'altra prospettiva, un processo I/O di scrittura necessita di una maggiore abilità del dispositivo a processare le richieste I/O rispetto al processo I/O di lettura.

5.4.2.2. Lettori/scrittori multipli

Un disco fisso che processa le richieste I/O provenienti da fonti multiple, presenta un carico diverso rispetto al disco fisso che fa fronte alle richieste I/O provenienti solo da una fonte. La ragione principale per questo è data dal fatto che le richieste I/O multiple hanno il potenziale di provvedere a carichi I/O più elevati, riguardanti un disco fisso invece di una richiesta I/O singola.

Questo perchè un richiedente I/O deve eseguire un certo numero di processazioni, prima che si possa verificare una operazione I/O. Dopo tutto, il richiedente deve determinare la natura della richiesta I/O prima che la stessa possa essere eseguita. Poichè la processazione necessaria per determinare quanto detto richiede del tempo, esiste un limite superiore sul carico I/O che un richiedente è in grado di generare — solo una CPU più veloce è in grado di elevare il suddetto limite. Questa limitazione è maggiormente evidenziata se il richiedente necessita un input umano prima di eseguire un I/O.

Tuttavia, con richieste multiple, i carichi I/O più elevati possono essere sostenuti. Fino a quando è disponibile una alimentazione CPU sufficiente a supportare il processo necessario a generare le richieste I/O, l'aggiunta di più richieste I/O aumenta il carico I/O risultante.

Tuttavia, vi è un altro fattore in grado di essere collegato al carico I/O risultante. Tale fattore viene affrontato nella seguente sezione.

6. Questo non risulta essere veritiero. Tutti i dischi fissi includono una certa quantità di memoria cache interna, che viene usata per migliorare la prestazione di lettura. Tuttavia ogni richiesta I/O di lettura dei dati, deve essere soddisfatta leggendo i dati dal mezzo utilizzato per lo storage. Ciò significa che, mentre la cache è in grado di alleviare i problemi inerenti la prestazione I/O di lettura, risulta impossibile eliminare il tempo necessario per leggere fisicamente i dati dal mezzo utilizzato per lo storage.

7. Alcune unità disco ottiche presentano questo comportamento, a causa della limitazione fisica delle tecnologie utilizzate per implementare la conservazione ottica dei dati.

5.4.2.3. Posizione delle letture/scritture

Anche se non forzato da un ambiente multi-richiedente, l'aspetto della prestazione del disco fisso tende ad essere più accentuato in tali ambienti. Il problema è se le richieste I/O fatte da un disco fisso, siano per dati fisicamente vicini ad altri dati già richiesti.

Il motivo per il quale esso risulta importante diventa apparente se si tiene in considerazione la natura elettromeccanica del disco fisso. La componente più lenta di qualsiasi disco fisso è rappresentata dal braccio d'accesso. Per questo motivo, se i dati interessati dalle richieste I/O in entrata non necessitano di alcun movimento da parte del braccio d'accesso, il disco fisso è in grado di servire molte più richieste I/O rispetto a quando i dati in questione si trovano in diverse posizioni nel drive, e quindi richiedono un movimento esteso da parte del braccio d'accesso stesso.

Ciò può essere illustrato guardando alle specifiche sulla prestazione del disco fisso. Queste specificazioni includono spesso i *tempi di ricerca adjacent cylinder* (dove il braccio d'accesso viene mosso lievemente — solo sul cilindro successivo), e *tempi di ricerca full-stroke* (dove il braccio d'accesso si muove dal primissimo cilindro fino all'ultimo). Per esempio, ecco i tempi di ricerca per un disco fisso ad alte prestazioni:

Adjacent Cylinder	Full-Stroke
0.6	8.2

Tabella 5-4. Tempi di ricerca Adjacent Cylinder e Full-Stroke (in Millisecondi)

5.5. Rendere utilizzabile lo storage

Una volta implementato un dispositivo mass storage lo si può utilizzare per pochi compiti. È vero, i dati possono essere scritti e letti, ma senza una struttura, l'accesso ai dati diventa possibile solo utilizzando gli indirizzi del settore (sia geometrico che logico).

Sono necessari metodi in grado di trasformare un raw storage in un disco fisso, questo processo fornisce degli 'usable' in modo più semplice. Le seguenti sezioni affrontano alcune tecniche molto comuni per eseguire tale processo.

5.5.1. Partizioni/Sezioni

La prima cosa che spesso si presenta ad un amministratore è quello della misura di un disco fisso, in alcuni casi essa risulta essere molto più grande del necessario. Come risultato, molti sistemi operativi hanno la capacità di dividere lo spazio di un disco fisso in diverse *partizioni* o *sezioni*.

Poichè esse sono separate le une dalle altre, le partizioni possono avere quantità diverse di spazio utilizzato, e lo stesso spazio non deve influenzare in alcun modo lo spazio utilizzato da altre partizioni. Per esempio, la partizione che gestisce i file che costituiscono il sistema operativo, non viene influenzata anche se la partizione che gestisce i file dell'utente è piena. Il sistema operativo ha ancora spazio disponibile al proprio uso.

Anche se potrebbe risultare alquanto semplicistico, pensate alle partizioni come se fossero simili alle unità disco individuali. Infatti alcuni sistemi operativi fanno riferimento alle partizioni come se fossero delle "unità". Tuttavia, questa prospettiva non è molto accurata; risulta quindi essere importante fare riferimento alle partizioni in modo più dettagliato.

5.5.1.1. Attributi della partizione

Le partizioni vengono definite dai seguenti attributi:

- Geometria della partizione
- Tipo di partizione
- Campo del tipo di partizione

Questi attributi saranno affrontati in dettaglio nelle seguenti sezioni.

5.5.1.1.1. Geometria

La geometria di una partizione fa riferimento al luogo fisico su di una unità disco. La suddetta geometria può essere intesa in termini di inizio e fine dei cilindri, delle testine, e dei settori, anche se spesso molte partizioni iniziano e finiscono ai limiti del cilindro. La misura di una partizione viene definita come la quantità di storage tra i cilindri di inizio e di fine.

5.5.1.1.2. Tipo di partizione

Il tipo di partizione si riferisce al rapporto che ha la partizione con le altre partizioni presenti sul disco fisso. Ci sono tre tipi di partizione:

- Partizioni primarie
- Partizioni estese
- Partizioni locali

Le seguenti sezioni descrivono ogni tipo di partizione.

5.5.1.1.2.1. Partizioni primarie

Le partizioni primarie sono quelle partizioni che richiedono uno dei quattro alloggiamenti delle partizioni primarie, presenti nella tabella della partizione del disco fisso.

5.5.1.1.2.2. Partizioni estese

Le partizioni estese sono state sviluppate a causa della necessità di avere più di quattro partizioni per disco fisso. Una partizione estesa è in grado di contenere a sua volta delle partizioni multiple, aumentando così il numero di partizioni presenti su di una unità singola. L'introduzione delle partizioni estese è stata resa necessaria dall'aumento delle capacità delle nuove unità disco.

5.5.1.1.2.3. Partizioni logiche

Le partizioni logiche sono quelle partizioni presenti all'interno di una partizione estesa; nel loro utilizzo, esse non sono diverse da una partizione primaria non-estesa.

5.5.1.1.3. Campo del tipo di partizione

Ogni partizione possiede un campo in grado di contenere un codice che indica l'utilizzo anticipato della partizione. Il suddetto campo potrebbe, o meno, riflettere il sistema operativo del computer. Esso potrebbe altresì riflettere il processo con il quale vengono conservati i dati all'interno della partizione. La seguente sezione contiene maggiori informazioni sul suddetto argomento.

5.5.2. File System

Anche con la presenza del dispositivo mass storage corretto, configurato e partizionato in modo efficiente, noi non saremo ancora in grado di conservare e richiamare le informazioni in modo semplice — ciò che manca è il modo di strutturare e organizzare le informazioni. Avremo quindi bisogno di un *file system*.

Il concetto di file system è così importante per l'utilizzo dei dispositivi mass storage, che l'utente normale spesso non fa alcuna differenza tra i due. Tuttavia, gli amministratori di sistema non possono permettersi di ignorare i file system ed il loro impatto nello svolgimento del lavoro quotidiano.

Il file system non è altro che un metodo di rappresentazione dei dati su di un dispositivo mass storage. I file system generalmente includono i seguenti contenuti:

- Data storage basato sul file
- Struttura della directory gerarchica (talvolta conosciuta come "cartella")
- Cronologia della creazione del file, dell'accesso e dei tempi di modifica
- Alcuni livelli di controllo sul tipo di accesso permesso ad un file specifico
- Alcuni concetti di file ownership
- Controllo dello spazio utilizzato

Non tutti i file system possiedono le suddette caratteristiche. Per esempio, un file system creato per un sistema operativo con utente singolo, utilizza un metodo di controllo d'accesso più semplice, offrendo anche un supporto per il file ownership.

Una cosa da ricordare è che il file system utilizzato, può influenzare la natura del vostro lavoro giornaliero. Assicurandovi che il file system utilizzato all'interno della vostra organizzazione sia idoneo a far fronte le necessità dell'organizzazione stessa, ne può risultare una gestione più semplice ed efficiente.

Tenendo a mente questa caratteristica, le seguenti sezioni affrontano le suddette caratteristiche in modo più dettagliato.

5.5.2.1. Storage basato sul file

Anche se i file system che utilizzano il file metaphor per lo storage dei dati sono quasi universali, ci sono ancora alcuni aspetti da tenere in considerazione.

Per prima cosa è importante sapere quali sono le restrizioni sui file name. Per esempio, sapere quali sono i caratteri permessi, la lunghezza massima dei file name ecc. Queste domande sono importanti, in quanto esse stabiliscono i file name da usare. Sistemi operativi più obsoleti, con file system più vecchi, spesso permettono l'uso solo di caratteri alfa numerici (ed in questo caso solo con lettere maiuscole), e solo con l'utilizzo dei file name 8.3 tradizionali (ciò significa un file name a otto caratteri, seguito da un file extension di tre caratteri).

5.5.2.2. Struttura della directory gerarchica

Anche se i file system in uso su alcuni sistemi operativi molto vecchi non includono il concetto di directory, tutti i file system oggi comunemente usati, includono questa caratteristica. Le directory vengono generalmente impiegate come dei file, ciò significa che non è necessario l'utilizzo di utility particolari per la loro gestione.

Altresì, poichè le directory sono per conto loro dei file in grado di contenere i file stessi, le directory possono contenere altre directory, rendendo possibile la costituzione di una loro gerarchia multi-livello. Questo rappresenta un concetto molto sviluppato, con il quale tutti gli amministratori di sistema devono avere una certa familiarità. Utilizzando le suddette gerarchie multi-livello, risulta più semplice per voi e per i vostri utenti gestire un file.

5.5.2.3. Cronologia della creazione del file, dell'accesso e dei tempi di modifica

La maggior parte dei file system sono in grado di conservare alcune informazioni, come ad esempio il momento della creazione di un file; altri invece sono in grado di conservare i tempi di modifica e di accesso. Dato per scontato che è molto utile essere in grado di determinare il tempo di creazione, di modifica e di accesso di un dato file, questi dati risultano essere importanti per l'operazione corretta per l'Incremental Backup.

Per maggiori informazioni su come i backup utilizzano le caratteristiche dei file system, consultate la Sezione 8.2.

5.5.2.4. Controllo dell'accesso

Il controllo dell'accesso rappresenta una delle aree dove i file system differiscono in modo drammatico tra loro. Alcuni file system non hanno alcun modello chiaro di controllo dell'accesso, mentre altri sono molto più sofisticati. In termini generali, molti file system moderni sono in grado di combinare due componenti in una metodologia coesiva di controllo:

- User identification
- Elenco di azioni abilitate

User identification significa che il file system (e il sistema operativo in questione) deve essere in grado di identificare in modo unico i singoli utenti. Ciò rende possibile essere responsabili delle operazioni nel livello del file system. Un'altra caratteristica molto utile è quella degli *user group* — creando un insieme corretto di utenti. I gruppi sono spesso utilizzati dalle organizzazioni dove gli utenti possono essere membri di uno o più progetti. Un'altra caratteristica supportata da alcuni file system è quella della creazione di identificatori generici, in grado di essere assegnati ad uno o più utenti.

Successivamente il file system deve essere in grado di mantenere un elenco di azioni che possono essere eseguite (o meno) per ogni file. Le più comuni sono:

- Lettura del file
- Scrittura del file
- Esecuzione del file

Alcuni file system sono in grado di estendere l'elenco in modo da includere azioni come ad esempio la cancellazione, o l'abilità di eseguire delle modifiche relative al controllo dell'accesso del file.

5.5.2.5. Controllo dello spazio utilizzato

Una costante presente nella vita di un amministratore di sistema è quella di non avere mai uno spazio sufficiente, ma se il suddetto spazio risulta essere presente, esso non resterà disponibile per molto tempo. Per questo un amministratore deve essere sempre in grado di determinare facilmente, il livello di spazio disponibile per ogni file system. In aggiunta, i file system che possiedono una capacità di identificazione dell'utente molto definita, spesso includono la possibilità di visualizzare la quantità di spazio consumato da un particolare utente.

Questa caratteristica risulta essere vitale in ambienti con diversi utenti, in quanto la regola 80/20 spesso viene applicata allo spazio del disco — 20 percento di utenti saranno responsabili del consumo dell'80 percento dello spazio disponibile. Facilitando la determinazione degli utenti che fanno parte del 20 percento, sarete in grado di gestire più facilmente gli asset relativi al vostro storage.

Andando oltre, alcuni file system includono la possibilità di impostare dei limiti per ogni utente (spesso conosciuti come *disk quota*) sulla quantità di spazio da poter consumare. Le specifiche variano a seconda del file system, ma in generale ad ogni utente può essere assegnata una quantità specifica di storage da utilizzare. Altri file system invece, permettono all'utente di eccedere il loro limite solo una

volta, mentre altri implementano un "grace period" durante il quale viene applicato un secondo limite, ma in questo caso più elevato.

5.5.3. Struttura della directory

Molti amministratori di sistema non si soffermano molto sull'utilizzo dello storage disponibile ai propri utenti. Tuttavia, soffermarsi un pò più a lungo su questo argomento, prima di rendere disponibile lo storage agli utenti, potrebbe evitarvi una perdita di tempo più in avanti.

La cosa migliore che gli amministratori possono fare, è quello di utilizzare delle directory e delle subdirectory per strutturare lo storage disponibile in modo da poter essere facilmente capito. Questo approccio riserva diversi benefici:

- Facilmente comprensibile
- Maggiore flessibilità in futuro

Imponendo alcuni livelli di struttura sul vostro storage, esso potrebbe essere capito più facilmente. Per esempio, considerate un sistema multi-utente. Invece di posizionare le directory di tutti gli utenti, all'interno di una directory più grande, potrebbe avere più senso utilizzare delle subdirectory che riflettono la struttura della vostra organizzazione. In questo modo, il personale che lavora nel settore della contabilità, avrà le proprie directory sotto una unica directory chiamata `accounting`, il personale che lavora nella sezione di ingegneria, avrà le proprie directory sotto `engineering`, e così via.

I benefici di tale approccio sono nella semplicità di controllo, in base giornaliera, dei requisiti dello storage (e del suo utilizzo), da parte di ogni singola sezione della vostra organizzazione. Ottenere un elenco dei file utilizzati dal personale facente parte alla `human resources`, è molto semplice. Risulta altresì molto semplice, l'esecuzione del backup di tutti i file utilizzati per esempio dall'ufficio legale.

Con una struttura adeguata, si è in grado di aumentare la flessibilità. Per continuare ad usare l'esempio precedente, assumete per un momento che la sezione di ingegneria sia in procinto di iniziare una serie di nuovi progetti. Per questo motivo, risulta necessario assumere numerosi nuovi ingegneri. Tuttavia, non risulta esserci spazio sufficiente nello storage, per poter supportare tale aggiunta di personale.

Tuttavia, poichè ogni persona nella sezione di ingegneria possiede i propri file sotto la directory `engineering`, potrebbe risultare semplice:

- Procurare dello storage aggiuntivo necessario per il supporto
- Eseguire il backup di ogni cosa sotto la directory `engineering`
- Ripristinare il backup sul nuovo storage
- Rinominare la directory `engineering` presente sullo storage originale, in qualcosa di simile a `engineering-archive` (prima di cancellarla interamente, dopo il suo uso per un mese con la nuova configurazione)
- Eseguire le modifiche necessarie in modo tale che il personale facente parte della sezione di ingegneria, possa accedere ai propri file presenti sul nuovo storage

Naturalmente tale approccio ha i suoi lati negativi. Per esempio, se il personale viene spesso trasferito tra le diverse sezioni, dovete essere in grado di esserne informati al più presto, e quindi sarà necessario modificare la struttura della directory in modo appropriato. In caso contrario tutto questo si tradurrà in maggiore — e non minore — lavoro.

5.5.4. Abilitare l'accesso allo storage

Una volta partizionato in modo corretto il dispositivo mass storage, e quindi aver scritto su di esso un file system, lo storage stesso sarà disponibile per un suo utilizzo generale.

Per alcuni sistemi operativi, quando viene rilevato il nuovo dispositivo mass storage, esso può essere formattato dall'amministratore di sistema in modo che il suo accesso avvenga senza particolari sforzi.

Altri sistemi operativi invece richiedono una fase aggiuntiva. Questa fase — spesso viene riferita come *mounting* — e cioè indica al sistema operativo come accedere allo storage. Il Mounting storage viene normalmente eseguito attraverso un programma di utility speciale o un comando, e necessita che il dispositivo mass storage (e possibilmente anche la partizione) venga esplicitamente identificato.

5.6. Tecnologie avanzate per lo storage

Anche se fino ad ora questo capitolo ha affrontato solo unità disco singole collegate direttamente ad un sistema, è bene sapere che sono anche disponibili altre opzioni più avanzate. Le seguenti sezioni descrivono alcuni degli approcci più comuni, in modo da ampliare le vostre opzioni riguardanti i mass storage.

5.6.1. Storage accessibile tramite la rete

La combinazione delle tecnologie mass storage e della rete può risultare in una maggiore flessibilità per gli amministratori di sistema. Questa configurazione permette di avere due possibili benefici:

- Consolidazione dello storage
- Gestione semplificata

Lo storage può essere consolidato impiegando server ad alte prestazioni con una connettività di rete molto elevata, e configurato con grosse quantità di fast storage. Data una configurazione corretta, è possibile fornire un accesso allo storage con delle velocità paragonabili allo storage collegato in modo locale. Altresì, la natura condivisa di tale configurazione, spesso rende possibile ridurre i costi, in quanto le spese associate nel fornire uno storage condiviso e centralizzato, possono essere minori rispetto alla fornitura di storage equivalente per ogni client. In aggiunta, si viene a consolidare dello spazio libero, invece di essere diffuso (e non utilizzabile) attraverso molti client.

I server storage centralizzati possono facilitare molti compiti amministrativi. Per esempio, il controllo dello spazio disponibile è più semplice quando lo storage da controllare è presente su di un sistema. I backup possono essere semplificati; è possibile eseguire un backup del tipo network-aware, ma tale tipo di backup necessita di un maggiore lavoro per eseguire la configurazione e la manutenzione.

Sono disponibili un certo numero di tecnologie storage networked; sceglierne una potrebbe risultare difficile. Quasi ogni sistema operativo presente nel mercato odierno, presenta dei mezzi di accesso ad uno storage accessibile attraverso la rete, le diverse tecnologie però, potrebbero risultare incompatibili tra di loro. Quindi, qual'è il miglior approccio per la scelta della tecnologia da impiegare?

L'approccio che generalmente fornisce il miglior risultato, è quello di lasciar decidere il problema alle capacità interne del client. Vi sono un certo numero di ragioni per questo:

- Problemi minimi sulla integrazione del client
- Lavoro minimo su ogni sistema del client
- Costo più basso della entry per-client

Tenete conto che qualsiasi problema riguardante il client, viene moltiplicato per il numero di client presenti nella vostra organizzazione. Utilizzando le capacità interne del client stesso, non avrete biso-

gno d'installare alcun software aggiuntivo nei client (ciò significa nessun costo addizionale). Tutto ciò vi permetterà di avere una buona integrazione e un buon supporto con il sistema operativo del client.

Vi è comunque un lato negativo. Cioè l'ambiente server deve essere in grado di fornire un buon supporto per le tecnologie storage, accessibili tramite la rete e necessarie per i client. In casi in cui i sistemi operativi del client e del server sono unici e uguali, normalmente non vi è alcun problema. In caso contrario, sarà necessario concentrarsi nel far sì che il server "parli" la stessa lingua del client. Tuttavia, questo compromesso è più che giustificato.

5.6.2. Storage basato su RAID

Una capacità che un amministratore dovrebbe coltivare, è l'abilità di orientarsi a configurazioni più complesse, e osservare i diversi aspetti inerenti ad ogni configurazione. Anche se a prima vista questo approccio potrebbe sembrare un pò demoralizzante, esso potrebbe essere d'aiuto per affrontare i problemi che si potrebbero verificare nei momenti meno opportuni, e quindi causare danni sostanziali.

Tenendo presente quanto detto, usiamo la nostra conoscenza per quanto riguarda lo storage basato sul disco, e vediamo se possiamo determinare i modi in cui l'unità disco può causare dei problemi. Innanzitutto consideriamo un problema riguardante l'hardware:

Una unità disco con quattro partizioni viene completamente a mancare: cosa succede ai dati presenti su quelle partizioni?

Tale unità non è disponibile (fino al momento in cui l'unità che presenta il problema non viene sostituita, ed i dati non vengono ripristinati tramite backup).

Una unità disco con una partizione singola stà operando al limite delle sue capacità a causa di carichi I/O molto elevati: che cosa succede alle applicazioni che necessitano di accedere ai dati presenti su quella partizione?

Le applicazioni rallentano perchè l'unità disco non è in grado di eseguire le procedure di lettura e scrittura più velocemente.

Se avete un file data molto grande che è in continua crescita; esso potrebbe divenire presto, più grande dell'unità disco maggiore disponibile per il vostro sistema. Cosa potrebbe succedere a questo punto?

L'unità disco diventa satura, il file data interrompe la sua crescita, e le applicazioni associate interrompono la propria esecuzione.

Proprio uno di questi problemi potrebbe intaccare ed invalidare il centro dati, purtroppo gli amministratori di sistema devono far fronte ogni giorno a questi problemi. Che cosa si potrebbe fare in questi casi?

Fortunatamente è disponibile una tecnologia in grado di affrontare queste problematiche. Il nome della suddetta tecnologia è **RAID**.

5.6.2.1. Concetti di base

RAID è l'acronimo di Redundant Array of Independent Disks⁸. Come implica il nome, RAID è un modo per le unità disco multiple, di comportarsi come delle unità disco singole.

Le tecniche RAID sono state sviluppate da alcuni ricercatori dell'Università della California, Berkeley nella metà degli anni 80. In quel periodo vi era un grosso divario tra le unità disco ad alte prestazioni, utilizzate nelle maggiori installazioni di computer, e le più piccole, e lente unità disco utilizzate nel settore del personal computer a quei tempi ancora molto giovane. RAID rappresentava a quei tempi, un metodo grazie al quale erano disponibili diverse unità disco meno costose, al posto di una unità più costosa.

8. Quando iniziarono le prime ricerche su RAID, l'acronimo era Redundant Array of *Inexpensive* Disks, ma con il passare del tempo i dischi "standalone", che RAID doveva sostituire, sono diventati sempre meno costosi, rendendo il paragone tra i due sempre meno interessante.

I RAID arrays potevano essere creati in modi diversi, riuscendo ad avere caratteristiche diverse a seconda della configurazione finale. Vediamo in dettaglio le diverse configurazioni (conosciute come *livelli RAID*).

5.6.2.1.1. Livelli RAID

I ricercatori di Berkeley definirono originariamente cinque diversi livelli RAID, numerandoli da "1" a "5." Successivamente sono stati aggiunti altri livelli RAID da parte di altri ricercatori e membri dello storage industry. Non tutti i livelli RAID si sono mostrati utili allo stesso modo, alcuni erano utili solo nel settore delle ricerche, altri invece erano troppo costosi per essere implementati.

Alla fine soltanto tre livelli RAID si sono rivelati, in un contesto generale, utili:

- Level 0
- Level 1
- Level 5

Le seguenti sezioni affrontano i suddetti livelli in modo più dettagliato.

5.6.2.1.1.1. RAID 0

La configurazione del disco conosciuta come RAID level 0 può trarre in inganno, in quanto esso rappresenta il solo livello RAID che non impiega alcuna ridondanza. Tuttavia, anche se RAID 0 non presenta alcun vantaggio dal punto di vista dell'affidabilità, esso presenta altri tipi di benefici.

Un RAID 0 array consiste di due o più unità disco. La capacità storage disponibile su ogni unità è suddivisa in diversi *blocchi*, i quali rappresentano i multipli della misura nativa del blocco delle unità. I dati scritti sull'array devono essere anche scritti, blocco per blocco, su ogni unità nell'array stesso. I blocchi possono essere concepiti come delle linee attraverso ogni unità presente nell'array; quindi ne deriva l'altro termine per RAID 0: *striping*.

Per esempio, con un array a due-unità ed una misura del blocco di 4KB, la scrittura di 12KB di dati sull'array, ne risulterebbe in una scrittura dei dati in blocchi da 4KB sulle seguenti unità:

- I primi 4KB saranno scritti sulla prima unità, nel primo blocco
- I secondi 4KB saranno scritti nella seconda unità, nel primo blocco
- Gli ultimi 4KB saranno scritti sulla prima unità, nel secondo blocco

Paragonato ad una unità disco singola, i vantaggi di RAID 0 includono:

- Misura totale più grande — si possono creare RAID 0 array più grandi di una unità disco singola, facilitando la conservazione di file data più grandi
- Migliori prestazioni di lettura/scrittura — Il carico I/O su di un RAID 0 array viene diffuso in modo più omogeneo su tutte le unità presenti nell'array (Assumendo che tutti gli I/O non sono concentrati su di un singolo blocco)
- Nessuna perdita di spazio — Lo storage disponibile su tutte le unità presenti nell'array, può essere usato per la conservazione dei dati

Paragonato ad una unità singola, RAID 0 presenta i seguenti svantaggi:

- Minore affidabilità — Per rendere disponibile un array, ogni unità in un RAID 0 array deve essere operativa; un singolo problema all'unità in un N -drive RAID 0 array, ne può conseguire la rimozione di $1/N$ di tutti i dati, rendendo l'array inutilizzabile

**Suggerimento**

Se avete dei problemi a mantenere i diversi livelli RAID, ricordate soltanto che RAID 0 ha una percentuale di ridondanza pari a *zero* per cento.

5.6.2.1.1.2. RAID 1

RAID 1 utilizza due (anche se alcune implementazioni ne supportano di più) unità disco identiche. Tutti i dati vengono scritti su entrambe le unità, rendendole simili tra loro. Ecco perchè RAID 1 è spesso conosciuto come *mirroring*.

Ogni volta che si scrivono i dati su di un RAID 1 array, si deve eseguire la scrittura fisica su due unità: una sulla prima unità, e una sulla seconda unità. La lettura dei dati può essere eseguita solo una volta, e su di una delle due unità presenti nell'array.

Paragonato ad una unità disco singola, un RAID 1 array presenta i seguenti vantaggi:

- Una ridondanza migliorata — Anche se una unità presente nell'array ha dei problemi, i dati resteranno sempre accessibili
- Migliori prestazioni di lettura — Con entrambe le unità in funzione, i processi di lettura possono essere ripartiti in modo equo, riducendo così i carichi I/O per ogni unità.

Paragonato ad una unità disco singola, un RAID 1 array presenta i seguenti svantaggi:

- La misura massima dell'array viene limitata all'unità singola più grande disponibile.
- Riduzione delle prestazioni di lettura — Poichè entrambe le unità devono essere sempre aggiornate, tutti i processi I/O di scrittura devono essere eseguiti da entrambe le unità, rallentando così il processo generale di scrittura dei dati sull'array
- Aumento dei costi — Con una unità interamente dedicata alla ridondanza, il costo di un RAID 1 array risulta essere il doppio rispetto a quello di una singola unità

**Suggerimento**

Se avete delle difficoltà nel gestire i diversi livelli RAID, ricordate che RAID 1 presenta una ridondanza pari al *cento* per cento.

5.6.2.1.1.3. RAID 5

RAID 5 tenta di combinare i benefici di RAID 0 e RAID 1, minimizzando i rispettivi svantaggi.

Come RAID 0, un RAID 5 array consiste di unità disco multiple, ognuna suddivisa in blocchi. Ciò permette ad un RAID 5 array di essere più grande di ogni singola unità. Come RAID 1 array, un RAID 5 array utilizza una parte di spazio del disco in modo ridondante, aumentando così l'affidabilità.

Tuttavia il funzionamento di RAID 5 è diverso da RAID 0 o 1.

Un RAID 5 array è composto da almeno tre unità disco di misura identica (è possibile utilizzare unità aggiuntive). Ogni unità è suddivisa in blocchi ed i dati vengono scritti nei blocchi. Tuttavia, non ogni blocco è dedicato alla conservazione dei dati come lo è per RAID 0. Al contrario, in un array con n unità disco, ogni n blocco è dedicato alla *parità*.

I blocchi che contengono la parità rendono possibile il ripristino dei dati nel caso in cui una delle unità presenti un problema. La parità nel blocco x viene calcolata attraverso una combinazione matematica dei dati di ogni blocco x su tutte le altre unità presenti nell'array. Se i dati presenti in un blocco vengono aggiornati, il blocco di parità corrispondente deve essere ricalcolato e aggiornato.

Questo significa anche che ogni qualvolta che i dati vengono scritti sull'array, almeno *due* unità vengono scritte su: una unità che contiene i dati, e sulla unità che contiene il blocco di parità.

Un punto molto importante da ricordare è la distribuzione dei blocchi di parità, infatti essi non sono concentrati su di una unità presente nell'array. Al contrario essi sono distribuiti in modo omogeneo attraverso tutte le unità. Anche se è possibile dedicare una unità specifica per contenere nient'altro che la parità (infatti questa configurazione è conosciuta come RAID level 4), il costante aggiornamento della parità stessa man mano che i dati vengono scritti sull'array, significherebbe che l'unità che presenta la parità potrebbe diventare una limitazione nei confronti della prestazione. Distribuendo le informazioni sulla parità in modo omogeneo sull'array, è possibile ridurre questo impatto.

Tuttavia, è importante tener presente l'impatto della parità sulla capacità generale dello storage dell'array. Anche se le informazioni sulla parità sono distribuite omogeneamente attraverso tutte le unità presenti nell'array, la quantità di storage disponibile viene ridotta nella misura di una unità.

Paragonato ad una unità disco singola, un RAID 5 array presenta i seguenti vantaggi:

- Miglioramento della ridondanza — Se si verifica un errore di una unità presente nell'array, le informazioni riguardanti la parità possono essere utilizzate per ricostruire i blocchi mancanti di dati, il tutto, mantenendo l'array disponibile all'uso⁹
- Prestazioni migliori nella lettura — A causa del modo in cui vengono suddivisi i dati tra le unità presenti nell'array in un RAID simile a 0, il processo I/O di lettura viene suddiviso in modo omogeneo tra tutte le unità
- Costi ragionevolmente buoni — Per un RAID 5 array con n unità, solo $1/n$ dell'intero storage disponibile, viene dedicato alla ridondanza

Paragonato ad una unità disco singola, un RAID 5 array presenta i seguenti svantaggi:

- Diminuzione della prestazione riguardante il processo di scrittura — Poichè un processo di scrittura su di un array risulta in almeno due processi di scrittura sulle unità fisiche (uno per i dati ed uno per la parità), le prestazioni riguardanti il processo di scrittura sono peggiori paragonati ad una unità singola¹⁰

5.6.2.1.1.4. Livelli RAID nidificati

Anche se sembra ovvio dopo aver affrontato i diversi livelli RAID, ogni livello risulta avere dei punti di forza e dei punti deboli. Subito dopo aver iniziato ad impiegare i diversi livelli RAID e combinandoli tra loro, gli utenti sono stati in grado di ottenere degli array che presentavano solo le caratteristiche migliori di ogni livello, e nessuna debolezza.

Per esempio, cosa potrebbe accadere se le unità disco in un RAID array erano invece dei RAID 1 array? Ciò presenterà i vantaggi della velocità RAID 0, insieme all'affidabilità di RAID 1.

Questo è proprio quello che si potrebbe fare. Ecco i livelli RAID nidificati più usati:

9. La prestazione I/O viene ridotta operando con una unità non disponibile, a causa dell'overhead presente nella ricreazione dei dati mancanti.

10. Vi è altresì un altro impatto derivato dai calcoli sulla parità necessari per ogni processo di scrittura. Tuttavia a seconda delle implementazioni RAID 5 specifiche (in modo particolare, dove vengono eseguiti i calcoli sulla parità), questo impatto può variare da una misura considerevole ad una quasi non esistente.

- RAID 1+0
- RAID 5+0
- RAID 5+1

Poichè i RAID nidificati vengono usati in ambienti specializzati, non entreremo nei particolari. Tuttavia se si pensa ai RAID nidificati, è necessario ricordare due punti importanti:

- L'ordine è importante — L'ordine con il quale i livelli RAID si 'nidificano' potrebbe avere un certo impatto sull'affidabilità. In altre parole, RAID 1+0 e RAID 0+1 *non* sono uguali.
- I costi potrebbero essere elevati — Se vi è uno svantaggio comune a tutte le implementazioni RAID nidificati, quello è il costo; per esempio, il RAID 5+1 array più piccolo consiste in sei unità disco (sono necessari un maggior numero di unità per array più grandi).

Ora che abbiamo esplorato i concetti riguardanti il RAID, vediamo ora come quest'ultimo possa essere implementato.

5.6.2.1.2. Implementazioni RAID

Risulta ovvio dalle precedenti sezioni, che RAID necessita di una "intelligenza" aggiuntiva, che vada oltre alla normale processazione I/O del disco, per le unità individuali. È necessario eseguire i seguenti compiti:

- Suddividere le richieste I/O in entrata per i dischi individuali nell'array
- Per RAID 5, calcolo della parità e la scrittura di quest'ultima sull'unità appropriata presente nell'array
- Controllo dei dischi individuali presenti nell'array, ed eseguire le azioni appropriate se uno dei dischi presenta dei problemi
- Controllo della ricostruzione di un disco individuale presente nell'array, quando il suddetto disco è stato sostituito o riparato
- Fornire un mezzo con il quale gli amministratori sono in grado di gestire l'array (rimozione e aggiunta delle unità, avviare e arrestare la ricreazione, ecc)

Ci sono due metodi principali in grado di eseguire questi compiti. Le seguenti sezioni descrivono i suddetti metodi in modo più dettagliato.

5.6.2.1.2.1. Hardware RAID

Una implementazione hardware RAID generalmente ha la forma di una scheda specializzata per il controller del disco. La scheda esegue tutte le funzioni relative a RAID, e controlla direttamente le unità individuali negli array collegati. Con l'unità corretta, gli array gestiti da una scheda hardware RAID, appaiono sul sistema operativo host come unità disco regolari.

Molte schede del RAID controller funzionano con le unità SCSI, anche se vi sono comunque alcuni RAID controller 'ATA-based'. In ogni caso, l'interfaccia amministrativa viene generalmente implementata in uno dei seguenti modi:

- Programmi utility specializzati che vengono eseguiti come applicazioni con il sistema operativo host, in modo da presentare una interfaccia software per la scheda del controller
- Una interfaccia 'on-board' che utilizza una porta seriale in grado di essere accessa usando un emulatore terminale
- Una interfaccia simile a BIOS che può essere accessa durante la prova di alimentazione del sistema

Alcuni RAID controller possiedono più di una interfaccia amministrativa. Per ovvie ragioni, una interfaccia software fornisce maggiore flessibilità, in quanto permette di eseguire funzioni amministrative mentre il sistema operativo è in funzione. Tuttavia se state eseguendo un avvio del sistema operativo stesso tramite un RAID controller, è necessario avere una interfaccia che non richiede un sistema operativo in esecuzione.

Poichè sono presenti sul mercato diverse schede per il RAID controller, risulta essere impossibile quindi scendere nei particolari. È fortemente consigliato leggere la documentazione fornita dal rivenditore in modo da poter ottenere maggiori informazioni.

5.6.2.1.2.2. Software RAID

Software RAID non è altro che una implementazione di RAID come se fosse un kernel- o un software driver-level per un sistema operativo particolare. Per questo, esso è in grado di fornire maggiore flessibilità in termini di supporto — fino a quando il sistema operativo è in grado di supportare l'hardware, i RAID array possono essere configurati ed impiegati. Questo riduce in modo sostanziale il costo per l'impiego di RAID, eliminando la necessità di utilizzare un hardware RAID specializzato e molto costoso.

Spesso l'alimentazione della CPU disponibile per i calcoli della parità per il software RAID, eccede l'alimentazione di processazione presente su di una scheda per il RAID controller. Quindi, alcune implementazioni per il software RAID hanno la capacità per una prestazione più elevata rispetto alle implementazioni hardware RAID.

Tuttavia, software RAID presenta alcune limitazioni che non sono presenti in hardware RAID. La più importante da considerare è il supporto per l'avvio da un software RAID array. In molti casi, solo RAID 1 array possono essere utilizzati per l'avvio, in quanto il BIOS del computer non riconosce RAID. Poichè l'unità singola di un RAID 1 array è indistinguibile da un dispositivo di avvio non-RAID, il BIOS è in grado di eseguire il processo di avvio in modo corretto; il sistema operativo a questo punto può smistarsi sulle operazioni del software RAID una volta avuto il controllo del sistema.

5.6.3. Logical Volume Management

Un'altra tecnologia molto avanzata riguardante lo storage, è rappresentata dal *logical volume management* (LVM). LVM rende possibile trattare i dispositivi fisici di mass storage come blocchi di basso livello sui quali sono create le diverse configurazioni per lo storage. Le capacità esatte variano a seconda dell'implementazione specifica, ma possono includere il grouping fisico dello storage, il logical volume resizing, e la migrazione dei dati.

5.6.3.1. Grouping fisico dello storage

Anche se il nome di questa funzione può differire, il grouping fisico dello storage rappresenta la base per tutte le implementazioni LVM. Come indicato dal nome, i dispositivi fisici mass storage possono essere raggruppati in modo tale da creare uno o più dispositivi logici mass storage. I suddetti dispositivi (o logical volume) possono avere una capacità maggiore di quella di qualsiasi dispositivo mass storage fisico indicato.

Per esempio, date due unità da 100GB, è possibile creare un logical volume di 200GB. Tuttavia, è possibile creare anche un logical volume di 50GB o di 150GB. Qualsiasi combinazione del logical volume che risulta essere uguale o minore alla capacità totale (in questo esempio 200GB) risulta essere possibile. Le scelte sono limitate solo dalle esigenze della vostra organizzazione.

Ciò rende possibile per un amministratore di sistema trattare l'intero storage come se fosse parte di un singolo gruppo, disponibile all'uso in qualsiasi quantità. Altresì, è possibile aggiungere al gruppo, in un secondo momento, delle unità, rendendo il processo molto semplice ed in grado di soddisfare la richiesta di storage degli utenti.

5.6.3.2. Logical Volume Resizing

Una caratteristica dell'LVM molto apprezzata da parte degli amministratori di sistema è la sua abilità a direzionare facilmente lo storage là dove se ne presenta la necessità. In una configurazione di un sistema non-LVM, esaurire lo spazio può significare — nelle ipotesi migliori — muovere i file dal dispositivo saturo ad uno con maggiore spazio. Spesso significa anche riconfigurare i dispositivi mass storage del vostro sistema; un compito che potrebbe essere eseguito dopo il normale orario d'ufficio.

Tuttavia, LVM rende possibile un facile aumento della misura di un logical volume. Considerate per un momento che il nostro gruppo di storage da 200GB sia stato usato per creare un logical volume da 150GB, con il rimanente 50GB tenuto come riserva. Se il logical volume da 150GB si riempie, LVM è in grado di aumentare la sua misura (diciamo di 10GB), senza alcuna configurazione fisica. A seconda dell'ambiente del sistema operativo, potrebbe essere possibile eseguire tale operazione sia dinamicamente oppure attraverso una brevissima interruzione 'downtime' in modo da eseguire il resizing.

5.6.3.3. Migrazione dei dati

Molti amministratori di sistema con una certa esperienza potranno essere impressionati dalla capacità dell'LVM, ma forse potrebbero domandarsi quanto segue:

Cosa succede se una delle unità che compongono il logical volume non riesce ad eseguire l'avvio?

La buona notizia è che la maggior parte delle implementazioni LVM offrono la possibilità di *migrare* i dati su di una unità particolare. Per funzionare, ci deve essere una sufficiente capacità di riserva, in modo da assorbire la perdita dell'unità. Una volta completata la migrazione, l'unità in questione può essere sostituita e aggiunta nel gruppo dello storage disponibile.

5.6.3.4. Perché usare RAID con l'LVM?

Tenuto conto che LVM presenta alcuni contenuti simili a RAID (e cioè la possibilità di sostituire dinamicamente le unità che presentano dei problemi), ed altre caratteristiche che non possono essere eguagliate da molte implementazioni RAID (come ad esempio la possibilità di aggiungere dinamicamente più storage ad un gruppo storage centrale), molte persone si possono domandare se RAID non ha più molta importanza.

Niente di più sbagliato. RAID e LVM sono delle tecnologie complementari che possono essere usate insieme (in modo simile ai livelli RAID nidificati), beneficiando così delle caratteristiche migliori di entrambi.

5.7. Gestione giornaliera dello storage

Gli amministratori di sistema devono prestare molta attenzione nei riguardi dello storage durante i propri compiti giornalieri. Ci sono diverse problematiche da tener presente:

- Controllo dello spazio disponibile
- Problematiche relative al disk quota

- Problematiche relative al file
- Problematiche relative alla directory
- Problematiche relative al backup
- Problematiche relative alla prestazione
- Aggiunta/rimozione dello storage

Le seguenti sezioni affrontano i suddetti punti in modo più dettagliato.

5.7.1. Controllo dello spazio disponibile

Assicurarsi che ci sia sufficiente spazio disponibile, dovrebbe rappresentare uno dei compiti giornalieri più importanti di un amministratore di sistema. Il motivo per il quale il suddetto controllo è così importante, è dovuto al fatto che il suddetto spazio è così dinamico, che potrebbe esserci spazio sufficiente in un momento, mentre potrebbe essere quasi nullo un secondo più tardi.

In generale, ci sono tre ragioni che determinano uno spazio insufficiente:

- Uso eccessivo da parte di un utente
- Uso eccessivo da parte di una applicazione
- Aumento normale nel suo utilizzo

I suddetti motivi verranno affrontati in dettaglio nelle seguenti sezioni.

5.7.1.1. Uso eccessivo da parte di un utente

Persone diverse possiedono differenti abilità. Alcune persone potrebbero inorridirsi nel vedere una patina di polvere ricoprire un tavolo, mentre altri potrebbero accatastare senza alcun problema, numerosi contenitori usati per la pizza in un angolo della cucina. Lo stesso discorso si potrebbe fare con lo storage:

- Alcune persone potrebbero essere molto meticolose nell'utilizzo dello storage, e non lasciare mai in giro i file non necessari.
- Altre invece non hanno mai tempo per riorganizzare i propri file, e di sbarazzarsi quindi di quelli non più necessari.

In molte circostanze in cui un utente è responsabile nell'uso di grosse quantità di storage, il responsabile viene sempre identificato come il secondo tipo di persona.

5.7.1.1.1. Come far fronte ad uno sfruttamento eccessivo da parte di un utente

Questa rappresenta un'area dove l'amministratore di sistema ha bisogno di unire tutta la diplomazia e le capacità interpersonali che possiede. Molto spesso le discussioni riguardanti lo spazio del disco possono diventare molto spinose, in quanto molte persone vedono nelle restrizioni sull'utilizzo del disco, un aumento delle difficoltà nell'adempimento del proprio lavoro, oppure che le restrizioni sono troppe, o semplicemente che essi non hanno tempo sufficiente per eliminare i file a loro non più utili.

I migliori amministratori di sistema considerano in queste situazioni molteplici fattori. Sono le suddette restrizioni ragionevoli e giuste per il tipo di lavoro svolto dalla persona in questione? L'utilizzo di spazio fatto dalla suddetta persona, è un utilizzo corretto? È possibile aiutare in alcun modo questa persona a ridurre l'utilizzo di spazio (creando un CD-ROM di backup di tutte le email più vecchie)? Il vostro compito durante la conversazione, è quello di scoprire quanto sopra elencato, assicurandovi che la persona che non necessita di tutto questo spazio, elimini i file non più utili.

In alcun modo, mantenete la conversazione a livelli professionali. Cercate di risolvere tali problematiche in modo chiaro e cortese ("Capisco che lei sia molto occupato, ma qualsiasi altro dipendente

nel suo dipartimento ha avuto le stesse istruzioni, ed il loro uso risulta essere la metà del suo."), e cercando di muovere la conversazione verso il punto desiderato. Assicuratevi di offrire assistenza se il problema risulta essere una carenza di esperienza/conoscenza.

Affrontando una conversazione in modo sensibile ma fermo, spesso risulta essere più efficace che usare la propria autorità di amministratore di sistema per ottenere quello che volete. Per esempio, potreste trovarvi nelle circostanze dove è necessario scendere a compromessi. Questi compromessi possono avere la seguente forma:

- Provvedere ad uno spazio temporaneo
- Creare dei backup di archivio
- Arrendersi

È possibile che gli utenti siano in grado di ridurre l'uso di spazio se essi hanno a disposizione una quantità di spazio temporaneo da poter usare senza restrizioni. Le persone che utilizzano questa soluzione, potranno lavorare senza alcun problema fino al raggiungimento di un limite logico, a questo punto essi potranno eseguire l'eliminazione di file non più necessari, e determinare quali sono i file utili presenti nello storage temporaneo.



Avvertenza

Se desiderate offrire questa alternativa ad un utente, *non* cadete nel tranello di trasformare il suddetto spazio temporaneo in permanente. Siate chiari quando offrite questo tipo di soluzione, chiarendo anche il fatto che non garantite alcuna memorizzazione dei dati; non è disponibile alcun backup di qualsiasi dato presente nello spazio temporaneo.

Infatti, molti amministratori spesso danno risalto a questo fattore, cancellando automaticamente qualsiasi file che presenti una data specifica, all'interno dello storage temporaneo (per esempio, una settimana).

Altre volte gli utenti possono avere file così vecchi, che risulta molto improbabile che essi abbiano necessità di un continuo accesso. Assicuratevi che quanto detto sia effettivamente la situazione sopra indicata. Talvolta alcuni utenti sono responsabili della gestione di un archivio di file molto vecchi; in questi casi, è vostro compito assistere questo tipo di utenti, fornendo backup multipli da trattare in modo simile ai backup di archivio del vostro centro dati.

Tuttavia, ci sono circostanze in cui i dati sono di dubbio valore. In questi casi potreste offrire l'utilizzo di un backup speciale. Eseguite allora il backup dei dati vecchi, e offrite all'utente il media di backup, spiegando loro che essi sono i responsabili del suo mantenimento, e che al momento di accedere ai dati, essi devono rivolgersi a voi (oppure al personale indicato della vostra organizzazione — se appropriato) per il loro ripristino.

Ci sono alcune cose da ricordare in modo tale da non far ricadere ogni responsabilità su di voi. Come prima cosa non includere i file che con tutta probabilità hanno bisogno di essere ripristinati, non selezionate file *troppo* recenti. Successivamente, assicuratevi di essere in grado di eseguire un ripristino di dati nel caso ce ne sia bisogno. Questo significa che il media di backup dovrebbe essere dello stesso tipo di quello che verrà utilizzato nel vostro centro dati.



Suggerimento

La scelta del media di backup dovrebbe tenere in considerazione quelle tecnologie che permettono all'utente stesso di eseguire un ripristino dei dati. Per esempio, anche se il backup di diversi gigabyte su di un media CD-R, risulta essere più impegnativo di una emissione di un comando singolo e trasmetterlo su di un nastro da 20GB, considerate il fatto che l'utente deve essere in grado di accedere ai dati sul CD-R ogni qualvolta essi lo desiderino — senza chiedere aiuto a voi.

5.7.1.2. Uso eccessivo da parte di una applicazione

Talvolta una applicazione può essere responsabile per un uso eccessivo di spazio. Le ragioni possono variare e possono includere:

- I miglioramenti della funzionalità di una applicazione necessitano maggiore spazio
- Un aumento del numero degli utenti che utilizzano la suddetta applicazione
- L'applicazione non riesce ad eliminare i file non più necessari, lasciando quest'ultimi nel disco
- L'applicazione non è più funzionante in modo corretto, ed il bug presente fa sì che essa utilizzi più spazio

Il vostro compito è quello di determinare, facendo riferimento all'elenco riportato, quale dei punti sopra riportati possa essere applicato nella vostra situazione. Facendo attenzione allo stato delle applicazioni utilizzate nel vostro centro dati, potrebbe facilitare il processo di eliminazione di alcuni dei punti sopra riportati. Ciò che rimane da fare è un lavoro di ricerca dello storage. Questa sequenza vi aiuterà ad assottigliare ulteriormente il campo.

A questo punto dovete seguire le fasi appropriate, e cioè l'aggiunta di storage in modo da supportare l'applicazione in questione, contattare gli sviluppatori dell'applicazione, in modo da discutere le caratteristiche di gestione del file, oppure scrivere alcuni script in grado di eliminare i file che non sono più desiderati.

5.7.1.3. Aumento normale nel suo utilizzo

Molte organizzazioni con il passare del tempo tendono ad aumentare il loro volume di lavoro. Per questo motivo è facile prevedere un aumento dell'utilizzo dello storage. In quasi tutte le occasioni, un controllo continuo è in grado di rivelare la percentuale di utilizzo dello storage della vostra organizzazione; questa percentuale può essere utilizzata per determinare il tempo necessario per ottenere dello storage aggiuntivo, prima che il vostro spazio disponibile venga consumato.

Se vi trovate nella situazione in cui non avete sufficiente spazio, ciò vorrà dire che non avete eseguito correttamente il vostro lavoro.

Tuttavia, è possibile che si verifichi una improvvisa ed elevata richiesta di storage. Per esempio se la vostra organizzazione si sia unita con un'altra entità, necessitando così di modifiche rapide nell'infrastruttura IT (tradotto quindi in storage). Oppure tale richiesta possa essere determinata dalla presenza di un progetto molto importante. Le modifiche ad una applicazione esistente potrebbero richiedere un aumento dello spazio del vostro storage.

Non ha importanza quale motivo possa causare un aumento della richiesta di spazio per il vostro storage, ci saranno sempre dei momenti in cui sarete presi di sorpresa, cercate quindi di configurare l'architettura del vostro storage in modo da ottenere sempre la massima flessibilità. Avere sempre dello spazio di riserva (se possibile), potrebbe alleviare l'impatto dovuto alla presenza di eventi non pianificati.

5.7.2. Problematiche riguardanti il disk quota

Spesso la prima cosa alla quale ci si riferisce pensando al disk quotas è quella di utilizzarlo per forzare gli utenti a mantenere le proprie directory in modo ordinato. Mentre quanto appena detto potrebbe

essere vero per alcuni siti, esso potrebbe essere utile per affrontare, da una prospettiva diversa, il problema riguardante lo spazio presente sul disco. Cosa fare per affrontare il problema delle applicazioni che per qualche ragione consumano troppo spazio? Purtroppo non è raro il verificarsi di tale problema, in questi casi i disk quota sono in grado di limitare il danno causato da tali applicazioni, arrestandole *prima* di esaurire lo spazio disponibile.

La parte più difficile per quanto riguarda l'implementazione e la gestione dei disk quota ruota intorno ai loro limiti. Quali dovrebbero essere? Un approccio alquanto semplicistico sarebbe quello di dividere lo spazio del disco per il numero degli utenti e/o gruppi che ne fanno uso, usando il numero risultante come il per-user quota. Per esempio, se il sistema possiede una unità disco di 100GB e 20 utenti, ad ogni utente deve essere conferito un disk quota che non sia maggiore di 5GB. In questo modo, verranno garantiti ad ogni utente 5GB (anche se il disco sarà 100%).

Per i sistemi operativi che lo supportano, i quota temporanei potrebbero essere impostati con un valore ancora più elevato — diciamo 7.5GB, con il quota permanente fisso a 5GB. Questo approccio presenta il beneficio di permettere agli utenti un consumo fisso pari alla percentuale del disco, pur avendo una certa flessibilità quando un utente raggiunge (ed eccede) il proprio limite. Quando si usano i disk quota in questo modo, significa impegnare oltremodo lo spazio disponibile. Il quota temporaneo è 7.5GB. Se tutti e 20 gli utenti eccedono il proprio quota permanente nello stesso momento, e cercano di usare il proprio quota temporaneo, il valore precedentemente indicato di 100GB del disco, dovrà essere invece di 150GB.

In pratica non tutti eccedono il proprio limite nello stesso istante, rendendo tale approccio molto utile. Naturalmente, la selezione dei quota permanenti e temporanei dipende dall'amministratore del sistema, in quanto ogni sito e ogni community sono diversi tra loro.

5.7.3. Problematiche relative al file

Gli amministratori di sistema spesso devono affrontare delle problematiche relative al file. Queste problematiche possono includere:

- L'accesso del file
- File Sharing

5.7.3.1. L'accesso del file

Le problematiche relative all'accesso del file si identificano generalmente in uno scenario — e cioè un utente non è in grado di accedere ad un file credendo invece di aver diritto all'accesso.

Spesso succede che l'utente #1 desidera dare una copia di un file all'utente #2. In molte organizzazioni, l'abilità di un utente di accedere ai file di un altro utente è un processo molto breve, che scaturisce il suddetto problema.

Per questo motivo sono disponibili tre diversi approcci:

- L'utente #1 esegue le modifiche necessarie per permettere all'utente #2 di accedere il file se lo stesso è esistente.
- Per tale scopo viene creata un'area di scambio; l'utente #1 posiziona lì una copia del file, il quale può essere copiata dall'utente #2.
- L'utente #1 utilizza l'email per dare all'utente #2 una copia del file.

Vi è comunque un problema con il primo approccio — a seconda di come si garantisce l'accesso; l'utente #2 potrebbe avere un accesso completo a tutti i file dell'utente #1. Ancora peggio, l'accesso potrebbe essere stato strutturato in modo da permettere a *tutti* gli utenti presenti nella vostra organizzazione, di accedere ai file dell'utente #1. Una situazione peggiore potrebbe essere quella in cui l'utente #2 non ha più bisogno di tale accesso, lasciando i file dell'utente #1 accessibili a tutti gli altri

utenti. Sfortunatamente, quando gli utenti sono responsabili di questo tipo di situazione, generalmente la sicurezza non ha per loro molta importanza.

Il secondo approccio elimina il problema precedentemente indicato, e cioè i file dell'utente #1 non sono più accessibili da parte di altri utenti. Tuttavia, quando i file si trovano nell'area di scambio, essi possono essere letti (e a seconda dei permessi, possono essere anche modificati) da tutti gli utenti. Questo tipo di approccio risalta la possibilità che l'area di scambio dei file venga saturata, in quanto gli utenti spesso si dimenticano di ripulirla dai propri file.

Il terzo approccio, anche se potrebbe sembrare un pò imbarazzante, potrebbe risultare quello più idoneo in molte situazioni. Con l'avvento dei protocolli standard delle email e con programmi molto più intelligenti, l'invio di file tramite email è divenuto una operazione molto semplice, che non necessita alcun coinvolgimento dell'amministratore. Naturalmente, si potrebbe verificare l'eventualità che un utente invii un file database da 1GB al personale del dipartimento delle finanze il quale è composto da 150 unità, per questo motivo è importantissimo educare i propri utenti (e possibilmente informarli dei limiti sugli allegati delle email).

5.7.3.2. File Sharing

Quando utenti multipli hanno la necessità di condividere la copia di un file, permettere l'accesso modificando i permessi del file, generalmente non rappresenta la cosa migliore da fare. È preferibile rendere pubblico lo stato del file condiviso. Ci sono diverse ragioni per fare quanto sopra descritto:

- I file condivisi che si trovano all'esterno della directory di un utente, possono scomparire quando, per esempio, un utente abbandona l'organizzazione oppure riordina i propri file.
- Mantenere l'accesso condiviso per più di uno o due utenti aggiuntivi diventa difficile, esso infatti può creare il problema di un lavoro aggiuntivo, richiesto quando gli utenti che condividono l'accesso modificano i loro compiti.

Quindi l'approccio preferito è il seguente:

- L'utente originale rinuncia all'ownership diretto del file
- Creare un gruppo che possiede il file
- Posizionare il file in una directory condivisa posseduta dal gruppo
- Rendere tutti gli utenti che hanno bisogno di accedere al file, parte del gruppo

Naturalmente, questo tipo di approccio funziona bene sia per file multipli e sia con file singoli, e può essere utilizzato per implementare lo storage condiviso per progetti più complessi.

5.7.4. Aggiunta/rimozione dello storage

Poichè la richiesta di spazio aggiuntivo è sempre una costante, l'amministratore di sistema spesso ha bisogno di reperire maggiore spazio, questo può essere eseguito rimuovendo unità più piccole e obsolete. Questa sezione fornisce una panoramica del processo di base per l'aggiunta e la rimozione dello storage.



Nota bene

Su molti sistemi operativi, i dispositivi mass storage sono identificati a seconda del loro collegamento fisico al sistema. Quindi, aggiungendo e rimuovendo i dispositivi mass storage, ne consegue una modifica dei nomi del dispositivo. Quando aggiungete o rimuovete il vostro storage, assicuratevi

sempre di rivedere (e aggiornare, se necessario) tutti i riferimenti al nome del dispositivo utilizzato dal vostro sistema operativo.

5.7.4.1. Aggiunta dello storage

Il processo per l'aggiunta dello storage ad un sistema, è relativamente semplice. Di seguito vengono riportate le fasi di base:

1. Installazione dell'hardware
2. Partizionamento
3. Formattazione delle partizioni
4. Aggiornamento della configurazione del sistema
5. Modifica del programma di backup

Le seguenti sezioni affrontano le suddette fasi in modo più dettagliato.

5.7.4.1.1. Installazione dell'hardware

Prima di tutto l'unità disco deve essere presente e accessibile. Anche se sono disponibili diverse configurazioni hardware, le seguenti sezioni affrontano le situazioni più comuni — aggiungendo una ATA o una unità disco SCSI. Anche con altre configurazioni, sono sempre valide le fasi di base qui riportate.



Suggerimento

Non ha importanza quale hardware utilizzate, dovete sempre tener presente il carico che la nuova unità disco aggiunge al sottosistema I/O del vostro computer. In generale, dovrete provare a distribuire il carico I/O del disco su tutti i canali/bus disponibili. Dal punto di vista della prestazione, questa operazione risulta essere migliore rispetto a quella di mettere tutte le unità disco su di un canale, lasciando un altro canale vuoto in posizione idle.

5.7.4.1.1.1. Aggiunta delle unità disco ATA

Le unità disco ATA sono molto usate nei desktop e nei sistemi server di tipo lower-end. Quasi tutti i sistemi in queste classi, possiedono dei controller ATA interni con canali ATA multipli — normalmente due o quattro.

Ogni canale è in grado di supportare due dispositivi — uno master ed uno slave. I due dispositivi sono collegati al canale per mezzo di un cavo singolo. Per questo motivo, la prima fase è quella di vedere quali canali possiedono uno spazio disponibile in grado di ospitare una unità disco aggiuntiva. È possibile il verificarsi di una delle seguenti situazioni:

- È presente un canale con una sola unità disco collegata ad esso
- È presente un canale con nessuna unità disco collegata ad esso
- Non vi è alcuno spazio disponibile

La prima situazione è quella generalmente più semplice, in quanto è possibile che il cavo già presente possieda un connettore non utilizzato nel quale la nuova unità disco può essere collegata. Tuttavia, se il cavo possiede solo due connettori (uno per il canale ed uno per l'unità disco precedentemente installata), allora risulta essere necessario sostituire il cavo esistente con un modello a tre connettori.

Prima di installare la nuova unità disco, assicuratevi che le due unità disco che condividono il canale, siano configurate in modo corretto (una come master, e l'altra come slave).

La seconda situazione risulta essere un pò più complessa, solo perchè è necessario procurare un cavo in modo da poter collegare una unità disco al canale. La nuova unità disco può essere configurata come master o come slave (anche se tradizionalmente la prima unità su di un canale, è configurata come master).

Nella terza situazione, non è disponibile alcuno spazio per poter implementare una unità aggiuntiva. Dovete quindi prendere una decisione. Preferite:

- Acquisire una scheda del controller ATA, e installarla
- Sostituire una delle unità disco installate con una più nuova e più larga

Aggiungere una scheda comporta il controllo della compatibilità dell'hardware e del software. Sostanzialmente la scheda deve essere compatibile con l'alloggiamento del bus del vostro computer, e deve essere disponibile un alloggiamento per la suddetta scheda, la quale deve essere supportata dal sistema operativo. La sostituzione di una unità disco presenta un unico problema: cosa fare con i dati presenti sul disco? In questo caso ci sono alcuni approcci possibili:

- Scrivere i dati su di un dispositivo di backup e ripristinarlo dopo l'installazione di una nuova unità disco
- Utilizzate la vostra rete per copiare i dati su di un altro sistema che presenta uno spazio sufficiente, ripristinando i dati installati sulla nuova unità disco
- Utilizzare lo spazio fisico occupato da una terza unità disco per:
 1. Rimuovere momentaneamente la terza unità disco
 2. Installare temporaneamente la nuova unità disco al suo posto
 3. Copiare i dati su di una nuova unità disco
 4. Rimuovere la vecchia unità disco
 5. Sostituirla con la nuova unità
 6. Installare nuovamente la terza unità disco momentaneamente rimossa
- Installare temporaneamente l'unità disco originale e la nuova unità su di un altro computer, copiare i dati su di una nuova unità, e poi installarla sul computer originale

Come potete vedere, talvolta bisogna impegnarsi leggermente per posizionare i dati (ed il nuovo hardware) nel posto corretto.

5.7.4.1.1.2. Aggiunta delle unità disco SCSI

Le unità disco SCSI vengono usate normalmente nelle workstation di tipo higher-end e sistemi server. Diversamente dai sistemi 'ATA-based', i sistemi SCSI possono o meno, avere dei controller SCSI interni; altri invece utilizzano una scheda del controller SCSI separata.

Le capacità dei controller SCSI (sia interni che esterni), possono variare. Possono fornire un bus SCSI di tipo 'wide' oppure di tipo 'narrow'. La velocità del bus può essere normale, veloce, ultra, ultra2, o ultra160.

Se questi termini vi risultano poco familiari (essi sono stati affrontati brevemente nella Sezione 5.3.2.2), allora dovrete determinare le capacità della vostra configurazione hardware, e selezionare l'unità disco appropriata. La migliore fonte per queste informazioni è rappresentata dalla documentazione del vostro sistema e/o dell'adattatore SCSI.

Successivamente è necessario determinare il numero di bus SCSI disponibili sul vostro sistema, e quali possiedono sufficiente spazio per una nuova unità disco. Il numero di dispositivi supportati da un bus SCSI, varia a seconda della larghezza del bus:

- Bus SCSI narrow (8-bit) — 7 dispositivi (più il controller)
- Bus SCSI Wide (16-bit) — 15 dispositivi (più il controller)

La prima fase è quella di controllare quale bus possiede uno spazio sufficiente per una unità disco aggiuntiva. Ecco le situazioni che si possono verificare:

- È presente un bus con un numero inferiore rispetto alla capacità massima, di unità disco ad esso collegate
- È presente un bus con nessuna unità disco collegata
- Non vi è alcuno spazio disponibile

La prima situazione è quella più semplice, in quanto è possibile che il cavo abbia un connettore non utilizzato nel quale è possibile collegare la nuova unità disco. Tuttavia, se il cavo in uso non possiede il suddetto connettore, allora è necessario sostituire il cavo esistente con uno che possiede almeno un connettore in più.

La seconda situazione è un po' più complessa, solo perché è necessario procurare un cavo in modo tale da poter collegare una unità disco al bus.

Se invece non vi è alcuno spazio disponibile per una unità disco aggiuntiva, allora dovete prendere una decisione. Desiderate:

- Acquistare e installare una scheda del controller SCSI
- Sostituire una unità installata con una nuova unità, più grande

Aggiungere una scheda del controller comporta un controllo della compatibilità hardware e della capacità fisica e software. Sostanzialmente la scheda deve essere compatibile con l'alloggiamento del bus del vostro computer, deve essere disponibile un alloggiamento per la suddetta scheda, la quale deve essere supportata dal sistema operativo.

La sostituzione di una unità disco precedentemente installata presenta un unico problema: cosa fare con i dati presenti sul disco? Vi sono alcuni approcci in tal senso:

- Scrivere i dati su di un dispositivo di backup, e ripristinarli dopo aver installato la nuova unità disco
- Usare la vostra rete in modo da copiare i dati su di un altro sistema che abbia uno spazio sufficiente, e ripristinarli dopo aver installato la nuova unità disco
- Utilizzare lo spazio fisico occupato da una terza unità disco per:
 1. Rimuovere momentaneamente la terza unità disco
 2. Installare temporaneamente la nuova unità disco al suo posto
 3. Copiare i dati su di una nuova unità disco
 4. Rimuovere la vecchia unità disco
 5. Sostituirla con la nuova unità
 6. Installare nuovamente la terza unità disco momentaneamente rimossa
- Installare temporaneamente l'unità disco originale e la nuova unità su di un altro computer, copiare i dati su di una nuova unità, e poi installarla sul computer originale

Una volta disponibile un connettore per mezzo del quale è possibile collegare la nuova unità disco, assicuratevi che l'ID SCSI dell'unità sia impostato in modo corretto. Per fare questo dovete essere a conoscenza dell'ID SCSI utilizzato (incluso il controller) dai dispositivi presenti sul bus. Il modo più semplice per fare quanto detto, è quello di accedere al BIOS del controller SCSI. Ciò può essere fatto piagiando una sequenza di tasti durante l'accensione del sistema. Potete così visualizzare la configurazione del controller SCSI, insieme con i dispositivi collegati a tutti i suoi bus.

Successivamente, dovrete considerare di terminare il bus in modo corretto. Quando aggiungete una nuova unità disco, la regola è molto semplice — se la nuova unità disco è l'ultimo dispositivo (oppure il solo dispositivo) presente sul bus, la sua terminazione deve essere abilitata, in caso contrario è necessario disabilitarla.

A questo punto, passate alla fase successiva — partizionare la vostra nuova unità disco.

5.7.4.1.2. Partizionamento

Una volta installata la nuova unità disco, è necessario creare una o più partizioni in modo da ottenere uno spazio idoneo per il vostro sistema operativo. Anche se i tool variano in base al sistema operativo, le fasi di base sono le stesse:

1. Selezionare la nuova unità disco
2. Visualizzare la nuova tabella di partizionamento dell'unità disco, in modo da assicurarsi che l'unità disco da partizionare sia quella corretta
3. Cancellate ogni partizione non desiderata che possa essere già presente nella nuova unità disco
4. Create la nuova partizione, assicuratevi di specificarne la dimensione ed il tipo
5. Salvate i vostri cambiamenti e uscite dal programma di partizionamento



Avvertenza

Quando partizionate la vostra nuova unità disco, è *necessario* assicurarsi che l'unità da partizionare sia quella corretta. In caso contrario, potreste partizionare una unità già in uso, causando una perdita di dati.

Assicuratevi altresì che abbiate precedentemente deciso la misura di partizionamento migliore. Affrontate questa fase in modo appropriato, in quanto la modifica della misura in un secondo momento, potrebbe risultare un processo molto difficoltoso.

5.7.4.1.3. Formattare le partizioni

A questo punto, la nuova unità disco può presentare una o più partizioni. Tuttavia prima di utilizzare lo spazio contenuto all'interno delle partizioni, esse devono essere formattate. Eseguendo tale operazione, non farete altro che selezionare il file system da utilizzare all'interno di ogni partizione. Tale fase risulta essere molto importante nella vita di una unità; le scelte fatte potranno essere modificate, ma ricordate che per fare ciò è richiesta una grandissima mole di lavoro.

Il processo vero e proprio di formattazione viene fatto eseguendo un programma utility; le fasi coinvolte per fare ciò, variano a seconda del sistema operativo. Una volta completato il processo di formattazione, l'unità disco è pronta all'uso.

Prima di continuare è consigliato controllare il lavoro fatto, per fare ciò accedere alle partizioni e assicuratevi che tutto sia in ordine.

5.7.4.1.4. Aggiornamento della configurazione del sistema

Se il vostro sistema operativo necessita delle modifiche riguardanti la sua configurazione, in modo da utilizzare il nuovo storage che avete appena aggiunto, adesso è il momento opportuno.

A questo punto potete essere relativamente sicuri che il sistema operativo sia configurato in modo corretto, in modo tale da poter rendere accessibile il nuovo storage in modo automatico ogni qualvolta il sistema viene avviato (anche se siete in grado di riavviare il vostro sistema in modo rapido, esso non presenta alcun problema —).

La sezione successiva affronta una delle fasi del processo per l'aggiunta di un nuovo storage, che molto spesso viene dimenticata.

5.7.4.1.5. Modifica del programma di backup

Considerando che lo storage contenga dati utili, questo è il momento giusto per eseguire le modifiche necessarie per quanto riguarda le vostre procedure di backup, e per assicurarsi che venga eseguito il backup del nuovo storage. La natura esatta di quello che è necessario eseguire, dipende dal modo con il quale il vostro sistema esegue i backup. Tuttavia ecco alcuni punti da considerare durante l'esecuzione delle modifiche necessarie:

- Considerate quale debba essere la frequenza di backup ottimale
- Determinate lo stile di backup più appropriato (solo backup completi, completi con incrementi, completi con differenziali, ecc)
- Considerate l'impatto dello storage aggiuntivo sull'utilizzo del vostro media di backup, in modo particolare nel corso della sua saturazione
- Giudicate se il backup aggiuntivo possa causare l'aumento del tempo necessario per eseguire un processo di backup, andando così oltre il limite della vostra finestra di backup
- Assicuratevi che i suddetti cambiamenti vengano comunicati al personale pertinente (altri amministratori di sistema, personale operativo ecc.)

Una volta fatto tutto questo, il vostro nuovo storage è pronto per l'uso.

5.7.4.2. Rimozione dello storage

La rimozione dello spazio presente sul disco è un processo molto semplice, con la maggior parte delle fasi molto simili alla sequenza d'installazione:

1. Rimuovete tutti i dati da salvare fuori dall'unità disco
2. Modificate il programma di backup in modo tale che non vi sia più un backup dell'unità disco
3. Aggiornate la configurazione del sistema
4. Cancellate i contenuti dell'unità disco
5. Rinuovere l'unità disco

Come potete notare, paragonato al processo di installazione, ci sono delle fasi aggiuntive. Queste fasi vengono affrontate nelle seguenti sezioni.

5.7.4.2.1. Rimuovete tutti i dati dall'unità disco

Se sono presenti sull'unità disco alcuni dati che necessitano di essere salvati, la prima cosa da fare è quella di determinare il luogo per la loro conservazione. Questa decisione dipende principalmente dal loro uso. Per esempio, se i dati non devono più essere usati, essi dovrebbero essere archiviati nello stesso modo dei backup del vostro sistema. Ciò significa che ora risulta essere il momento più opportuno, per considerare i periodi di memorizzazione più appropriati per questo processo di backup finale.



Suggerimento

Ricordate che in aggiunta alle direttive di qualsiasi memorizzazione dei dati presenti nella vostra organizzazione, ci potrebbero essere anche dei requisiti legali per la memorizzazione dei dati stessi. Per questo motivo, assicuratevi di consultare la sezione responsabile durante l'utilizzo dei suddetti dati; tale sezione dovrebbe essere a conoscenza di tale periodo.

D'altro canto, se i dati sono ancora in uso, gli stessi dovrebbero essere conservati sul sistema più appropriato. Naturalmente, se questo è il caso, risulterebbe essere più semplice muovere i dati, reinstallando l'unità disco sul nuovo sistema. Se decidete di seguire questo approccio, è consigliato eseguire un backup completo dei dati — molte persone camminando attraverso un centro dati, hanno perso dati utilissimi, solo perchè hanno urtato violentemente alcune unità disco.

5.7.4.2.2. Cancellare i contenuti dell'unità disco

Non ha importanza se l'unità conservi dati più o meno importanti, è sempre consigliato cancellare i contenuti prima di riassegnare o rinunciare al loro controllo. Mentre la ragione più ovvia è di assicurarsi di non lasciare informazioni sensibili sull'unità, esso risulta essere il momento migliore per controllare lo stato dell'unità, eseguendo una prova di lettura-scrittura per la ricerca di blocchi non corretti attraverso l'intera unità.



Importante

Molte compagnie (e agenzie governative), presentano metodi specifici per la cancellazione dei dati dalle unità disco, da altri media e per la loro conservazione. In questo caso assicuratevi *sempre* di aver capito e seguito questi requisiti; in molti casi se non si osservano i suddetti requisiti, si può andare incontro a problemi legali. L'esempio sopra riportato non deve essere considerato il metodo ultimo per la cancellazione dei dati da una unità.

In aggiunta le organizzazioni che lavorano con dati classificati, per poter disinstallare l'unità disco, devono osservare delle procedure legali ben specifiche (come ad esempio la sua distruzione fisica). In queste circostanze l'ufficio responsabile per la sicurezza delle informazioni della vostra organizzazione, dovrebbe essere in grado di assistervi in questo processo.

5.8. Una parola sui Backup...

Uno dei fattori più importanti nella considerazione dello storage di un disco, è quello dei backup. Purtroppo non abbiamo potuto affrontare questo tema in questa sezione, in quanto una sezione più dettagliata (la Sezione 8.2) è stata dedicata agli stessi backup.

5.9. Informazioni specifiche su Red Hat Enterprise Linux

A seconda della vostra esperienza come amministratore di sistema, la gestione dello storage con Red Hat Enterprise Linux può risultare familiare oppure una cosa totalmente sconosciuta. Questa sezione affronta gli aspetti riguardanti la gestione dello storage con Red Hat Enterprise Linux.

5.9.1. Convenzioni del device naming

Come con tutti i sistemi operativi simili a Linux, Red Hat Enterprise Linux utilizza i device file per accedere all'hardware (incluso le unità disco). Tuttavia, le convenzioni del naming per i dispositivi storage collegati, possono variare tra le diverse implementazioni Linux o implementazioni simili. Ecco come vengono chiamati i suddetti device file sotto Red Hat Enterprise Linux.



Nota bene

I nomi del dispositivo con Red Hat Enterprise Linux vengono determinati al momento dell'avvio

Quindi, le modifiche eseguite sulla configurazione hardware del sistema, possono risultare in una modifica dei device name al momento del riavvio del sistema. Per questo motivo si possono verificare alcuni problemi se i device name, presenti nei file di configurazione del sistema, non sono aggiornati in modo corretto.

5.9.1.1. File del dispositivo

Con Red Hat Enterprise Linux, i file del dispositivo per le unità disco appaiono nella directory `/dev/`. Il formato per ogni file name dipende da diversi aspetti dell'hardware, e da come esso sia stato configurato. Ecco i punti più importanti:

- Tipo di dispositivo
- Unità
- Partizione

5.9.1.1.1. Tipo di dispositivo

Le prime due lettere di un file name di un dispositivo si riferiscono al tipo specifico di dispositivo. Per le unità disco, ci sono due tipi di dispositivi più comuni:

- `sd` — Il dispositivo è basato su SCSI
- `hd` — Il dispositivo è di tipo 'ATA-based'

Maggiori informazioni su ATA e SCSI sono disponibili nella Sezione 5.3.2.

5.9.1.1.2. Unità

A seguire le due lettere inerenti il tipo di dispositivo, vengono riportate le due lettere che identificano l'unità specifica. Il designatore dell'unità inizia con "a" per la prima unità, "b" per la seconda, e così via. Per questo motivo, il primo disco fisso presente sul vostro sistema, potrebbe apparire come `hda` o `sda`.



Suggerimento

L'abilità di SCSI di indirizzare un grosso numero di dispositivi, necessita l'aggiunta di un secondo carattere per poter supportare i sistemi con più di 26 dispositivi SCSI collegati. Per questo motivo, i primi 26 dischi fissi SCSI presenti su di un sistema, verranno chiamati `sda` fino a `sdz`, i successivi 26 verranno chiamati `sdaa` fino a `sdaz` e così via.

5.9.1.1.3. Partizione

La parte finale del file name del dispositivo, è un numero in grado di rappresentare una partizione specifica sul dispositivo stesso, iniziando con "1." Il suddetto numero può essere composto da uno o due cifre, in base al numero di partizioni scritte sul dispositivo specifico. Una volta conosciuto il formato per i file name del dispositivo, è più facile capire a cosa ci si riferisce. Ecco alcuni esempi:

- `/dev/hda1` — La prima partizione sulla prima unità ATA
- `/dev/sdb12` — La dodicesima partizione sulla seconda unità SCSI
- `/dev/sdad4` — La quarta partizione sulla tredicesima unità SCSI

5.9.1.1.4. Accesso all'intero dispositivo

Ci possono essere delle circostanze in cui è necessario accedere non solo una partizione, bensì l'intero dispositivo. Questa operazione viene normalmente eseguita quando il dispositivo non viene partizionato o non supporta partizioni standard (come ad esempio una unità CD-ROM). In questi casi, il numero della partizione viene omissso:

- `/dev/hdc` — Il terzo dispositivo ATA
- `/dev/sdb` — Il secondo dispositivo SCSI

Tuttavia, molte unità disco utilizzano le partizioni (per maggiori informazioni sul partizionamento con Red Hat Enterprise Linux, consultate la Sezione 5.9.6.1).

5.9.1.2. Alternative ai file name del dispositivo

Poichè l'aggiunta o la rimozione dei dispositivi mass storage può causare delle modifiche ai file name per i dispositivi esistenti, è possibile che si verifichi il caso in cui lo storage non sia disponibile al momento del riavvio del sistema. Ecco un esempio della sequenza di eventi che portano a tale problema:

1. L'amministratore di sistema aggiunge un nuovo controller SCSI in modo tale che due nuove unità SCSI possono essere aggiunte al sistema (il bus SCSI esistente è completamente pieno)
2. Le unità SCSI originali (incluso la prima unità presente sul bus: `/dev/sda`) non vengono modificate in alcun modo
3. Questo sistema viene riavviato
4. L'unità SCSI formalmente conosciuta come `/dev/sda` possiede ora un nuovo nome, poichè la prima unità SCSI sul nuovo controller è ora `/dev/sda`

In teoria potrebbe essere un grosso problema. Tuttavia in pratica lo è di rado. È raramente un problema per un certo numero di fattori. Come prima cosa le riconfigurazioni hardware di questo tipo accadono molto raramente. In secondo luogo è probabile che l'amministratore di sistema ha previsto

un downtime per eseguire le modifiche necessarie; i periodi di downtime richiedono una pianificazione molto meticolosa, in quanto il lavoro fatto non deve richiedere un periodo superiore al tempo previsto. Questa pianificazione presenta il beneficio di evidenziare qualsiasi problematica relativa alle modifiche del nome del dispositivo.

Tuttavia alcune organizzazioni e configurazioni di sistema con molta probabilità si troveranno di fronte a tale problema. Le organizzazioni che necessitano delle riconfigurazioni frequenti dello storage, in modo da soddisfare le proprie esigenze, spesso utilizzano un hardware in grado di eseguire una riconfigurazione senza la necessità di un periodo di downtime. Questo tipo di hardware *hotpluggable* facilita la rimozione o l'aggiunta di storage. In tali circostanze le problematiche riguardanti il naming del dispositivo possono diventare un vero e proprio problema. Fortunatamente Red Hat Enterprise Linux presenta dei contenuti che rende meno problematico le modifiche del nome del dispositivo.

5.9.1.2.1. File System Label

Alcuni file system (i quali verranno affrontati più in avanti nella Sezione 5.9.2) hanno la possibilità di conservare un *label* — un carattere che può essere utilizzato per identificare unicamente i dati contenuti dal file system. I label possono essere utilizzati quando si monta il file system, eliminando la necessità di utilizzare il nome del dispositivo.

I label di un file system funzionano molto bene; tuttavia, essi devono essere unici. Se sono presenti più di un file system con lo stesso label, è possibile che l'accesso al file system desiderato non sia possibile. Notate che anche le configurazioni del sistema che non utilizzano un file system (per esempio alcuni database) non possono trarre vantaggio dai label di un file system.

5.9.1.2.2. Utilizzo di *devlabel*

Il software *devlabel* cerca di far fronte al problema del naming del dispositivo in modo diverso rispetto ai label del file system. Il software *devlabel* viene eseguito da Red Hat Enterprise Linux ogni qualvolta che viene eseguito il riavvio del sistema (e ogni qualvolta i dispositivi *hotpluggable* vengono inseriti o rimossi)

Quando viene eseguito *devlabel*, esso legge i propri file di configurazione (*/etc/sysconfig/devlabel*) in modo da ottenere un elenco di dispositivi per i quali è responsabile. Per ogni dispositivo presente nell'elenco, è disponibile un link simbolico (scelto dall'amministratore di sistema) e un UUID del dispositivo (Universal Unique Identifier).

Il comando *devlabel* assicura che il link simbolico si riferisca sempre al dispositivo specificato originariamente — anche se il nome del dispositivo è stato modificato. In questo modo, un amministratore di sistema è in grado di configurare un sistema in modo da far riferimento a */dev/projdisk* invece di */dev/sda12*.

Poiché l'UUID viene ottenuto direttamente dal dispositivo, *devlabel* deve solo eseguire una ricerca per l'UUID corrispondente, e aggiornare il link simbolico in modo appropriato.

Per maggiori informazioni su *devlabel*, consultate *Red Hat Enterprise Linux System Administration Guide*.

5.9.2. Informazioni di base sui file system

Red Hat Enterprise Linux include un supporto per numerosi file system, rendendo possibile così l'accesso ai file system presenti su altri sistemi operativi.

Questo è particolarmente utile in scenari di tipo dual-boot, e quando si migrano file da un sistema operativo ad un altro.

I file system supportati includono (ma non si limitano al seguente elenco):

- EXT2
- EXT3
- NFS
- ISO 9660
- MSDOS
- VFAT

Le seguenti sezioni affrontano in modo più dettagliato i file system più comuni.

5.9.2.1. EXT2

Fino a poco tempo fa i file system ext2 rappresentavano i file system standard di Linux. Per questo motivo, il suddetto file system è stato provato in modo molto accurato ed è considerato oggi uno dei più robusti in uso.

Tuttavia, non vi è un file system perfetto, e ext2 non fa eccezione. Un problema che viene comunemente riportato è quello relativo ad un controllo molto lungo sull'integrità di un file system, se il sistema non è stato terminato in modo corretto. Mentre questa caratteristica non è unica di ext2, la popolarità di ext2, combinata con l'avvento di unità disco più grandi, ha portato ad un controllo sull'integrità del file system sempre maggiore. Per questo motivo, era necessario trovare un'alternativa.

La seguente sezione descrive l'approccio preso per risolvere questo problema presente in Red Hat Enterprise Linux.

5.9.2.2. EXT3

Il file system ext3 viene creato usando ext2 e aggiungendo le capacità del journaling al codice base di ext2. Come file system journaling, ext3 mantiene costante lo stato del file system, eliminando la necessità dei controlli sull'integrità.

Questo può essere eseguito scrivendo tutte le modifiche del file system su di un journal su disco, il quale viene rimosso a intervalli regolari. Dopo il verificarsi di un evento inaspettato (come ad esempio un crash del sistema o un abbassamento di tensione), l'unica operazione che necessita di essere eseguita prima di rendere disponibile il file system, è quella di processare i contenuti del journal; in molti casi il tempo necessario per tale operazione è pari ad un secondo.

Poiché il formato dei dati di ext3 presenti sul disco è basato su ext2, è possibile accedere ad un file system ext3 su qualsiasi sistema in grado di leggere e scrivere un file system ext2 (senza il beneficio del journaling). Questo processo può rappresentare un grande beneficio per le organizzazioni che presentano alcuni sistemi ext3 ed altri ext2.

5.9.2.3. ISO 9660

Nel 1987, l'International Organization for Standardization (conosciuta come ISO), ha rilasciato il 9660 standard. ISO 9660 definisce come rappresentare i file su CD-ROM. Gli amministratori di sistema di Red Hat Enterprise Linux potranno vedere i dati formattati utilizzando ISO 9660 in due luoghi:

- CD-ROM
- File (generalmente riferiti come *immagini ISO*) contenenti file system ISO 9660 completi, da scrivere su media CD-R o CD-RW

L'ISO 9660 standard di base è limitato nelle sue funzionalità, maggiormente se confrontato a file system più moderni. I file name possono avere un massimo di otto caratteri ed è permessa una estensione

massima di tre caratteri. Tuttavia, diverse estensioni sono diventate molto comuni col trascorrere degli anni:

- Rock Ridge — Utilizza i campi unificati in ISO 9660, per fornire un supporto per contenuti del tipo file name 'long mixed-case', link simbolici, e directory nidificate (in altre parole, directory in grado di contenere altre directory)
- Joliet — Una estensione dell'ISO 9660 standard, sviluppato da Microsoft per permettere ai CD-ROM di contenere file name molto lunghi, utilizzando l'impostazione con caratteri di tipo Unicode

Red Hat Enterprise Linux è in grado di interpretare correttamente i file system ISO 9660 utilizzando le estensioni Rock Ridge e Joliet

5.9.2.4. MSDOS

Red Hat Enterprise Linux è in grado di supportare i file system di altri sistemi operativi. Come indicato dal nome per i file system di msdos, il sistema operativo originale in grado di supportare questo file system era l'MS-DOS® di Microsoft. Come in MS-DOS, un sistema Red Hat Enterprise Linux che accede un file system di msdos, è limitato a 8.3 file name. Allo stesso modo altri attributi come i permessi e l'ownership, non possono essere cambiati. Tuttavia, dal punto di vista di un file, il file system msdos è più che sufficiente per adempiere al proprio compito.

5.9.2.5. VFAT

Il file system vfat è stato utilizzato per la prima volta dal sistema operativo Windows® 95 di Microsoft. Un miglioramento rispetto al file system di msdos, i file name su di un sistema vfat possono essere più lunghi dei file name 8.3 di msdos. Tuttavia i permessi e l'ownership non possono essere cambiati.

5.9.3. Montaggio dei file system

Per accedere ogni file system è necessario prima *montarlo*. Montando un file system, indicherete a Red Hat Enterprise Linux di creare una partizione specifica (su di un dispositivo specifico) disponibile per il sistema. Allo stesso modo, quando l'accesso ad un file system particolare non è più desiderato, è necessario eseguire l'operazione di *umount* dello stesso.

Per poter montare qualsiasi file system è necessario specificare quanto segue:

- Un modo per identificare unicamente il disco fisso e la partizione desiderata, come ad esempio il file name del dispositivo, il label del file system, o il `devlabel-managed` symbolic link
- Una directory per mezzo della quale il file system montato deve essere disponibile (conosciuto anche come un *mount point*)

Le seguenti sezioni affrontano i mount point in modo dettagliato.

5.9.3.1. Mount Point

Se non avete una certa familiarità con i sistemi operativi di Linux (o simili), il concetto di mount point potrebbe sembrare un pò strano. Tuttavia, esso rappresenta uno dei metodi più flessibili e potenti per la gestione dei file system. Come molti altri sistemi operativi, una specificazione completa include il nome del file, diversi modi per l'identificazione della directory specifica nella quale risiede il file, e alcuni mezzi per identificare il dispositivo fisico sul quale può essere trovato il file.

Con Red Hat Enterprise Linux, viene utilizzato un approccio leggermente diverso. Come con altri sistemi operativi, una specificazione completa di un file include il nome e la directory nella quale il file risiede. Tuttavia non vi è alcun indicatore esplicito del dispositivo.

Il motivo è dato dal mount point. Su altri sistemi operativi, vi è una gerarchia della directory per ogni partizione. Tuttavia, sui sistemi simili a Linux c'è solo *una* gerarchia per l'intero sistema, e la suddetta gerarchia è valida anche per partizioni multiple. La chiave è il mount point. Quando viene montato un file system, lo stesso file system viene reso disponibile come un insieme di subdirectory sotto il mount point specificato.

Tale apparente difetto, in realtà risulta essere uno dei lati positivi. Infatti è possibile l'espansione di un file system di Linux, con ogni directory in grado di agire come un mount point per uno spazio aggiuntivo del disco.

Per esempio, considerate un sistema Red Hat Enterprise Linux che contiene una directory `/foo` nella propria directory root; il percorso completo per la directory sarà `/foo/`. Successivamente considerate un sistema con una partizione da montare, e che il suo mount point sia `/foo/`. Se la suddetta partizione presenta un file chiamato `bar.txt` nella sua directory superiore, una volta montata la partizione, sarete in grado di accedere il file con il seguente:

```
/foo/bar.txt
```

In altre parole, una volta montata la partizione, qualsiasi file letto oppure scritto in qualsiasi luogo sotto la directory `/foo/`, sarà letto da o scritto sulla stessa partizione.

Un mount point comunemente usato su molti sistemi Red Hat Enterprise Linux è `/home/` — questo perchè le directory di login per tutti gli account dell'utente, si trovano normalmente sotto `/home/`. Se si utilizza `/home/` come mount point, tutti i file degli utenti vengono scritti su di una partizione apposita, e non riempiranno il file system del sistema operativo.



Suggerimento

Poichè un mount point rappresenta solo una directory ordinaria, è possibile scrivere i file in una directory che viene utilizzata più avanti come mount point. Se ciò accade, cosa potrà accadere ai file che sono presenti nella directory originale?

Fino a quando la partizione è montata sulla directory, i file risultano non accessibili (il file system montato appare al posto dei contenuti della directory). Tuttavia, i file non verranno intaccati e saranno accessibili dopo aver eseguito una procedura di unmount della partizione.

5.9.3.2. Controllare cosa è stato montato

In aggiunta al montaggio e smontaggio dello spazio del disco, è possibile controllare anche cosa è stato montato. Sono disponibili diversi modi per eseguire questo processo:

- Visualizzando `/etc/mtab`
- Visualizzando `/proc/mounts`
- Emettendo il comando `df`

5.9.3.2.1. Visualizzando `/etc/mtab`

Il file `/etc/mtab` è un file normale che viene aggiornato dal programma `mount` ogni qualvolta i file system vengono montati o smontati. Ecco un esempio di `/etc/mtab`:

```
/dev/sda3 / ext3 rw 0 0
```

```

none /proc proc rw 0 0
usbdevfs /proc/bus/usb usbdevfs rw 0 0
/dev/sda1 /boot ext3 rw 0 0
none /dev/pts devpts rw,gid=5,mode=620 0 0
/dev/sda4 /home ext3 rw 0 0
none /dev/shm tmpfs rw 0 0
none /proc/sys/fs/binfmt_misc binfmt_misc rw 0 0

```



Nota bene

Il file `/etc/mtab` deve essere utilizzato per visualizzare lo stato dei file system attualmente montati. Non deve essere modificato manualmente.

Ogni riga rappresenta un file system attualmente montato e contiene i seguenti campi (da destra a sinistra):

- La specificazione del dispositivo
- Il mount point
- Il tipo di file system
- Sia che il file system sia stato montato in sola lettura (`ro`) o lettura e scrittura (`rw`), insieme con altre opzioni di montaggio
- Due campi non utilizzati che riportano degli zeri (per una compatibilità con `/etc/fstab`¹¹)

5.9.3.2.2. Visualizzando `/proc/mounts`

Il file `/proc/mounts` fa parte del `proc` virtual file system. Come con altri file sotto `/proc/`, il "file" `mounts` non risulta esistente su qualsiasi unità disco presente sul vostro sistema Red Hat Enterprise Linux.

Infatti, esso non risulta essere un file; ma risulta essere una rappresentazione dello stato del sistema, resa disponibile (dal kernel di Linux) sotto forma di file.

Usando il comando `cat /proc/mounts`, possiamo visualizzare lo stato di tutti i file system montati:

```

rootfs / rootfs rw 0 0
/dev/root / ext3 rw 0 0
/proc /proc proc rw 0 0
usbdevfs /proc/bus/usb usbdevfs rw 0 0
/dev/sda1 /boot ext3 rw 0 0
none /dev/pts devpts rw 0 0
/dev/sda4 /home ext3 rw 0 0
none /dev/shm tmpfs rw 0 0
none /proc/sys/fs/binfmt_misc binfmt_misc rw 0 0

```

Come possiamo vedere dall'esempio sopra riportato, il formato di `/proc/mounts` risulta essere molto simile a quello di `/etc/mtab`. Vi sono un certo numero di file system montati che non hanno nulla a che fare con le unità disco. Tra di essi si annoverano lo stesso file system `/proc/` (insieme con altri sue file system montati sotto `/proc/`), pseudo-ttys e la memoria condivisa.

Mentre il formato risulta essere non molto semplice, riferirsi a `/proc/mounts`, risulta essere il modo migliore per essere sicuri al 100% di sapere cosa è stato montato sul vostro sistema Red Hat Enterprise

11. Per maggiori informazioni consultate la Sezione 5.9.5.

Linux, in quanto il kernel è in grado di fornire queste informazioni. Altri metodi possono, in alcuni casi, essere non accurati.

Tuttavia, molto spesso vi troverete ad usare un comando in grado di fornire output più semplici (e utili) da leggere. La sezione successiva descrive il suddetto comando.

5.9.3.2.3. Emissione del comando `df`

Usando `/etc/mtab` o `/proc/mounts` sarete in grado di sapere i file system attualmente montati. Molto spesso sarete più interessati in un aspetto particolare riguardante i file system attualmente montati — la quantità di spazio disponibile su di essi.

Per questo motivo è possibile usare il comando `df`. Di seguito vengono riportati alcuni esempi di output di `df`:

Filesystem	1k-blocks	Used	Available	Use%	Mounted on
/dev/sda3	8428196	4280980	3719084	54%	/
/dev/sda1	124427	18815	99188	16%	/boot
/dev/sda4	8428196	4094232	3905832	52%	/home
none	644600	0	644600	0%	/dev/shm

Alcune differenze tra `/etc/mtab` e `/proc/mount` sono facilmente visibili:

- Viene visualizzato un testo facile da leggere
- Con l'eccezione del file system della memoria condivisa, vengono mostrati solo i file system basati sul disco
- Vengono visualizzati la misura totale, lo spazio utilizzato, lo spazio disponibile e la percentuale in uso

L'ultimo punto risulta probabilmente essere il più importante, in quanto ogni amministratore di sistema deve eventualmente far fronte al problema della carenza di spazio. Con `df` risulta più semplice identificare il problema.

5.9.4. Storage accessibile dalla rete con Red Hat Enterprise Linux

Sono presenti due tipi di tecnologie per l'implementazione di uno storage accessibile attraverso la rete con Red Hat Enterprise Linux:

- NFS
- SMB

Le seguenti sezioni descrivono le suddette tecnologie.

5.9.4.1. NFS

Come indicato dal nome, il Network File System (conosciuto come NFS) è un file system che può essere accesso tramite un collegamento di rete. Con altri file system, il dispositivo storage deve essere collegato direttamente al sistema locale. Tuttavia, con NFS questo processo non rappresenta una necessità, rendendo possibile così l'implementazione di diverse configurazioni, da server centralizzati del file system ai sistemi computerizzati senza disco.

Tuttavia diversamente da altri file system, NFS non indica alcun formato specifico su-disco. Esso si affida invece al supporto nativo del file system del sistema operativo, in modo da controllare l'I/O

attuale per l'unità del disco locale. NFS rende il file system disponibile a qualsiasi sistema operativo che esegue un client NFS compatibile.

Anche se originario della tecnologia UNIX e Linux, vale la pena notare che le implementazioni client NFS esistono anche per altri sistemi operativi, rendendo NFS una tecnica utile per condividere i file con diverse piattaforme.

Il server NFS utilizza il file di configurazione `/etc/exports` per controllare quale file system sia disponibile per il client. Per maggiori informazioni controllare la pagina `man exports(5)` e *Red Hat Enterprise Linux System Administration Guide*.

5.9.4.2. SMB

SMB è l'acronimo di *Server Message Block* ed è il nome per il protocollo di comunicazione utilizzato da vari sistemi operativi prodotti da Microsoft. SMB rende possibile la condivisione dello storage attraverso una rete. Le implementazioni odierne spesso utilizzano TCP/IP come mezzi di trasporto; precedentemente lo era NetBEUI.

Red Hat Enterprise Linux supporta SMB tramite il programma del server Samba. *Red Hat Enterprise Linux System Administration Guide* include le informazioni su come configurare Samba.

5.9.5. Montaggio automatico del file system con `/etc/fstab`

Se avete installato un nuovo sistema Red Hat Enterprise Linux, tutte le partizioni del disco definite e/o create durante l'installazione, sono configurate in modo da essere montate automaticamente ogni qualvolta il sistema esegue un riavvio. Tuttavia, cosa può accadere quando le unità disco aggiuntive vengono aggiunte ad un sistema dopo aver eseguito l'installazione? La risposta è "niente" in quanto il sistema non è stato configurato per montarle in modo automatico. Tuttavia questo processo è facilmente modificabile.

La risposta si trova nel file `/etc/fstab`. Questo file viene usato per controllare i file system montati al momento dell'avvio del sistema, ed anche per fornire i valori di default per altri file system che possono essere montati manualmente. Ecco un esempio di file `/etc/fstab`:

```

LABEL=/                /                ext3    defaults        1 1
/dev/sda1              /boot            ext3    defaults        1 2
/dev/cdrom             /mnt/cdrom       iso9660 noauto,owner,kudzu,ro 0 0
/dev/homedisk          /home            ext3    defaults        1 2
/dev/sda2              swap             swap    defaults        0 0

```

Ogni riga rappresenta un file system e contiene i seguenti campi:

- **Indicatore del file system** — Esso può essere per i file system basati sul disco, sia un file name del dispositivo (`/dev/sda1`), una specificazione del label del file system (`LABEL=/`), oppure un link simbolico gestito da `devlabel` (`/dev/homedisk`)
- **Mount point** — Ad eccezione delle partizioni swap, questo campo specifica il mount point utilizzato quando il file system è montato (`/boot`)
- **Tipo di file system** — Il tipo di file system presente sul dispositivo specificato (notate che `auto` può essere specificato in modo da selezionare una individuazione automatica del file system da montare, tale procedura può essere molto utile per le unità di media in grado di essere rimossi, come ad esempio i dischetti)
- **Opzioni mount** — Un elenco di opzioni separato da una virgola, in grado di essere usato per controllare il comportamento di `mount` (`noauto,owner,kudzu`)

- **Dump frequency** — Se viene usata la utility di backup `dump`, il numero presente in questo campo controlla la gestione di `dump` del file system specificato
- **Ordine di controllo dei file system** — Controlla l'ordine con il quale il controllore del file system `fsck` controlla l'integrità dei file system

5.9.6. Aggiunta/rimozione dello storage

Mentre molte delle fasi necessarie per aggiungere o rimuovere lo storage dipendono più dall'hardware del sistema che dal software, ci sono alcuni aspetti della procedura specifici per il vostro ambiente operativo. Questa sezione esplora le fasi necessarie per aggiungere e rimuovere lo storage specifico per Red Hat Enterprise Linux.

5.9.6.1. Aggiunta dello storage

Il processo per l'aggiunta di storage ad un sistema Red Hat Enterprise Linux è relativamente molto semplice. Ecco le fasi specifiche di Red Hat Enterprise Linux:

- Partizionamento
- Formattazione delle partizioni
- Aggiornamento di `/etc/fstab`

Le seguenti sezioni affrontano le fasi in modo più dettagliato.

5.9.6.1.1. Partizionamento

Una volta installata l'unità disco, è tempo di creare una o più partizioni per rendere disponibile lo spazio presente su Red Hat Enterprise Linux.

Ci sono molteplici modi per eseguire quanto detto:

- Usare il programma utility della linea di comando `fdisk`
- Utilizzando `parted`, un altro programma utility della linea di comando

Anche se i tool possono essere diversi, le fasi di base sono le stesse. Nel seguente esempio vengono riportati i comandi necessari per eseguire queste fasi usando `fdisk`:

1. Selezionare la nuova unità disco (il nome dell'unità può essere identificato seguendo le convenzioni per il naming del dispositivo affrontate nella Sezione 5.9.1). Utilizzando `fdisk`, è necessario includere il nome del dispositivo quando iniziate `fdisk`:

```
fdisk /dev/hda
```

2. Visualizzate la tabella riguardante la partizione dell'unità disco, in modo da assicurarvi che l'unità da partizionare sia quella corretta. Nel nostro esempio, `fdisk` visualizza la suddetta tabella utilizzando il comando `p`:

```
Command (m for help): p
```

```
Disk /dev/hda: 255 heads, 63 sectors, 1244 cylinders
Units = cylinders of 16065 * 512 bytes
```

Device	Boot	Start	End	Blocks	Id	System
/dev/hda1	*	1	17	136521	83	Linux
/dev/hda2		18	83	530145	82	Linux swap
/dev/hda3		84	475	3148740	83	Linux
/dev/hda4		476	1244	6176992+	83	Linux

3. Cancellate le partizioni non desiderate presenti sulla nuova unità disco. Questa procedura può essere eseguita utilizzando il comando `d` in `fdisk`:

```
Command (m for help): d
Partition number (1-4): 1
```

Tale processo sarà ripetuto per tutte le partizioni non necessarie presenti sull'unità disco.

4. Create la nuova partizione, assicuratevi di specificare la misura desiderata ed il tipo di file system. Utilizzando `fdisk` è necessario seguire un processo a due fasi — la prima è quella di creare la partizione (usando il comando `n`):

```
Command (m for help): n
Command action
e   extended
p   primary partition (1-4)
```

P

```
Partition number (1-4): 1
First cylinder (1-767): 1
Last cylinder or +size or +sizeM or +sizeK: +512M
```

La seconda è quella di impostare il tipo di file system (utilizzando il comando `t`):

```
Command (m for help): t
Partition number (1-4): 1
Hex code (type L to list codes): 82
```

Il tipo di partizione 82 rappresenta una partizione swap di Linux.

5. Salvate le vostre modifiche e uscite dal programma di partizione. Questo può essere fatto in `fdisk` utilizzando il comando `w`:

```
Command (m for help): w
```



Avvertenza

Quando partizionate la vostra nuova unità disco, è *necessario* assicurarsi che l'unità da partizionare sia quella corretta. In caso contrario, potreste partizionare una unità già in uso, causando una perdita di dati.

Assicuratevi altresì che abbiate precedentemente deciso la misura di partizionamento migliore. Affrontate questa fase in modo appropriato, in quanto la modifica della misura in un secondo momento, potrebbe risultare un processo molto difficoltoso.

5.9.6.1.2. Formattare le partizioni

La formattazione delle partizioni con Red Hat Enterprise Linux viene eseguita usando il programma utility `mkfs`. Tuttavia, `mkfs` non esegue il lavoro di scrittura delle informazioni specifiche del file system su di una unità disco; al contrario esso passa il controllo ad uno dei programmi che in effetti crea il file system stesso.

Questo è il momento giusto per dare uno sguardo alla pagina `man` di `mkfs.<fstype>` per il file system che avete selezionato. Per esempio, controllate la pagina `man` di `mkfs.ext3`, per poter consultare le diverse opzioni disponibili quando si crea un nuovo file system. In generale i programmi `mkfs.<fstype>` forniscono dei default per la maggior parte delle configurazioni; ecco alcune delle opzioni che gli amministratori di sistema cambiano più di frequente:

- Impostare un 'volume label' per un utilizzo futuro in `/etc/fstab`
- Su dischi fissi molto grandi, impostare una percentuale più bassa di spazio riservata per il superutente
- Impostare una misura del blocco non-standard e/o byte per un inode per quelle configurazioni che devono supportare sia file molto grandi che file molto piccoli
- Controllare la presenza di blocchi non corretti prima di eseguire la formattazione

Una volta creati i file system su tutte le partizioni appropriate, l'unità disco è configurata correttamente ed è pronta all'uso.

Successivamente, è sempre consigliabile controllare quello che avete svolto, montando manualmente le partizioni e assicurandovi che tutto sia in ordine. Una volta controllato il tutto, è tempo di configurare il vostro sistema Red Hat Enterprise Linux in modo da montare automaticamente il nuovo file system, ogni qualvolta che si esegue una procedura di riavvio.

5.9.6.1.3. Aggiornamento di `/etc/fstab`

Come mostrato in la Sezione 5.9.5, dovete aggiungere le righe necessarie su `/etc/fstab`, in modo da assicurarvi che il nuovo file system venga montato ogni volta che si esegue una procedura di riavvio. Una volta aggiornato `/etc/fstab`, testate il vostro lavoro emettendo un `mount` "incompleto", specificando solo il dispositivo o il mount point. Qualcosa di simile ai seguenti comandi risulta essere sufficiente:

```
mount /home
mount /dev/hda3
```

(Sostituire `/home` o `/dev/hda3` con il mount point o con il dispositivo idoneo alla vostra situazione.)

Se la voce `/etc/fstab` risulta essere corretta, `mount` è in grado di ottenere le informazioni, completando le operazioni di montaggio.

A questo punto potete essere sicuri che `/etc/fstab` sia configurato correttamente in modo da montare automaticamente il nuovo storage, ogni volta che il sistema esegue un avvio (ma se siete in grado di eseguire un riavvio rapido, eseguitelo — solo per essere maggiormente sicuri).

5.9.6.2. Rimozione dello storage

Il processo di rimozione dello storage da un sistema Red Hat Enterprise Linux risulta essere molto semplice. Ecco le fasi specifiche per Red Hat Enterprise Linux:

- Rimuovere le partizioni dell'unità disco da `/etc/fstab`
- Eseguire una procedura di `umount` delle partizioni attive dell'unità
- Cancellare i contenuti dell'unità disco

Le seguenti sezioni affrontano questi argomenti in modo più dettagliato

5.9.6.2.1. Rimuovere le partizioni dell'unità disco da `/etc/fstab`

Usando l'editor di testo, rimuovete la riga corrispondente alla partizione dell'unità disco presente nel file `/etc/fstab`. Potete identificare le righe corrette seguendo uno dei metodi riportati:

- Far corrispondere il mount point della partizione con le directory presenti nella seconda colonna di `/etc/fstab`

- Far corrispondere il file name del dispositivo con i nomi presenti nella prima colonna di `/etc/fstab`



Suggerimento

Assicuratevi di andare alla ricerca di qualsiasi riga in `/etc/fstab`, in grado di identificare le partizioni swap da rimuovere presenti sull'unità disco; esse possono facilmente sfuggire al vostro controllo.

5.9.6.2.2. Interruzione dell'accesso con `umount`

Successivamente, è necessario interrompere gli accessi all'unità disco. Per le partizioni che presentano dei file system attivi, ciò viene eseguito usando il comando `umount`. Se esiste sull'unità disco una partizione di swap, quest'ultima deve essere disattivata con il comando `swapoff`, oppure è necessario riavviare il sistema.

Per eseguire una procedura di `umount` con il comando `umount`, è necessario specificare il file name del dispositivo, oppure il mount point della partizione:

```
umount /dev/hda2
umount /home
```

È possibile eseguire una procedura di `umount` per una partizione, solo se quest'ultima non è in uso. Se invece non è possibile eseguire una tale procedura in un runlevel normale, eseguite l'avvio in modalità rescue, e rimuovere la voce `/etc/fstab` della partizione.

Se usate `swapoff` per disabilitare lo swapping su di una partizione, è necessario specificare il file name del dispositivo che rappresenta la partizione di swap:

```
swapoff /dev/hda4
```

Se non è possibile disabilitare lo swapping per una partizione di swap usando `swapoff`, eseguite l'avvio in modalità rescue, e rimuovete la voce `/etc/fstab` della partizione.

5.9.6.2.3. Cancellare i contenuti dell'unità disco

Cancellare i contenuti di una unità disco con Red Hat Enterprise Linux è un processo molto semplice.

Dopo avere eseguito una procedura di `umount` delle partizioni dell'unità disco, emettere il seguente comando (come utente root):

```
badblocks -ws <device-name>
```

Dove `<device-name>` rappresenta il file name dell'unità disco che desiderate cancellare, escluso il numero della partizione. Per esempio, `/dev/hdb` per la seconda unità ATA.

Il seguente output viene visualizzato mentre viene eseguito `badblocks`:

```
Writing pattern 0xaaaaaaaa: done
Reading and comparing: done
Writing pattern 0x55555555: done
Reading and comparing: done
Writing pattern 0xffffffff: done
Reading and comparing: done
Writing pattern 0x00000000: done
```

Reading and comparing: done

Ricordate che `badblocks` è in grado di scrivere quattro pattern di dati su ogni blocco presente sull'unità disco. Per unità molto grandi, questo processo può richiedere del tempo — molto spesso delle ore.



Importante

Molte compagnie (e agenzie governative), presentano metodi specifici per la cancellazione dei dati dalle unità disco, da altri media e per la loro conservazione. In questo caso assicuratevi *sempre* di aver capito e seguito questi requisiti; in molti casi se non si osservano i suddetti requisiti, si può andare incontro a problemi legali. L'esempio sopra riportato non deve essere considerato il metodo ultimo per la cancellazione dei dati da una unità.

Tuttavia, risulta essere molto più efficace che usare il comando `rm`. Questo perché quando cancellate un file usando `rm`, esso viene segnato solo come cancellato — in effetti il contenuto del file *non* viene cancellato.

5.9.7. Implementazione dei Disk Quota

Red Hat Enterprise Linux è in grado di monitorare l'utilizzo dello spazio del disco in base all'utente e al gruppo, utilizzando i disk quota. La seguente sezione fornisce una panoramica dei contenuti presenti nei disk quota con Red Hat Enterprise Linux.

5.9.7.1. Alcune informazioni sui disk quota

Con Red Hat Enterprise Linux, i disk quota presentano le seguenti caratteristiche:

- Implementazione per file system
- Controllo dello spazio per ogni utente
- Controllo dello spazio per ogni gruppo
- Controllo sull'utilizzo del blocco del disco
- Controllo sull'utilizzo dell'inode del disco
- Limiti hard
- Limiti soft
- Grace period

Le seguenti sezioni descrivono le suddette caratteristiche in modo più dettagliato.

5.9.7.1.1. Implementazione per file system

Con Red Hat Enterprise Linux i disk quota possono essere usati in base al sistema. In altre parole, i disk quota possono essere abilitati o disabilitati in modo individuale per ogni file system.

Ciò garantisce una certa flessibilità all'amministratore di sistema. Per esempio, se la directory `/home/` era presente sul proprio file system, i disk quota possono essere abilitati nel file system stesso, imponendo un uso equo del disco da parte di tutti gli utenti. Tuttavia il file system root potrebbe essere lasciato senza disk quotas, eliminando la complessità della gestione dei disk quota stessi su di un file system dove risiede il solo sistema operativo.

5.9.7.1.2. Controllo dello spazio in base all'utente

Disk quotas è in grado di controllare lo spazio per ogni utente. Ciò significa che l'utilizzo fatto da ogni utente è controllato in modo individuale. Inoltre le limitazioni relative all'uso (questo argomento viene affrontato nelle sezioni successive) viene stabilito in base ad ogni utente.

Avendo la possibilità di controllare e imporre ad ogni utente un determinato uso del disco, è possibile per ogni amministratore di sistema, assegnare limiti diversi a seconda della responsabilità e delle necessità dello storage di ogni utente.

5.9.7.1.3. Controllo dello spazio per ogni gruppo

I disk quota sono anche in grado di eseguire un controllo sull'utilizzo dello spazio del disco in base al gruppo. Questa operazione risulta essere ideale per ogni organizzazione che utilizza i gruppi per combinare utenti diversi all'interno di una risorsa singola di un progetto.

Impostando i group-wide disk quota, l'amministratore di sistema è in grado di gestire più facilmente l'utilizzo dello storage, dando agli utenti individuali solo il disk quota necessario per il loro uso privato, mentre saranno in grado d'impostare disk quota più grandi, idonei per progetti che necessitano un numero di utenti maggiore. Questo processo risulta essere particolarmente utile per quelle organizzazioni che usano un meccanismo "chargeback", per assegnare i costi del centro dati a quelle sezioni ed ai team, che fanno uso delle risorse del centro dati.

5.9.7.1.4. Controllo sull'utilizzo del blocco del disco

Il disk quota è in grado di controllare l'utilizzo del blocco del disco. Poichè tutti i dati conservati su di un file system sono conservati in blocchi, i disk quotas sono in grado di correlare direttamente i file creati e cancellati su di un file system con la quantità di storage richiesto dai file in questione.

5.9.7.1.5. Controllo sull'utilizzo dell'inode del disco

In aggiunta al controllo sull'uso del blocco del disco, il disk quota è in grado di controllare anche l'utilizzo dell'inode. Con Red Hat Enterprise Linux, gli inode vengono usati per conservare diverse parti del file system, ma quello che più conta è che gli inode contengono le informazioni per ogni file. Per questo motivo, controllando l'utilizzo dell'inode, è possibile monitorare la creazione di nuovi file.

5.9.7.1.6. Limiti hard

Il limite hard rappresenta il numero massimo di blocchi (o inodi) che possono essere usati momentaneamente da un utente (o gruppo). Ogni tentativo di utilizzare un blocco o un inode singolo superiore a tale limite è destinato a fallire.

5.9.7.1.7. Limiti soft

Un limite soft è il numero massimo di blocchi (o inodi) che possono essere usati permanentemente da un utente (o gruppo).

Il limite soft viene impostato generalmente al di sotto del limite hard. Ciò permette a tutti gli utenti di eccedere temporaneamente il loro limite soft, permettendo così di portare a termine il proprio compito, e conferendo loro un tempo sufficiente durante il quale essi possono controllare i propri file scendendo così al di sotto del proprio limite soft.

5.9.7.1.8. Grace Period

Come precedentemente accennato, qualsiasi uso riguardante il disco che vada oltre il limite soft, è temporaneo. È il così detto ‘grace period’ che determina il periodo nel quale un utente (o gruppo), può estendere il proprio uso al di sopra del limite soft, avvicinandosi al proprio limite hard.

Se un utente continua a fare un uso maggiore a quello imposto dal proprio limite soft, e se il grace period è scaduto, non gli verrà permesso alcun utilizzo addizionale del disco, questo fino a quando l’utente (o il gruppo), ha ridotto il proprio uso fino a scendere al di sotto del limite soft.

Il grace period può essere espresso in secondi, minuti, ore, giorni, settimane o mesi, dando all’amministratore una certa libertà nel determinare la quantità di tempo necessario agli utenti per scendere al di sotto del proprio limite soft sull’uso del proprio disco.

5.9.7.2. Abilitare il disk quota



Nota bene

Le seguenti sezioni forniscono una breve panoramica delle fasi necessarie per abilitare i disk quota con Red Hat Enterprise Linux. Per maggiori informazioni, consultare il capitolo riguardante i disk quota presente nella *Red Hat Enterprise Linux System Administration Guide*.

Per usare i disk quotas, è necessario abilitarli. Questo processo è composto dalle seguenti fasi:

1. Modifica di `/etc/fstab`
2. Rimontaggio del file system
3. Esecuzione del `quotacheck`
4. Assegnazione dei quota

Il file `/etc/fstab` controlla il montaggio del file system con Red Hat Enterprise Linux. Poiché i disk quota vengono implementati in base al file system, sono possibili due opzioni — `usrquota` e `grpquota` — le quali devono essere aggiunti a file in questione per abilitare i disk quota.

L’opzione `usrquota` abilita i disk quotas basati sull’utente, mentre l’opzione `grpquota` abilita i quotas basati sul gruppo. Una o entrambe queste opzioni, possono essere abilitate posizionandole nel campo delle opzioni per il file system desiderato.

Bisogna eseguire quindi una procedura di unmount del file system in questione, per poi rimontarlo in modo da abilitare le opzioni relative al disk quota.

Successivamente, il comando `quotacheck` viene usato per creare i file del disk quota, e raccogliere le informazioni riguardanti l’uso corrente dai file già esistenti. I file del disk quota (chiamati `aquota.user` e `aquota.group` per quota basati sull’utente e sul gruppo), contengono le informazioni necessarie relative al quota, e risiedono nella directory root del file system.

Per assegnare i disk quota, usare il comando `edquota`.

Questo programma utility utilizza un editor di testo per visualizzare le informazioni del quota per l’utente ed il gruppo specificato come parte del comando `edquota`. Ecco un esempio:

```
Disk quotas for user matt (uid 500):
Filesystem      blocks      soft      hard      inodes      soft      hard
/dev/md3        6618000      0         0         17397        0         0
```

Ciò mostra che l'utente matt utilizza 6GB di spazio, e oltre 17000 inode. Nessun valore di quota (soft o hard) è stato impostato per i blocchi del disco e per gli inodi, ciò significa che non vi è alcun limite per lo spazio del disco e per gli inode riguardante l'utente in questione..

Utilizzando l'editor di testo che visualizza le informazioni sul disk quota, l'amministratore di sistema è in grado di modificare i limiti soft e hard:

```
Disk quotas for user matt (uid 500):
Filesystem      blocks      soft      hard      inodes      soft      hard
/dev/md3        6618000    6900000    7000000    17397        0        0
```

In questo esempio, è stato conferito all'utente matt un limite soft di 6.9GB e un limite hard di 7GB. Non è stato impostato per il suddetto utente alcun limite soft o hard sull'inode.



Suggerimento

Il programma `edquota` può essere usato per impostare il grace period per file system, utilizzando l'opzione `-t`.

5.9.7.3. Gestione dei disk quota

È necessario un certo metodo per la gestione del supporto dei disk quota con Red Hat Enterprise Linux. Tutto ciò che è richiesto è il seguente:

- Generare dei rapporti sull'utilizzo dello spazio ad intervalli regolari (e seguendo gli utenti che mostrano di avere delle particolari difficoltà nella gestione dello spazio del disco a loro riservato)
- Assicurarsi che i disk quota siano accurati

Creare un rapporto sull'uso del disco, comporta l'esecuzione del programma utility `repquota`. Usando il comando `repquota /home` si ottiene il seguente output:

```
*** Report for user quotas on device /dev/md3
Block grace time: 7days; Inode grace time: 7days
                                Block limits
User                             soft  hard  grace  used  soft  hard  grace
-----
root      --   32836      0      0      4      0      0
matt      -- 6618000 6900000 7000000 17397      0      0
```

Per maggiori informazioni su `repquota`, consultate *Red Hat Enterprise Linux System Administration Guide*, nel capitolo riguardante i disk quota.

Ogni qualvolta si esegue una procedura di unmount non corretta di un file system (per esempio a causa di un crash del sistema), è necessario eseguire `quotacheck`. Tuttavia, molti amministratori di sistema consigliano di eseguire `quotacheck` regolarmente, anche se non si è verificato il suddetto crash.

Il processo è simile all'uso iniziale di `quotacheck` quando si abilitano i disk quota.

Ecco un esempio del comando `quotacheck`:

```
quotacheck -avug
```

Il modo più semplice per eseguire `quotacheck` regolarmente, è quello di usare `cron`. Molti amministratori di sistema eseguono `quotacheck` una volta alla settimana, questo però dipende dalle vostre specifiche condizioni.

5.9.8. Creazione dei RAID array

In aggiunta al supporto delle soluzioni hardware RAID, Red Hat Enterprise Linux supporta anche il software RAID. Per creare gli array RAID software sono disponibili due metodi:

- Durante una installazione di Red Hat Enterprise Linux
- Dopo aver installato Red Hat Enterprise Linux

Le seguenti sezioni affrontano questi metodi in modo più dettagliato.

5.9.8.1. Durante una installazione di Red Hat Enterprise Linux

Durante il processo normale di installazione di Red Hat Enterprise Linux, è possibile creare i RAID array. Questo processo può essere eseguito durante la fase di partizionamento del disco.

Per iniziare, partizionate manualmente le vostre unità utilizzando **Disk Druid**. Come prima cosa, dovete creare una nuova partizione del tipo "software RAID." Successivamente selezionate le unità che desiderate implementare nel RAID array nel campo delle **Unità Autorizzate**. Continuate selezionando la misura scelta e se desiderate la partizione come partizione primaria.

Una volta create tutte le partizioni necessarie per il RAID array, allora potete utilizzare il pulsante **RAID** per creare gli array. Vi verrà presentata una finestra di dialogo dove è possibile selezionare il mount point dell'array, il tipo di file system, il nome del dispositivo RAID, il livello di RAID, e le partizioni "software RAID" sulle quali viene basato l'array in questione.

Una volta creati gli array desiderati, il processo d'installazione continua normalmente.



Suggerimento

Per maggiori informazioni sulla creazione di software RAID array durante il processo di installazione di Red Hat Enterprise Linux, consultate *Red Hat Enterprise Linux System Administration Guide*.

5.9.8.2. Dopo aver installato Red Hat Enterprise Linux

Creare un RAID array dopo aver installato Red Hat Enterprise Linux risulta essere un po' più complesso. Come per l'aggiunta di qualsiasi disk storage, l'hardware necessario deve essere prima installato e configurato in modo corretto.

Il partizionamento è leggermente diverso per RAID rispetto alle singole unità disco. Invece di selezionare un tipo di partizione di "Linux" (tipo 83) o "Linux swap" (tipo 82), tutte le partizioni che devono far parte di RAID array devono essere impostate su "Linux raid auto" (tipo fd).

Successivamente è necessario creare il file `/etc/raidtab`. Questo file è responsabile per la configurazione corretta di tutti i RAID array presenti sul vostro sistema. Il formato del file (il quale viene documentato nella pagina `man raidtab(5)`), risulta essere relativamente semplice. Ecco un esempio riguardante la voce `/etc/raidtab` per un RAID 1 array:

```
raiddev                /dev/md0
raid-level              1
nr-raid-disks          2
chunk-size             64k
persistent-superblock  1
nr-spare-disks         0
    device              /dev/hda2
    raid-disk           0
    device              /dev/hdc2
```

```
raid-disk      1
```

Alcune delle sezioni più evidenti in questa voce sono:

- `raiddev` — Mostra il file name del dispositivo per il RAID array¹²
- `raid-level` — Definisce il livello RAID che deve essere usato da questo RAID array
- `nr-raid-disks` — Indica il numero delle partizioni del disco fisico che devono far parte di questo array
- `nr-spare-disks` — Il software RAID con Red Hat Enterprise Linux permette la definizione di una o più partizioni di riserva del disco; queste partizioni possono sostituire automaticamente un disco difettoso
- `device, raid-disk` — Insieme essi definiscono le partizioni del disco fisico che costituiscono il RAID array

Successivamente è necessario creare il RAID array. Questo può essere fatto con il programma `mkraid`. Usando il nostro file `/etc/raidtab`, siamo in grado di creare il RAID array `/dev/md0` usando il seguente comando:

```
mkraid /dev/md0
```

Il RAID array `/dev/md0` è ora pronto per essere formattato e montato. Il processo a questo punto non è diverso dal processo di formattazione e montaggio di una unità singola del disco.

5.9.9. Gestione giornaliera dei RAID Array

Il RAID array non ha bisogno di molta manutenzione per mantenere la propria operatività. Se non si verifica alcun problema hardware, l'array dovrebbe funzionare proprio come se fosse presente una unità singola. Tuttavia proprio come un amministratore di sistema dovrebbe controllare periodicamente lo stato di tutte le unità disco presenti sul sistema, anche lo stato del RAID array dovrebbe essere controllato ad intervalli regolari.

5.9.9.1. Controllo dello stato dell'array con `/proc/mdstat`

Il file `/proc/mdstat` rappresenta il modo migliore per controllare lo stato di tutti i RAID array presenti su di un sistema particolare. Ecco un esempio di `mdstat` (da visualizzare con l'ausilio del comando `cat /proc/mdstat`):

```
Personalities : [raid1]
read_ahead 1024 sectors
md1 : active raid1 hda3[0] hdc3[1]
      522048 blocks [2/2] [UU]

md0 : active raid1 hda2[0] hdc2[1]
      4192896 blocks [2/2] [UU]

md2 : active raid1 hda1[0] hdc1[1]
      128384 blocks [2/2] [UU]

unused devices: <none>
```

12. Nota bene, poichè il RAID array è composto da uno spazio partizionato, il file name del dispositivo di un RAID array, non riflette alcuna informazione riguardante il livello della partizione.

Sono presenti su questo sistema tre RAID array (tutti RAID 1). Ogni RAIDarray possiede la propria sezione in `/proc/mdstat` e contiene le seguenti informazioni:

- Il nome del dispositivo del RAID array (non includendo la parte `/dev/`)
- Lo stato del RAID array
- Il RAID-level del RAID array
- Le partizioni fisiche che compongono l'array (seguite dal numero dell'unità array della partizione)
- La misura dell'array
- Il numero dei dispositivi configurati contro il numero dei dispositivi operativi presenti nell'array
- Lo stato di ogni dispositivo configurato nell'array (U lo stato del dispositivo è OK, e _ che indica la presenza di alcuni problemi nel dispositivo)

5.9.9.2. Creazione di un RAID array con `raidhotadd`

Se `/proc/mdstat` mostra la presenza di un errore con uno dei RAID array, il programma utility `raidhotadd` deve essere usato per creare nuovamente l'array. Ecco le fasi necessarie per tale processo:

1. Determinare quale unità disco contiene la partizione in questione
2. Correggere il problema (molto probabilmente sostituendo l'unità)
3. Partizionare la nuova unità in modo che le partizioni presenti su di essa, siano *identici* a quelle presenti sull'altra unità dell'array
4. Emettere il seguente comando:
`raidhotadd <raid-device> <disk-partition>`
5. Controllare `/proc/mdstat` in modo da verificare il processo



Suggerimento

Ecco un comando che può essere utilizzato per controllare il suddetto processo:

```
watch -n1 cat /proc/mdstat
```

Questo comando visualizza i contenuti di `/proc/mdstat`, aggiornandolo ogni secondo.

5.9.10. Logical Volume Management

Red Hat Enterprise Linux include il supporto per LVM. LVM può essere configurato mentre Red Hat Enterprise Linux viene installato, oppure può essere configurato dopo il completamento dell'installazione. LVM sotto Red Hat Enterprise Linux supporta il grouping fisico dello storage, il processo di ridimensionamento del logical volume, e la migrazione dei dati al di fuori di un volume fisico specifico.

Per maggiori informazioni su LVM, consultate *Red Hat Enterprise Linux System Administration Guide*.

5.10. Risorse aggiuntive

Questa sezione include le risorse aggiuntive che possono essere utilizzate per ottenere maggiori informazioni sulle tecnologie inerenti lo storage e gli argomenti specifici di Red Hat Enterprise Linux affrontati in questo capitolo.

5.10.1. Documentazione installata

Le seguenti risorse vengono installate nel corso di una installazione tipica di Red Hat Enterprise Linux, e possono essere di aiuto per approfondire gli argomenti affrontati in questo capitolo.

- Pagina man di `exports(5)` — Per maggiori informazioni sul formato del file di configurazione NFS.
- Pagina man di `fstab(5)` — Per maggiori informazioni sul formato del file di configurazione del file system.
- Pagina man di `swapoff(8)` — Per maggiori informazioni su come disabilitare le partizioni swap.
- Pagina man di `df(1)` — Per maggiori informazioni su come visualizzare le informazioni riguardanti l'uso dello spazio del disco su file system montati
- Pagina man di `fdisk(8)` — Per saperne di più sul programma utility di gestione per la tabella della partizione.
- Pagine man di `mkfs(8)`, `mke2fs(8)` — Per saperne di più sui programmi utility per la creazione di questi file system.
- Pagina man di `badblocks(8)` — Per saperne di più su come eseguire un test di un dispositivo per la verifica di blocchi non corretti.
- Pagina man di `quotacheck(8)` — Imparate a verificare l'utilizzo dei blocchi e dell'inode da parte degli utenti e dei gruppi, e creare facoltativamente i file del disk quota.
- Pagina man di `edquota(8)` — Per saperne di più sul programma utility per la gestione di questo disk quota.
- Pagina man di `repquota(8)` — Per saperne di più sul programma utility per il riporto di questo disk quota.
- Pagina man di `raidtab(5)` — Per saperne di più sul formato del file di configurazione del software RAID.
- Pagina man di `mkraid(8)` — Per saperne di più sul programma utility per la creazione del software RAID array.
- Pagina man di `mdadm(8)` — Per saperne di più sul programma utility per la gestione di questo software RAID array.
- Pagina man di `lvm(8)` — Per saperne di più su Logical Volume Management.
- Pagina man di `devlabel(8)` — Per saperne di più sull'accesso continuo del dispositivo storage.

5.10.2. Siti web utili

- <http://www.pcguides.com/> — Un sito molto utile con molte informazioni sulle varie tecnologie storage.
- <http://www.tldp.org/> — Il Linux Documentation Project presenta HOWTOs e mini-HOWTOs in grado di fornire delle panoramiche sulle tecnologie storage relative a Linux.

5.10.3. Libri correlati

I seguenti libri affrontano le diverse problematiche relative allo storage e rappresentano delle risorse importanti per gli amministratori di Red Hat Enterprise Linux.

- *Red Hat Enterprise Linux Installation Guide*; Red Hat, Inc. — Contiene le istruzioni sulla partizione delle unità disco durante il processo di installazione di Red Hat Enterprise Linux, e una panoramica sulle partizioni del disco.
- *Red Hat Enterprise Linux Reference Guide*; Red Hat, Inc. — Contiene informazioni dettagliate sulla struttura della directory usata in Red Hat Enterprise Linux, e una panoramica di NFS.
- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — Include i capitoli riguardanti i file system, RAID, LVM, `devlabel`, `partitioning`, `disk quota`, NFS, e Samba.
- *Linux System Administration: A User's Guide* di Marcel Gagne; Addison Wesley Professional — Contiene le informazioni riguardanti i permessi sul gruppo e sull'utente, informazioni sui file system, sul `disk quota`, NFS e Samba.
- *Linux Performance Tuning and Capacity Planning* di Jason R. Fink and Matthew D. Sherer; Sams — Contiene le informazioni sul disco, su RAID, e sulle prestazioni NFS.
- *Linux Administration Handbook* di Evi Nemeth, Garth Snyder, and Trent R. Hein; Prentice Hall — Contiene informazioni sui file system, sulla gestione delle unità disco, NFS e Samba.

Capitolo 6.

Gestione degli User Account e dell'accesso alle risorse

La gestione degli *user account* e dei *gruppi* rappresenta una parte essenziale della gestione del sistema all'interno di una organizzazione. Per eseguire tale operazione in modo efficace, un buon amministratore di rete deve prima capire cosa sono gli user account ed i gruppi ed il loro funzionamento.

Lo scopo principale sull'uso di un user account è quello di verificare l'identità di ogni individuo tramite l'uso di un computer. Quello secondario (ma anch'esso importante) è quello di permettere di eseguire un tailoring individuale delle risorse e dei privilegi d'accesso.

Le risorse possono includere file, directory e dispositivi. Il controllo dell'accesso alle suddette risorse, rappresenta una fase importante del lavoro giornaliero di un amministratore di rete, spesso l'accesso ad una risorsa viene controllata dai gruppi. I gruppi sono delle costruzioni logiche che possono essere usati per eseguire un cluster degli user account per un unico scopo. Per esempio, se una organizzazione possiede amministratori multipli di sistema, gli stessi possono essere messi insieme in un unico gruppo di amministratori. Al gruppo, successivamente, può essere garantito il permesso di accedere alle risorse più importanti del sistema. In questo modo i gruppi possono rappresentare un tool molto potente per la gestione delle risorse e degli accessi.

Le seguenti sezioni affrontano gli user account ed i gruppi in modo più dettagliato.

6.1. Gestione degli User Account

Come precedentemente detto, gli user account rappresentano il metodo grazie al quale un individuo viene identificato ed autenticato. Gli user account possiedono diversi componenti. In primo luogo vi è il nome utente o 'username', poi la password seguita dalle informazioni di controllo per l'accesso.

Le seguenti sezioni affrontano ogni componente in modo più dettagliato.

6.1.1. Nome utente

Sotto il punto di vista del sistema, il nome utente rappresenta la risposta alla domanda "Chi sei?" Esso richiede una caratteristica molto importante —, deve essere unico. In altre parole, ogni utente deve avere un nome utente diverso da tutti gli altri utenti presenti nel sistema.

Per questo motivo, è importante determinare — in anticipo — come creare un nome utente. Altrimenti potreste trovarvi nella posizione di reagire ogni qualvolta un nuovo utente richieda un account.

Avrete bisogno quindi di avere una naming convention per i vostri user account.

6.1.1.1. Naming Convention

Creando una naming convention per il nome utente, sarete in grado di evitare molti grattacapi. Invece d'inventare nomi durante la prosecuzione del vostro lavoro (e trovarsi in difficoltà nella creazione di un nome accettabile), pianificatevi prima ideando una convenzione da usare per gli user account. La vostra naming convention può essere molto semplice, ma è possibile anche che la sola descrizione potrebbe richiedere diverse pagine del documento.

La natura esatta della vostra naming convention dovrebbe considerare diversi fattori:

- La dimensione della vostra organizzazione
- La struttura della vostra organizzazione

- La natura della vostra organizzazione

La dimensione della vostra organizzazione è importante in quanto indica il numero di utenti che la vostra naming convention è in grado di supportare. Per esempio, una organizzazione molto piccola potrebbe far sì che gli utenti siano in grado di usare il proprio nome. Per una organizzazione più grande la suddetta naming convention potrebbe non funzionare.

La struttura di una organizzazione potrebbe anche orientarsi facilmente sulla naming convention più appropriata. Per organizzazioni con una struttura molto rigida e definita, potrebbe risultare idoneo includere elementi della stessa struttura all'interno della naming convention. Per esempio potreste includere i codici di una sezione della vostra organizzazione come parte di ogni nome utente.

A causa della natura della vostra organizzazione, si potrebbe verificare che alcune naming convention possano risultare più appropriate di altre. Una organizzazione che tratta con dati riservati potrebbe scegliere una naming convention in grado di fare riferimento all'individuo con parte del proprio nome. In tali organizzazioni il nome utente di Maggie McOmie potrebbe divenire LUH3417.

Ecco altre naming convention che altre organizzazioni hanno usato in precedenza:

- Nome (john, paul, george, ecc.)
- Cognome (smith, jones, brown, ecc.)
- Iniziale del proprio nome seguita dal cognome (jsmith, pjones, gbrown, ecc.)
- Cognome seguito dal codice del reparto (smith029, jones454, brown191, ecc.)



Suggerimento

Siate a conoscenza che se la vostra naming convention include diversi tipi di dati, ne potrebbe scaturire un nome utente con contenuto offensivo o ridicolo. Per questo motivo, anche se possedete un metodo di creazione automatizzato, è suggeribile avere una forma di revisione.

In comune con la naming convention descritta qui di seguito, è la possibilità, in accordo con la naming convention stessa, di conferire a due individui lo stesso nome utente. Ciò viene chiamato *collision*. A causa della unicità di ogni nome utente, è necessario risolvere la suddetta problematica. La sezione seguente mostra come fare.

6.1.1.1.1. Come affrontare le Collision

Le collisioni sono una prassi — non ha importanza quanto proviate, è sicuro che prima o poi vi troverete di fronte ad una collision. È importante pianificare questa eventualità durante la creazione di una naming convention. Ci sono diversi modi per fare questo:

- Aggiungere una sequenza numerica al nome utente es.(smith, smith1, smith2, ecc.)
- Aggiungere dati specifici dell'utente al nome utente in questione (smith, esmith, eksmith, ecc.)
- Aggiungere informazioni inerenti l'organizzazione al nome utente (smith, smith029, smith454, ecc.)

Implementare un metodo per risolvere le problematiche delle collision, rappresenta una necessità nel contesto della naming convention. Tuttavia, per un individuo esterno all'organizzazione, il lavoro di determinazione del nome utente può risultare più difficoltoso. Per questo motivo è possibile che si possa verificare l'invio di email ad indirizzi sbagliati.

6.1.1.2. Come affrontare i cambiamenti dei nomi

Se la vostra organizzazione usa una naming convention basata sul nome di ogni utente, è probabile che dobbiate affrontare la problematica dovuta al cambiamento del nome. Anche se il nome di una persona non varia, potrebbe essere necessario di tanto in tanto, modificare il nome utente. Il motivo potrebbe essere dato dall'insoddisfazione di un dirigente dell'organizzazione, che usa la sua autorità per ottenere un nome utente "più appropriato".

Non ha importanza il motivo per quale si debba cambiare il nome utente, ma nel fare questo tenete presente le seguenti problematiche:

- Eseguire i cambiamenti su *tutti* i sistemi interessati
- Mantenere sempre aggiornati e costanti le informazioni sull'identificazione dell'utente
- Cambiare l'ownership di tutti i file e altre risorse specifiche dell'utente (se necessario)
- Gestire le problematiche relative alle email

Innanzitutto, è importante assicurarsi che il nuovo nome utente venga inoltrato a tutti i sistemi, dove il nome utente originale era precedentemente in uso. Se questa operazione non viene eseguita, qualsiasi funzione del sistema operativo che si affida al nome utente potrebbe funzionare in alcuni sistemi, ma non funzionare in altri. Alcuni sistemi operativi usano un controllo d'accesso basato sul nome utente, tali sistemi potrebbero essere particolarmente vulnerabili ai problemi dovuti alla variazione del nome utente stesso.

Molti sistemi operativi usano dei numeri d'identificazione utente per molte funzioni 'user-specific'. Per minimizzare i problemi dovuti al cambiamento del nome utente, cercate di mantenere costante il numero d'identificazione tra il nuovo utente ed il vecchio. Se non tenete presente quanto detto, l'utente potrebbe non essere in grado di accedere i propri file e altre risorse precedentemente disponibili.

Se il suddetto numero d'identificazione deve essere cambiato, è necessario cambiare l'ownership per tutti i file e per le risorse del tipo 'user-specific', in modo da riflettere le nuove informazioni sull'identificazione dell'utente. Questo può rappresentare un processo error-prone, in quanto potreste sempre tralasciare qualcosa.

Le problematiche riguardanti le email sono con tutta probabilità, l'area dove il cambiamento del nome utente potrebbe essere più complicato. Il motivo principale potrebbe essere quello dove le email inviate al vecchio nome utente, potrebbero non essere recapitate a quello nuovo.

Sfortunatamente, le problematiche che si verificano quando si effettua un cambiamento del nome utente sono, per quanto riguarda le email, molteplici. La più semplice potrebbe essere quella dove le persone non sono a conoscenza del nuovo nome utente. A prima vista potrebbe non risultare un problema serio —, in quanto risulta molto semplice informare tutto il personale interno di una organizzazione. Ma cosa dire di coloro esterni all'organizzazione stessa che hanno inviato email alla persona in questione? Come si potrebbe effettuare la notifica di quanto accaduto? E per quanto riguarda le mailing list (sia quelle interne che esterne)? Come possono essere aggiornate?

Purtroppo non vi è una risposta semplice a questo. La risposta migliore potrebbe essere quella di creare una email alias in modo tale che tutte le email inviate, possano essere automaticamente inoltrate al nuovo nome utente. L'utente a sua volta dovrebbe informare coloro che inviano tali email, che il nome utente è cambiato. Con il tempo, sempre meno email saranno inviate all'email alias, permettendo così la rimozione dello stesso.

Mentre l'uso degli alias, in alcuni casi, potrebbe dare luogo a interpretazioni errate (come ad esempio un utente conosciuto ora come esmith 'nuovo nome utente', viene sempre indicato come ejones 'vecchio nome utente'), gli stessi alias rappresentano l'unico modo per garantire che le email arrivino alla persona corretta.

**Importante**

Se usate gli alias, assicuratevi di seguire tutte le fasi necessarie per evitare il riuso del vecchio nome utente. Se non seguite le suddette fasi, ed un nuovo utente riceve un vecchio nome utente, la consegna delle email (sia per l'utente originale che per l'utente nuovo), potrebbe risultare molto confusa. L'esatta natura di tale problema dipende dall'implementazione, sul vostro sistema operativo, della consegna delle email. I due sintomi più comuni potrebbero essere:

- Il nuovo utente potrebbe non ricevere le email — le stesse andranno all'utente originale.
- L'utente originale improvvisamente non riceve più le email — le stesse vengono inoltrate al nuovo utente.

6.1.2. Password

Se il nome utente è in grado di fornire una risposta alla domanda, "chi sei?", la password è la risposta alla domanda seguente:

"Provalo!"

In altri termini, una password fornisce il mezzo tramite il quale viene riconosciuta una persona. L'efficacia dello schema di autenticazione basato sulla password si affida pesantemente su diversi aspetti:

- Segretezza della password
- La difficoltà ad indovinare la password
- La resistenza della password ad un attacco forzato

Le password che rispettano i suddetti presupposti vengono chiamate *forti*, mentre ovviamente quelle che non seguono lo schema sopra riportato vengono chiamate *deboli*. La creazione di password forti risulta essere molto importante per la sicurezza dell'organizzazione, in quanto più forti esse sono, minore è la possibilità di indovinarle. Sono disponibili due opzioni per attuare l'uso di password forti:

- Creazione delle password per tutti gli utenti da parte dell'amministratore del sistema.
- L'amministratore potrebbe lasciare il compito della creazione della password direttamente all'utente, verificando però che le password create siano sufficientemente forti.

La creazione delle password per tutti gli utenti, assicura che le stesse siano forti, ma tale approccio potrebbe assumere aspetti negativi con la crescita dell'organizzazione. Tale metodo potrebbe anche aumentare il rischio che gli utenti possano scrivere le proprie password su di un foglietto.

Per queste ragioni, molti amministratori preferiscono autorizzare i propri utenti alla creazione delle proprie password. Tuttavia, un buon amministratore verifica sempre che le stesse siano sufficientemente forti.

Per una guida sulla creazione delle password, consultare il capitolo *Workstation Security* nella *Red Hat Enterprise Linux Security Guide*.

La necessità da parte di un amministratore di mantenere segrete le password, deve essere un elemento indispensabile durante la gestione dei sistemi. Tuttavia, questo punto non viene percepito da molti utenti, infatti essi non riescono a capire la differenza che intercorre tra il nome utente e le password. Constatato ciò, è vitale educare gli utenti, in modo che gli stessi capiscano che le proprie password devono essere mantenute segrete.

Le password devono essere difficili da indovinare, Una password forte è quella che risulta essere impossibile da indovinare, anche se l'aggressore conosce molto bene l'utente.

Un attacco tipo 'brute-force' ad un password comporta un tentativo metodico (generalmente tramite un programma chiamato *password-cracker*) di ogni combinazione possibile di caratteri, nella speranza che la corretta password venga trovata. Una password forte deve essere creata in modo tale che l'aggressore sia obbligato a trascorrere molto tempo prima di poter indovinare quella corretta.

Sia le password forti che quelle deboli verranno affrontate in dettaglio nelle sezioni seguenti.

6.1.2.1. Password deboli

Come precedentemente accennato, una password debole non rispetcia le seguenti regole:

- È segreta
- È difficile da indovinare
- È resistente agli attacchi che usano 'forza-bruta'

Le seguenti sezioni mostrano come una password possa essere debole.

6.1.2.1.1. Password corte

Una password corta è debole in quanto potrebbe essere più facilmente soggetta ad attacchi che usano 'forza-bruta'. Per illustrare quanto detto, considerate la seguente tabella, dove il numero delle potenziali password da provare in un attacco è conosciuto. (Le password in questione hanno tutte lettere minuscole.)

Lunghezza della password	Password potenziali
1	26
2	676
3	17,576
4	456,976
5	11,881,376
6	308,915,776

Tabella 6-1. Lunghezza della password contro il numero delle potenziali password

Come potete vedere, il numero delle possibili password aumenta sensibilmente in accordo con la lunghezza.



Nota bene

Anche se questa tabella termina con sei caratteri, ciò non significa che le password a sei caratteri sono sufficientemente lunghe ed idonee per la sicurezza del sistema. In generale, più la password è lunga meglio è.

6.1.2.1.2. Set di caratteri limitati

Il numero di diversi caratteri che costituisce una password, ha un impatto sostanziale nel modo di condurre un attacco con l'uso di forza bruta da parte di un aggressore. Per esempio, invece di usare solo 26 caratteri, nel caso in cui usassimo solo caratteri minuscoli, che ne dite di usare anche i nu-

meri? Questo significherebbe che ogni carattere all'interno di una password, potrebbe essere uno di 36 caratteri invece di 26. Nel caso di una password a sei caratteri, questo aumenterebbe il numero di password possibili da 308,915,776 a 2,176,782,336.

C'è ancora molto che può essere fatto. Se includessimo anche un mix di caratteri alfanumerici con lettere minuscole e maiuscole (per sistemi operativi che lo supportano), il numero di possibili combinazioni per una password aumenterebbe a 56,800,235,584. Aggiungendo altri caratteri (come ad esempio la punteggiatura), tale combinazione aumenterebbe sensibilmente, rendendo un attacco con forza bruta ancora più difficoltoso.

Tuttavia è da ricordare che non tutti gli attacchi usano forza bruta. Le sezioni seguenti descrivono altre cause che d'anno origine a password deboli.

6.1.2.1.3. Parole comuni

La maggior parte degli attacchi alle password si basano sul fatto che molte persone usano parole semplici da ricordare. E per molte persone, le password facili da ricordare contengono parole comuni. Per questo motivo, molti attacchi vengono eseguiti usando parole presenti nel dizionario. In altre parole, l'aggressore usa le suddette parole per riuscire ad indovinare la parola chiave che compone una password.



Nota bene

Molti programmi d'attacco basati sulle parole presenti nel dizionario, usano diverse lingue. Per questo motivo, non pensate di avere una password forte solo perchè usate una parola di lingua diversa dalla vostra.

6.1.2.1.4. Informazioni personali

Le password che contengono informazioni personali (ad esempio il nome o la data di nascita di una persona cara, il nome di un animale, o il numero d'identificazione personale) potrebbe essere individuato da un attacco basato con una parola presente nel dizionario. Tuttavia se l'aggressore vi conosce personalmente (oppure se è abbastanza motivato da ricercare nella vostra vita privata), potrebbe essere in grado di indovinare la vostra password con poco o nessuno sforzo.

In aggiunta alle parole presenti nel dizionario, molti aggressori potrebbero usare anche nomi comuni, date e altre informazioni di questo genere. Per questo motivo, anche se l'aggressore non conosce il nome del vostro cane, Gracie, egli potrebbe essere in grado di sapere che la vostra password è "mydogisgracie", utilizzando un buon programma di attacco della password.

6.1.2.1.5. Trucchi semplici

Usando una qualsiasi delle informazioni precedentemente affrontate come base per la vostra password, e invertendo l'ordine dei caratteri, non sarete comunque in grado di tramutare una password debole in una forte. Molti programmi password-cracker sono in grado di eseguire molti trucchi per ottenere le possibili password. Essi possono includere la sostituzione di numeri con lettere ecc. Ecco alcuni esempi:

- drowssaPdaB1
- R3allyP00r

6.1.2.1.6. *La stessa password per sistemi multipli*

Anche se possedete una password forte, non è consigliabile usare la stessa password per più di un sistema. Ovviamente non si può fare diversamente se si è in presenza di sistemi configurati in modo tale da usare un server centrale di autenticazione, ma in qualsiasi altro caso, si consiglia vivamente l'uso di una password diversa per ogni sistema.

6.1.2.1.7. *Trascrizione della password su foglietti*

Un altro modo per trasformare una password forte in una debole, è quello di trascriverla su di un pezzo di carta. Trascrivendola su carta, annullerete la segretezza della vostra password, in cambio avrete così un problema di sicurezza 'fisica' — ciò significa che dovrete fare in modo di mantenere quel pezzettino di carta al sicuro. Per questo motivo, trascrivere la vostra password su di un pezzettino di carta, non è consigliabile.

Tuttavia, alcune organizzazioni hanno la necessità di avere delle password scritte. Per esempio, alcune organizzazioni adottano una password scritta come procedura di ripristino dovuta alla perdita di personale importante (ad esempio un amministratore del sistema). In questi casi, il pezzettino di carta contenente le password, è conservato in un luogo fisico sicuro che necessita della cooperazione di diverse persone per ottenere il suo accesso. Per questo motivo vengono spesso usate piccole casseforti con serrature multiple, o cassette bancarie di sicurezza.

Qualsiasi organizzazione che adotta questo tipo di sistema, tenga sempre in considerazione che l'esistenza di una password scritta presenta sempre un certo rischio, anche se la stessa viene conservata in un luogo molto sicuro. Ciò aumenterebbe il rischio se si è a conoscenza che la password è stata trascritta (ed il luogo nella quale essa viene conservata).

Sfortunatamente le password trascritte non fanno parte di una procedura di recupero e non sono conservate in casseforti, ma esse sono password di utenti normali, e vengono conservate nei seguenti luoghi:

- In un cassetto (sotto chiave o senza alcun tipo di protezione)
- Sotto una tastiera
- All'interno di un portafoglio
- Attaccata ai lati di un monitor

Nessuno dei luoghi indicati è idoneo a conservare una password.

6.1.2.2. Password forti

Precedentemente abbiamo descritto le password deboli, le sezioni seguenti descriveranno tutte le caratteristiche di una password forte.

6.1.2.2.1. *Password più lunghe*

Più lunghe sono le password, e maggiore sarà la possibilità che un attacco del tipo 'forza-bruta' non abbia successo. Per questo motivo, se il vostro sistema operativo lo supporta, impostate per i vostri utenti, una password relativamente lunga.

6.1.2.2.2. *Set esteso di caratteri*

Incoraggiare l'uso di lettere minuscole e maiuscole, password alfanumeriche, e l'aggiunta di almeno un carattere non alfanumerico:

- t1Te-Bf,te
- Lb@lbhom

6.1.2.2.3. Da ricordare

Una password è forte solo se si può ricordare. Tuttavia facile da ricordare e facile da indovinare coincidono molto spesso tra loro. Per questo motivo, divulgate ai vostri utenti alcuni suggerimenti per la creazione di password facili da ricordare ma difficili da indovinare.

Per esempio, prendete un detto o una frase, e usate le prime lettere di ogni parola come punto di partenza per la creazione della nuova password. Il risultato sarà facile da ricordare (in quanto la frase o il detto in sè per sè sono semplici da ricordare), ed il risultato non conterrà alcuna parola comune.



Nota bene

Tenete presente che non basta usare le prime lettere di ogni parola per ottenere una password sufficientemente sicura. Assicuratevi di aumentare sempre il numero dei caratteri che compongono la password stessa, usando caratteri alfanumerici, lettere minuscole e maiuscole e almeno un carattere speciale.

6.1.2.3. Invecchiamento della password

Se possibile, implementate il metodo d'invecchiamento della password. Tale metodo è un contenuto (disponibile su molti sistemi operativi), che imposta i limiti entro i quali una password viene considerata valida. Alla fine del ciclo di validità di una password, verrà richiesto all'utente di inserire una nuova password, che potrà essere usata fino a quando la stessa non raggiunge il tempo limite impostato.

La domanda chiave alla quale molti amministratori devono trovare risposta, è quella di dare un tempo di validità alla password. Ma quanto tempo deve durare la password?

Ci sono due problematiche diametralmente opposte per quanto riguarda la durata della validità della password:

- Comodità per l'utente
- Sicurezza

Una password con una durata di 99 anni rappresenterebbe, in modo minimo o nullo, alcun problema per l'utente. Tuttavia, non sarà in grado di fornire alcun miglioramento della sicurezza.

D'altro canto, una durata di 99 minuti potrebbe presentare delle inconvenienze da parte degli utenti. Tuttavia, la sicurezza sarebbe migliorata sensibilmente.

L'idea è quella di trovare un equilibrio tra l'agevolazione dei vostri utenti, e le necessità inerenti la sicurezza della vostra organizzazione. Per molte organizzazioni, la durata delle password, tradotte in settimane e mesi, rappresentano le prassi più comuni.

6.1.3. Informazioni per il controllo dell'accesso

Insieme con il nome utente e la password, gli user account contengono anche le informazioni sul controllo dell'accesso. Queste informazioni prendono forma a seconda dei sistemi operativi in uso. Tuttavia, le informazioni spesso includono:

- L'identificazione specifica dell'utente nel sistema
- L'identificazione specifica del gruppo nel sistema
- Elenchi di gruppi aggiuntivi/capacità alle quali è membro l'utente
- Informazioni d'accesso di default da applicare a tutti i file creati dagli utenti e alle risorse

In alcune organizzazioni, non vi è la necessità di 'toccare' le informazioni per il controllo dell'accesso dell'utente. Questo è solitamente il caso con standalone, e personal workstation. Altre organizzazioni, in particolare quelle che fanno uso continuo di risorse della rete e condividono con gruppi e utenti diversi, richiedono invece una modifica continua delle suddette informazioni.

Il carico di lavoro necessario per gestire in modo corretto le informazioni per il controllo dell'accesso del vostro utente, variano a seconda dell'utilizzo fatto dalla vostra organizzazione dei suddetti contenuti. Anche se non è una prassi scorretta affidarsi così fortemente a questi contenuti (infatti, potrebbe risultare inevitabile), ciò potrebbe significare che l'ambiente del vostro sistema, potrebbe necessitare di più attenzione nella gestione, in aggiunta, ogni user account potrebbe presentare molteplici modi per essere configurato in modo non corretto.

Per questo motivo, se la vostra organizzazione necessita di questo tipo di ambiente, è consigliato creare una guida dove vengono riportate le varie fasi necessarie per creare e configurare correttamente un user account. Infatti, se sono presenti diversi tipi di user account, dovrete documentarli tutti (creando un nuovo finance user account, un nuovo operations user account, ecc.).

6.1.4. Gestione degli account e dell'accesso alle risorse giorno per giorno

Un vecchio detto dice, l'unica cosa costante è il cambiamento. In questo contesto non vi è alcuna differenza con gli utenti della vostra community. Le persone vanno e vengono, molte cambiano la loro area di responsabilità. Così, gli amministratori, devono essere in grado di far fronte ai cambiamenti, in quanto essi fanno parte della vita giornaliera della vostra organizzazione.

6.1.4.1. Nuovi dipendenti

Quando viene assunta una nuova persona, gli si dà la possibilità di accedere a diverse risorse (a seconda della loro carica). Verranno resi disponibili una scrivania, un telefono e la chiave d'ingresso.

Potrebbero anche aver accesso ad uno o più computer. Come amministratore di sistema, è vostra responsabilità che ciò venga fatto in modo responsabile e corretto. Quindi come farlo in modo corretto?

Prima di tutto dovete essere a conoscenza dell'arrivo di una nuova persona. Questo è gestito in modo differente a seconda dell'organizzazione. Ecco alcuni esempi:

- Attuate una procedura dove ogni qualvolta si assume una nuova persona, voi verrete subito informati.
- Create un modulo da far completare dal supervisore del nuovo impiegato con la richiesta di un nuovo account.

A seconda dell'organizzazione vi sono approcci diversi. Tuttavia è necessario farlo, è importante avere un processo molto affidabile, in grado di informarvi sugli interventi da apportare nei confronti degli account dei vostri utenti.

6.1.4.2. Termine di un rapporto lavorativo

È normale che alcune persone decidano di terminare il loro rapporto lavorativo con la vostra organizzazione. Talvolta si può terminare questo rapporto in modo consenziente e talvolta in modo un pò più brusco. In entrambi i casi, è importante essere informati di questi eventi, per poter apportare da parte vostra le azioni più appropriate.

È importante che le suddette azioni includano:

- Disabilitare la possibilità d'accesso da parte dell'utente, a tutti i sistemi e le relative risorse (generalmente cambiando/invalidando la password dell'utente stesso)
- Eseguire il back up di tutti i file, nel caso si presenti la necessità di usarli in futuro
- Coordinare l'accesso ai file dell'utente da parte del suo manager

La cosa più importante è quella di rendere sicuri i vostri sistemi nei confronti del personale che ha terminato un rapporto lavorativo. Ciò è particolarmente importante se vi è stato un licenziamento e che tale azione abbia scaturito nella persona licenziata, uno stato d'animo di rancore nei confronti della vostra organizzazione. Tuttavia, anche se la persona abbia volontariamente deciso di interrompere il proprio rapporto con l'organizzazione, rientra negli interessi dell'organizzazione stessa disabilitare la possibilità di accesso ai sistemi e alle risorse.

Questo indica la necessità di essere informati ogni volta che si verifica una cessazione nel rapporto lavorativo — e possibilmente anche prima che questo accada. Ciò implicherà un maggiore coinvolgimento da parte vostra con il personale della vostra organizzazione, in modo tale da essere avvertiti in tempi molto brevi.



Suggerimento

Durante la gestione di un "lock-down" del sistema come risposta ad una cessazione del rapporto lavorativo, tenete presente che la tempestività delle vostre azioni è molto importante. Se il lock down si manifesta dopo che il processo di terminazione è stato completato, si potrebbe verificare un accesso non autorizzato. Tuttavia se il suddetto lock down si manifesta prima che il processo sia stato avviato, il personale in questione potrebbe essere informato, rendendo difficile così la prassi di terminazione del rapporto lavorativo.

Tale cessazione ha inizio con un meeting tra la persona interessata, il suo manager, e una rappresentanza della vostra sezione. Tuttavia, la presenza di una procedura che vi metta a conoscenza dell'inizio di questo procedimento è molto importante, in quanto fa sì che il periodo di lock down sia appropriato.

Una volta disabilitato l'accesso, è necessario eseguire un back up dei file del personale in questione. Tale procedura potrebbe far parte di uno standard seguito dalla vostra organizzazione, il modo con il quale eseguire il suddetto back up viene regolato da leggi precise che garantiscono la privacy.

Ad ogni modo, eseguire un back up è una pratica molto valida in quanto si potrebbero cancellare (accesso da parte del manager ai file della persona 'terminata') accidentalmente file importanti. In queste circostanze, avere un back up aggiornato, rende possibile far fronte alle problematiche della suddetta cancellazione, facilitando così il processo.

A questo punto dovete determinare il tipo di accesso ai file da conferire al manager del personale appena 'terminato'. A seconda dell'organizzazione e della natura dei compiti del suddetto personale, è possibile che non sia necessario conferire alcun tipo di accesso, oppure permettere un accesso completo.

Se il personale era in grado di utilizzare i vostri sistemi per inviare non solo le email riguardanti gli incidenti verificatisi, ma anche email più complesse, allora è necessario che anche il manager in questione sia a conoscenza del contenuto delle email, e determinare quindi le email da cancellare e quelle da conservare. Alla fine di questo processo è possibile inoltrare ai sostituti, i file posseduti dal

personale 'terminato'. Potrebbe essere necessaria la vostra presenza durante l'esecuzione di questo processo, tutto dipende dalla natura e dal contenuto dei file che la vostra organizzazione deve gestire.

6.1.4.3. Modifiche degli incarichi

La creazione degli account per nuovi utenti e la gestione delle fasi necessarie per eseguire un lock-down dell'account di un utente 'terminato', sono procedimenti molto semplici. Tuttavia il suddetto procedimento potrebbe rappresentare qualche ostacolo se il personale assume un incarico diverso pur restando all'interno della stessa organizzazione. Talvolta il personale può richiedere la variazione del proprio account.

Per accertarsi che il nuovo account sia riconfigurato in modo corretto, è necessaria la presenza di almeno tre persone:

- Voi
- Il precedente manager dell'utente in questione
- Il nuovo manager

Insieme, dovrete essere in grado di determinare le azioni per terminare in modo corretto le responsabilità precedenti, e intraprendere le azioni appropriate per creare un nuovo account con nuovi incarichi. In molti casi questo processo potrebbe essere simile alla chiusura di un vecchio account e alla creazione di uno nuovo. Molte organizzazioni adottano quest'ultima procedura.

Tuttavia è molto probabile che l'account in questione venga modificato per adattarsi alle nuove responsabilità. Questo tipo di approccio richiede molta attenzione da parte vostra nel rivedere l'account e assicurarvi che lo stesso abbia nuove risorse e privilegi idonei ai nuovi incarichi.

Per complicare la situazione si potrebbe verificare l'opportunità che una persona sia chiamata a ricoprire due incarichi contemporaneamente. In questo caso il precedente manager insieme con quello nuovo, potranno aiutarvi garantendovi un periodo di tempo per far fronte a questa eventualità.

6.2. Gestione delle risorse dell'utente

La creazione di nuovi account fa parte del lavoro di un amministratore di sistema. La gestione delle risorse di un utente rappresenta un compito essenziale. Per questo motivo, è necessario considerare tre punti:

- Chi è in grado di accedere ai dati condivisi
- Da dove si possono accedere questi dati
- I limiti presenti per prevenire l'abuso di queste risorse

Le seguenti sezioni affrontano in modo abbreviato questi aspetti.

6.2.1. Personale in grado di accedere ai dati condivisi

L'accesso da parte di un utente di applicazioni, file o directory è determinato dai permessi in vigore.

In aggiunta, risulta essere utile se vengono applicati permessi diversi a seconda delle diverse classi di utenti. Per esempio, un 'temporary storage' condiviso dovrebbe essere in grado di prevenire la cancellazione accidentale (o maliziosa) dei file di un utente, da parte di altri utenti, permettendo un accesso completo dei file da parte del proprietario.

Un altro esempio è l'accesso assegnato alla home directory di un utente. Solo il proprietario della home directory dovrebbe essere in grado di creare o visualizzare i file. Gli altri utenti non dovrebbero essere in grado di ottenere l'accesso (se non diversamente richiesto dall'utente). Questa procedura aumenta la privacy dell'utente e previene la possibilità di entrare in possesso di file personali.

Vi sono molteplici situazioni dove utenti multipli potrebbero aver bisogno di accedere alle stesse risorse presenti sulla macchina. In questo caso potrebbe essere necessario creare, con molta cautela, gruppi condivisi.

6.2.1.1. Gruppi condivisi e dati

Come accennato nella parte introduttiva, i gruppi sono delle entità logiche che possono essere usati per raggruppare gli user account per determinati compiti.

Nella gestione degli utenti all'interno di una organizzazione, è sempre utile identificare i dati necessari a determinate sezioni, quali invece possono essere resi non disponibili e quali possono essere condivisi da tutti. Implementare questa struttura durante la creazione di un gruppo, insieme ai permessi appropriati per i dati condivisi, rappresenta una fase molto importante.

Per esempio, supponiamo che vi siano due account, uno rappresenta la sezione dei Crediti, che deve mantenere un elenco degli account che hanno pagamenti in sospeso. Tale sezione deve condividere lo stesso elenco con la sezione responsabile alle Riscossioni dei pagamenti. Se il personale di entrambi gli account viene inserito in un gruppo chiamato *accounts*, le informazioni possono essere posizionate all'interno di una directory condivisa (dove il gruppo *accounts* risulta essere proprietario), con il permesso di scrittura e lettura della directory per il gruppo stesso.

6.2.1.2. Determinazione della struttura del gruppo

Alcune delle sfide che si possono presentare agli amministratori durante la creazione di gruppi condivisi sono:

- I gruppi da creare
- Chi inserire nei gruppi
- Il tipo di permessi che devono avere le risorse condivise

Un approccio ragionevole potrebbe essere molto utile. Una possibilità potrebbe essere quella di riflettere la struttura della vostra organizzazione durante la creazione dei gruppi. Per esempio, si vi è una sezione responsabile alle finanze, allora potreste creare un gruppo chiamato *finance*, ed inserire tutto il personale della suddetta sezione all'interno di questo gruppo. Se le informazioni hanno un livello molto riservato, allora potete garantire al personale dirigenziale un permesso del tipo 'group-level', in modo da accedere alle directory e dati, inserendo il solo personale dirigenziale nel gruppo *finance*.

È consigliabile fare molta attenzione nel garantire i permessi agli utenti. In questo modo la possibilità che le informazioni sensibili possano cadere in mani sbagliate risulta essere minore.

Seguendo questo consiglio durante la creazione della struttura del gruppo della vostra organizzazione, è possibile far fronte alla necessità di accedere i dati condivisi in modo sicuro ed efficace.

6.2.2. Da dove si possono accedere i dati condivisi

Quando si condividono i dati tra utenti è normale usare un server centrale (o un gruppo di server) che rendono alcune directory disponibili ad altre macchine presenti nella rete. In questo modo i dati vengono conservati in un unico luogo; non è necessario sincronizzare i dati tra macchine multiple.

Prima di intraprendere questo approccio, dovete determinare i sistemi che possono accedere i dati conservati in modo centralizzato. Nel fare questo, prendete nota dei sistemi operativi usati dai sistemi.

Queste informazioni hanno attinenza sulla vostra abilità di implementare tale approccio, in quanto il vostro storage server deve essere in grado di servire i propri dati ad ogni sistema operativo in uso nella vostra organizzazione.

Sfortunatamente, quando i dati vengono condivisi tra diversi computer presenti su di una rete, è possibile che si verifichi un conflitto nella ownership dei file.

6.2.2.1. Problematiche inerenti la ownership globale

Vi sono diversi benefici nel conservare i dati in uno storage centralizzato, rendendo così possibile l'accesso degli stessi ai diversi computer presenti sulla rete. Tuttavia, immaginate per un momento che ogni computer possiede un elenco di user account gestito in modo locale. Cosa potrebbe accadere se gli elenchi degli utenti presenti sui sistemi, non coincidono con gli elenchi presenti sul server centrale? Oppure, se i suddetti elenchi presenti su ogni sistema non coincidono tra loro?

Tutto questo dipende soprattutto da come vengono implementati su ogni sistema i permessi per gli accessi e gli utenti, ma in alcuni casi è possibile che l'utente A su di un sistema, potrebbe essere conosciuto come utente B su di un altro sistema. Questo potrebbe rappresentare un problema quando si condividono dati tra i sistemi in questione, l'utente potrebbe accedere ai dati su entrambi i sistemi sia come utente A che come utente B.

Per questa ragione, molte organizzazioni usano una specie di database centrale. Ciò assicura che non si verifichino sovrapposizioni tra elenchi degli utenti presenti sui vari sistemi.

6.2.2.2. Home Directory

Un'altro problema che un amministratore di sistema può affrontare, è quello di decidere se gli utenti possono avere home directory centralizzate.

Il vantaggio primario nel centralizzare le home directory su di un server collegato ad una rete, è rappresentato dal fatto che se un utente esegue un log in su di una macchina, egli è saràn grado di accedere ai file nella propria home directory.

Lo svantaggio invece è rappresentato dal fatto che se la rete subisce una interruzione, gli utenti presenti nell'organizzazione saranno impossibilitati ad accedere i prori file. In alcune situazioni (come quella che si potrebbe verificare in una organizzazione che fa ampio uso di computer portatili), avere delle home directory centralizzate potrebbe non rappresentare la migliore soluzione. Tale soluzione però potrebbe risultare la migliore alternativa, per questo motivo implementatela, anche perchè essa-semplificherebbe molto il vostro lavoro.

6.2.3. Limiti implementati per prevenire l'abuso delle risorse

L'accurata organizzazione dei gruppi e l'assegnazione dei permessi per le risorse condivise, rappresenta una delle fasi molto delicate che un amministratore di sistema deve affrontare, in modo da prevenire l'abuso delle risorse da parte degli utenti all'interno di una organizzazione. Così facendo, gli utenti non autorizzati ad accedere informazioni sensibili, verrà negato l'accesso.

Non importa come si comporta la vostra organizzazione, la migliore difesa è rappresentata dalla continua vigilanza del sistema da parte dell'amministratore di sistema. Mantenere alto il controllo vi eviterà spiacevoli sorprese.

6.3. Informazioni specifiche su Red Hat Enterprise Linux

Le seguenti sezioni descrivono i diversi contenuti specifici di Red Hat Enterprise Linux in relazione all'amministrazione degli user account e delle risorse associate.

6.3.1. User Account, Gruppi, e Permessi

Con Red Hat Enterprise Linux, un utente è in grado di eseguire il log in nel sistema e usare qualsiasi applicazione o file, ai quali essi sono abilitati per l'accesso, dopo la creazione di un normale user account. Red Hat Enterprise Linux determina se un utente o un gruppo può accedere a queste risorse in base ai loro permessi.

Sono presenti tre diversi tipi di permessi, essi si riferiscono ai file, alle directory e alle applicazioni. Questi permessi sono usati per un controllo dei diversi tipi di accesso. Per descrivere ogni permesso in un elenco della directory, vengono utilizzati diversi caratteri. Ecco riportati i seguenti esempi:

- `r` — Indica una categoria di utenti in grado di leggere un file.
- `w` — Indica una categoria di utenti in grado di scrivere in un file.
- `x` — Indica una categoria di utenti in grado di eseguire un file.

Un quarto simbolo (`-`) indica che non è consentito alcun accesso.

Ogni permesso viene assegnato a tre diverse categorie di utenti. Le categorie sono:

- *owner* — Il proprietario del file o dell'applicazione
- *group* — Il gruppo che possiede il file o l'applicazione
- *everyone* — Tutti gli utenti che hanno accesso al sistema.

Come precedentemente affrontato, è possibile visualizzare i permessi per un file, richiamando un elenco completo tramite il comando `ls -l`. Per esempio, se l'utente `juan` crea un file eseguibile chiamato `foo`, l'output del comando `ls -l foo` sarà il seguente:

```
-rwxrwxr-x  1 juan    juan          0 Sep 26 12:25 foo
```

I permessi per questo file verranno elencati all'inizio della riga, iniziando con `rwx`. Questo insieme di simboli definisce l'accesso del proprietario — in questo esempio, il proprietario `juan` detiene un accesso completo, ed è in grado di leggere, scrivere ed eseguire il file. Il set successivo `rwx`, definisce l'accesso del gruppo (viene riportato ancora un accesso completo), mentre l'ultimo set di simboli definisce il tipo di accesso autorizzato per tutti gli altri utenti. Qui, tutti gli altri utenti possono leggere ed eseguire il file, ma potrebbero non essere in grado di apportare alcuna modifica.

Una cosa da ricordare per quanto riguarda i permessi e gli user account, è che ogni applicazione eseguita su Red Hat Enterprise Linux viene eseguita nel contesto di un utente specifico. Generalmente, questo significa che se l'utente `juan` lancia un'applicazione, la stessa viene eseguita usando il contesto dell'utente `juan`. Tuttavia in alcuni casi l'applicazione potrebbe aver bisogno di un livello d'accesso più privilegiato per poter portare a termine un compito. Tali applicazioni includono anche le applicazioni in grado di modificare le impostazioni del sistema oppure il log in dell'utente. Per questa ragione, sono stati creati permessi speciali.

Sono presenti tre permessi speciali all'interno di Red Hat Enterprise Linux. Essi sono:

- *setuid* — usato solo per le applicazioni, questo permesso indica che l'applicazione deve essere eseguita come proprietario del file, e non come utente che esegue l'applicazione. Questo viene indicato dal carattere `s` al posto di `x` nella categoria del proprietario. Se il proprietario del file non possiede i permessi di esecuzione, comparirà il carattere `s`.
- *setgid* — usato principalmente per le applicazioni, questo permesso indica che l'applicazione deve essere eseguita come il gruppo proprietario del file e non come il gruppo dell'utente che esegue l'applicazione.

Se applicato ad una directory, tutti i file creati all'interno della directory sono posseduti dal gruppo che detiene la directory e non dal gruppo dell'utente che crea il file. Il permesso `setgid` viene indi-

cato dal carattere `s` al posto di `x` nella categoria del gruppo. Se il gruppo proprietario del file o della directory non possiede i permessi di esecuzione, verrà visualizzato il carattere `s`.

- *sticky bit* — usato principalmente sulle directory, la parola bit indica che il file creato nella directory può essere rimosso solo dall'utente che ha creato il file. Viene indicato dal carattere `t` al posto di `x` in ogni categoria. Se ogni categoria non possiede i permessi di esecuzione, verrà visualizzata la lettera `T`.

Con Red Hat Enterprise Linux, sticky bit viene impostato per default sulla directory `/tmp/` per questa ragione.

6.3.1.1. Nome utente e UID, Gruppi e GID

Con Red Hat Enterprise Linux, l'user account ed i group name sono principalmente per comodità delle persone. Internamente, il sistema usa identificatori numerici. Per gli utenti, questo identificatore è conosciuto come *UID*, mentre per i gruppi è *GID*. I programmi che rendono disponibili le informazioni inerenti l'utente o il gruppo sono in grado di tradurre i valori di UID/GID in valori facilmente leggibili dalle persone.



Importante

Se desiderate condividere file e risorse tramite la rete all'interno della vostra organizzazione, UID e GID devono essere unici. Altrimenti nonostante implementiate qualsiasi tipo di controllo per l'accesso, esso potrebbe non funzionare correttamente in quanto il suddetto controllo è basato sugli UID e GID, e non sul nome utente o quello dei gruppi.

In modo più specifico, se i file `/etc/passwd` e `/etc/group` su di un file server ed una workstation presentano UID e GID diversi, i permessi inerenti l'applicazione potrebbero compromettere la sicurezza.

Per esempio, se l'utente `juan` ha un UID pari a 500 su di un computer desktop, i file creati da `juan` su di un file server, saranno creati con un UID di 500. Tuttavia, se l'utente `bob` esegue un log in in modo locale al file server (o su di un altro computer), e l'account di `bob` possiede un UID di 500, `bob` otterrà un accesso completo ai file di `juan`, e viceversa.

Per questo motivo è sempre consigliabile evitare i conflitti tra UID e GID.

Ci possono essere due casi dove il valore numerico di un UID e GID hanno un significato specifico. Un UID e GID pari a zero (0) vengono generalmente usati per l'utente `root`, e sono trattati in modo particolare da Red Hat Enterprise Linux —, l'accesso completo viene garantito in modo automatico.

Il secondo è quello dove tutti gli UID e GID hanno un valore inferiore a 500, e sono riservati per l'uso del sistema. A differenza di UID/GID con il valore zero, gli UID e GID sotto i 500 non vengono trattati in modo particolare da Red Hat Enterprise Linux. Tuttavia, questi UID/GID non devono mai essere assegnati ad un utente, in quanto potrebbe essere probabile che un componente del sistema usa o userà questi UID/GID. Per maggiori informazioni su questi utenti e gruppi standard, controllare il capitolo *Utenti e Gruppi* nella *Red Hat Enterprise Linux Reference Guide*.

Quando vengono aggiunti i nuovi user account usando i tool standard di Red Hat Enterprise Linux per la creazione degli utenti, vengono assegnati ai nuovi user account gli UID e GID disponibili il cui valore inizia da 500. All'user account successivo viene assegnato un UID/GID pari a 501, seguito da UID/GID 502, e così via.

Più avanti nel capitolo verrà affrontato brevemente l'argomento inerente i vari tool disponibili con Red Hat Enterprise Linux, per la creazione dell'utente. Ma prima di affrontare tale argomento, la sezione successiva affronta i file usati da Red Hat Enterprise Linux per definire gli account ed i gruppi del sistema.

6.3.2. File che controllano gli User Account ed i Gruppi

Su Red Hat Enterprise Linux, le informazioni sugli user account e gruppi sono conservate in diversi file di testo all'interno della directory `/etc/`. Quando un amministratore di sistema crea nuovi user account, questi file devono essere modificati manualmente oppure è necessario usare alcune applicazioni per eseguire i suddetti cambiamenti.

La seguente sezione riporta i file presenti nella directory `/etc/`, che conservano le informazioni sugli utenti e sui gruppi presenti su Red Hat Enterprise Linux.

6.3.2.1. `/etc/passwd`

Il file `/etc/passwd` può essere letto da tutti e contiene un elenco di utenti, ognuno dei quali riportato in righe separate. Ogni riga rappresenta un elenco delimitato dai due punti e contiene le seguenti informazioni:

- *Nome utente* — Il nome che l'utente digita quando esegue il log in nel sistema.
- *Password* — Contiene la password cifrata (o una serie di `x` se viene usata la password shadow — più avanti tratteremo questo argomento in modo più approfondito).
- *User ID (UID)* — L'equivalente numerico del nome utente riferito dal sistema e dalle applicazioni quando si determinano i privilegi d'accesso.
- *Group ID (GID)* — L'equivalente numerico del gruppo primario riferito dal sistema e dalle applicazioni quando si determinano i privilegi d'accesso.
- *GECOS* — Nominato per ragioni storiche, il campo `GECOS`¹ è facoltativo e viene usato per conservare informazioni supplementari (come ad esempio il nome completo dell'utente). Entry multiple possono essere conservate in un elenco delimitato da un virgola. Le utility come ad esempio `finger`, possono accedere questo campo per fornire informazioni aggiuntive sull'utente.
- *Home directory* — Path assoluta per la home directory dell'utente, come ad esempio `/home/juan/`.
- *Shell* — Il programma viene lanciato automaticamente ogni qualvolta un utente esegue il log in. Questo è generalmente chiamato command interpreter (oppure *shell*). Con Red Hat Enterprise Linux, il valore di default è `/bin/bash`. Se invece viene lasciato vuoto, viene usato `/bin/sh`. Se viene impostato su di un file non esistente, allora l'utente non sarà in grado di eseguire il log in nel sistema.

Ecco un esempio di una entry `/etc/passwd`:

```
root:x:0:0:root:/root:/bin/bash
```

Questa riga mostra che l'utente `root` possiede una password shadow, e una UID e GID pari a 0. L'utente `root` ha `/root/` come home directory, e usa `/bin/bash` come shell.

Per maggiori informazioni su `/etc/passwd`, consultare la pagina `man di passwd(5)`.

1. `GECOS` è l'acronimo di General Electric Comprehensive Operating Supervisor. Questo campo è stato usato da Bell Labs, nell'implementazione originale di UNIX. Il lab presentava diversi computer, incluso un `GECOS` funzionante. Questo campo è stato usato per conservare informazioni quando il sistema UNIX generalmente inviava batch e lavori di stampa al sistema `GECOS`.

6.3.2.2. `/etc/shadow`

A causa della necessità che il file `/etc/passwd` sia letto da tutti (la ragione principale è che questo file viene usato per eseguire un cambiamento da UID a nome utente), potrebbe essere rischioso conservare le password in `/etc/passwd`. È vero le password sono cifrate. Tuttavia è possibile eseguire attacchi contro le password se le stesse sono disponibili.

Se un aggressore è in grado di ottenere una copia di `/etc/passwd`, è possibile che egli stesso possa eseguire un attacco in tutta segretezza. Invece di rischiare di essere scoperto usando la password fornita dal password-cracker, l'aggressore può seguire le seguenti fasi:

- Un attacco con l'utilizzo del password-cracker è in grado di generare password potenziali
- Ogni password viene cifrata usando lo stesso algoritmo del sistema
- La password viene paragonata con le password cifrate in `/etc/passwd`

L'aspetto più pericoloso di questo attacco è rappresentato dal fatto che esso può essere portato a termine su di un sistema esterno alla vostra organizzazione. Per questo motivo, l'aggressore è in grado di usare un hardware disponibile ad altissime-prestazioni, in modo tale da analizzare un numero elevatissimo di parole molto velocemente.

Per questo motivo, il file `/etc/shadow` è leggibile solo da parte dell'utente root e contiene la password (e informazioni aggiuntive sull'invecchiamento della password) per ogni utente. Come riportato nel file `/etc/passwd`, le informazioni di ogni utente sono riportate su righe separate. Ogni riga rappresenta un elenco delimitato da due punti, e racchiude le seguenti informazioni:

- *Nome utente* — Il nome digitato dall'utente quando esegue il log in nel sistema. Questo abilita l'applicazione **login** a recuperare la password dell'utente (e relative informazioni).
- *Password cifrata* — Password da 13 a 24 caratteri. La password viene cifrata usando la funzione della libreria `crypt(3)` o l'algoritmo md5 hash. In questo campo, i valori diversi da quelli cifrati e formattati in modo valido o password di tipo criptata o 'hashed', vengono usati per controllare i log in degli utenti e per mostrare lo stato della password. Per esempio, se il valore è `! o *`, l'account viene bloccato e di conseguenza l'utente non è autorizzato ad eseguire il log in. Se il valore è `!!` ciò sta ad indicare che non è mai stata impostata una password (e l'utente, non avendo impostato una password, non sarà in grado di eseguire il log in).
- *Ultimo cambiamento della password* — Il numero di giorni dall'1 Gennaio 1970 (anche chiamato *epoch*) che la password è stata modificata. Questa informazione viene usata insieme ai campi seguenti d'invecchiamento della password.
- *Numero di giorni prima dei quali è possibile modificare una password* — Numero minimo di giorni da trascorrere prima che la password possa essere cambiata.
- *Numero di giorni prima che sia necessario modificare la password* — Numero di giorni da trascorrere necessari prima di modificare la password.
- *Numero di giorni di preavviso prima che la password venga cambiata* — Il numero di giorni prima della scadenza della password, durante i quali l'utente viene avvisato.
- *Il numero di giorni prima che l'account venga disabilitato* — Il numero di giorni dopo i quali si verifica la scadenza della password e prima che l'account venga disabilitato.
- *Data da quando l'account è stato disabilitato* — La data (conservata come il numero di giorni da quando si è verificata la disabilitazione) da quando l'user account è stato disabilitato.
- *Un campo riservato* — Rappresenta un campo ignorato da Red Hat Enterprise Linux.

Ecco un esempio da `/etc/shadow`:

```
juan:$1$.QKDPc5E$SWlkjRWexrXYgc98F.:12825:0:90:5:30:13096:
```

Questa riga mostra le seguenti informazioni per l'utente `juan`:

- L'ultima volta che la password è stata modificata è stato 11 Febbraio 2005
- Non vi è alcun limite minimo di tempo prima del quale si può cambiare la password
- La password deve essere cambiata ogni 90 giorni
- L'utente riceverà un preavviso cinque giorni prima di cambiare la password
- L'account verrà disabilitato 30 giorni dopo la scadenza della password se nessun tentativo di log in è stato eseguito
- L'account scadrà il 9 Novembre 2005

Per maggiori informazioni sul file `/etc/shadow`, consultare la pagina `man 5 shadow`.

6.3.2.3. `/etc/group`

Il file `/etc/group` può essere letto da tutti e contiene un elenco di gruppi, ognuno dei quali su di una riga separata. Ogni riga possiede quattro campi, separati da due punti, includendo le seguenti informazioni:

- *Group name* — Il nome del gruppo. Usato da diversi programmi utility come identificatore del tipo 'human-readable' per il gruppo.
- *Group password* — Se impostata, permette agli utenti che non fanno parte del gruppo, di unirsi al gruppo stesso usando il comando `newgrp` e digitando la relativa password. Se è presente in questo campo una `x` minuscola, allora significa che sono state usate le `group password shadow`.
- *Group ID (GID)* — L'equivalente numerico del `group name`. Viene usato dal sistema operativo e dalle applicazioni quando vengono determinati i privilegi per l'accesso.
- *Member list* — Un elenco di utenti, delimitato da una virgola, che appartengono al gruppo.

Ecco un esempio da `/etc/group`:

```
general:x:502:juan,shelley,bob
```

Questa riga mostra l'utilizzo da parte del gruppo `general` delle `password shadow`, detiene un GID di 502, e che `juan`, `shelley`, e `bob` sono i membri.

Per maggiori informazioni su `/etc/group`, consultare la pagina `man 5 group`.

6.3.2.4. `/etc/gshadow`

Il file `/etc/gshadow` è leggibile solo dall'utente `root` e contiene una password cifrata per ogni gruppo, contiene anche il `membership` del gruppo e le informazioni sull'amministratore. Proprio come nel file `/etc/group`, ogni informazione inerente il gruppo viene riportata su di una riga separata. Ogni riga è un elenco delimitato da due punti e include le seguenti informazioni:

- *Group name* — Il nome del gruppo. Usato da diversi programmi utility come identificatore del tipo 'human-readable' per il gruppo.
- *Password cifrata* — La password cifrata per il gruppo. Se impostata, i membri che non fanno parte del gruppo possono unirsi al gruppo stesso digitando la password e usando il comando `newgrp`. Se il valore di questo campo è `!`, nessun utente è in grado di unirsi al gruppo usando il comando `newgrp`. Il valore `!!` verrà trattato in egual modo di `!` — tuttavia, indica anche che non è mai stata impostata una password. Se il valore è nullo, solo i membri del gruppo possono eseguire il log in all'interno del gruppo stesso.
- *Amministratori del gruppo* — I membri del gruppo presenti in questo elenco (delimitato da una virgola), possono aggiungere o sottrarre membri usando il comando `gpasswd`.

- *Membri del gruppo* — I membri del gruppo presenti in questo elenco (delimitato da una virgola) sono membri così detti normali, e non sono membri amministrativi del gruppo.

Ecco un esempio da `/etc/gshadow`:

```
general:::shelley:juan,bob
```

Questa riga mostra che il gruppo `general` non possiede password e accetta tutti i membri che vogliono unirsi al gruppo stesso usando il comando `newgrp`. In aggiunta, `shelley` è un amministratore del gruppo, mentre `juan` e `bob` sono membri normali, non amministrativi.

Poichè modificando i file manualmente si corre il rischio di avere errori nella sintassi, è consigliato usare le applicazioni fornite da Red Hat Enterprise Linux. La sezione successiva riporta i diversi tool usati per eseguire questi compiti.

6.3.3. User Account e Group Application

Sui sistemi Red Hat Enterprise Linux vi sono due tipi di applicazione basiche, in grado di essere utilizzate durante la gestione degli user account e dei gruppi.

- L'applicazione grafica **Utente Manager**
- Una suite di tool a linea di comando

Per istruzioni dettagliate sull'uso di **Utente Manager**, consultate il capitolo *User e Group Configuration*

Mentre l'applicazione **Utente Manager** e le utility della linea di comando eseguono essenzialmente lo stesso compito, i tool della linea di comando hanno il vantaggio di essere personalizzabili e quindi più facilmente automatizzabili.

La seguente tabella descrive alcuni dei tool più comuni della linea di comando, usati per creare e gestire gli user account e i gruppi:

Applicazione	Funzione
<code>/usr/sbin/useradd</code>	Aggiunge gli user account. Questo tool viene usato per specificare il group membership primario e secondario.
<code>/usr/sbin/userdel</code>	Cancella gli user account.
<code>/usr/sbin/usermod</code>	Modifica gli attributi dell'account incluso alcune funzioni relative all'invecchiamento della password. Per un controllo migliore, usare il comando <code>passwd</code> . <code>usermod</code> viene usato anche per specificare i group membership primari e secondari.
<code>passwd</code>	Imposta le password. Anche se usato principalmente per cambiare la password dell'utente, esso controlla anche tutti gli aspetti inerenti l'invecchiamento della password.
<code>/usr/sbin/chpasswd</code>	Legge in un file composte da coppie di nome utente e password, ed aggiorna di conseguenza la password di ciascun utente.
<code>chage</code>	Cambia la policy sull'invecchiamento della password di ogni utente. Il comando <code>passwd</code> può essere usato anche per questo scopo.
<code>chfn</code>	Cambia le informazioni GECOS dell'utente.
<code>chsh</code>	Cambia la shell di default dell'utente.

Tabella 6-2. Tool della linea di comando per la gestione dell'utente

La seguente tabella descrive alcuni dei tool della linea di comando più comuni, usati per creare e gestire i gruppi:

Applicazione	Funzione
<code>/usr/sbin/groupadd</code>	Aggiunge i gruppi, ma non assegna gli utenti a quei gruppi. I programmi <code>useradd</code> e <code>usermod</code> dovrebbero essere usati per assegnare utenti ad un determinato gruppo.
<code>/usr/sbin/groupdel</code>	Cancella i gruppi.
<code>/usr/sbin/groupmod</code>	Modifica i GID oppure i nomi del gruppo, ma non cambia il group membership. I programmi <code>useradd</code> e <code>usermod</code> dovrebbero essere usati per assegnare gli utenti ad un determinato gruppo.
<code>gpasswd</code>	Cambia il group membership ed imposta le password per permettere ai membri che non appartengono al gruppo, ma che sono a conoscenza della password, di unirsi al gruppo stesso. Esso viene usato per specificare gli amministratori del gruppo.
<code>/usr/sbin/grpck</code>	Controlla l'integrità dei file <code>/etc/group</code> e <code>/etc/gshadow</code>

Tabella 6-3. Tool della linea di comando per la gestione del gruppo

I tool elencati pocanzi conferiscono agli amministratori di sistema una grossa flessibilità nel controllo di tutti gli aspetti inerenti gli user account ed i group membership. Per saperne di più sul loro funzionamento, consultate le relative pagine man.

Queste applicazioni, tuttavia, non determinano a quali risorse gli utenti ed i gruppi possono esercitare il proprio controllo. Per questo, l'amministratore deve usare le applicazioni riguardanti il permesso di un file.

6.3.3.1. Applicazioni per il permesso di un file

I permessi di un file all'interno di una organizzazione, sono parte integrante per la gestione delle risorse. La seguente tabella descrive alcuni dei tool della linea di comando più comuni usati per questo scopo.

Applicazione	Funzione
<code>chgrp</code>	Cambia il gruppo proprietario di un file
<code>chmod</code>	Cambia i permessi d'accesso per un file. È in grado altresì di assegnare permessi speciali.
<code>chown</code>	Cambia l'ownership di un file (ed è in grado di cambiare il gruppo)

Tabella 6-4. Tool della linea di comando per la gestione dei permessi

È possibile alterare questi attributi negli ambienti grafici GNOME e KDE. Fare clic con il tasto destro del mouse sull'icona del file (per esempio, quando l'icona viene mostrata in un file manager grafico o su di un desktop), e selezionare **Proprietà**.

6.4. Risorse aggiuntive

Questa sezione include diverse risorse che possono essere usate per informazioni aggiuntive sulla gestione delle risorse e dell'account, e sugli argomenti riguardanti Red Hat Enterprise Linux affrontati in questo capitolo.

6.4.1. Documentazione installata

Le seguenti risorse vengono installate nel corso di una installazione tipica di Red Hat Enterprise Linux, e possono essere utili per avere informazioni aggiuntive sugli argomenti affrontati in questo capitolo.

- Entry del menu **Utente Manager Aiuto** — Imparate a gestire gli user account ed i gruppi.
- Pagina man di `passwd(5)` — Per saperne di più sul formato del file `/etc/passwd`.
- Pagina man di `group(5)` — Per saperne di più sul formato del file `/etc/group`.
- Pagina man di `shadow(5)` — Per saperne di più sul formato del file `/etc/shadow`.
- Pagina man di `useradd(8)` — Imparate a creare ed aggiornare gli user account.
- Pagina man di `userdel(8)` — Imparate a cancellare gli user account.
- Pagina man di `usermod(8)` — Imparate a modificare gli user account.
- Pagina man di `passwd(1)` — Imparate ad aggiornare la password dell'utente.
- Pagina man di `chpasswd(8)` — Imparate ad aggiornare contemporaneamente diverse password dell'utente.
- Pagina man di `chage(1)` — Imparate a modificare le informazioni sull'invecchiamento della password dell'utente.
- Pagina man di `chfn(1)` — Imparate a modificare le informazioni GECOS di un utente (*finger*).
- Pagina man di `chsh(1)` — Imparate a modificare la shell di log in di un utente.
- Pagina man di `groupadd(8)` — Imparate a creare un nuovo gruppo.
- Pagina man di `groupdel(8)` — Imparate a cancellare un gruppo.
- Pagina man di `groupmod(8)` — Imparate a modificare un gruppo.
- Pagina man di `gpasswd(1)` — Imparate ad amministrare i file `/etc/group` e `/etc/gshadow`.
- Pagina man di `grpck(1)` — Imparate a verificare l'integrità dei file `/etc/group` e `/etc/gshadow`.
- Pagina man di `chgrp(1)` — Imparate a modificare l'ownership del group-level.
- Pagina man di `chmod(1)` — Imparate a modificare i permessi per l'accesso al file.
- Pagina man di `chown(1)` — Imparate a modificare il proprietario del file ed il gruppo.

6.4.2. Siti web utili

- <http://www.bergen.org/ATC/Course/InfoTech/passwords.html> — Un buon esempio per la comunicazione delle informazioni inerenti la sicurezza della password agli utenti di una organizzazione.
- <http://www.crypticide.org/users/alecm/> — Homepage dell'autore di uno dei programmi più conosciuti per il password-cracking (Crack). Potete scaricare (Crack) direttamente da questa pagina, e controllare gli utenti che possiedono una password debole.

- <http://www.linuxpowered.com/html/editorials/file.html> — una buona panoramica sui permessi dei file di Linux.

6.4.3. Libri correlati

I seguenti libri affrontano le problematiche relative alla gestione delle risorse e degli account, e possono coadiuvare il lavoro degli amministratori di sistema di Red Hat Enterprise Linux.

- *Red Hat Enterprise Linux Security Guide*; Red Hat, Inc. — Fornisce una panoramica sugli aspetti riguardanti la sicurezza degli account, e sulla scelta di password forti.
- *Red Hat Enterprise Linux Reference Guide*; Red Hat, Inc. — Contiene informazioni dettagliate sugli utenti e sui gruppi presenti in Red Hat Enterprise Linux.
- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — Include un capitolo sulla configurazione dei gruppi e degli utenti.
- *Linux Administration Handbook* di Evi Nemeth, Garth Snyder, e Trent R. Hein; Prentice Hall — Fornisce un capitolo sulla gestione di un user account, una sezione sulla sicurezza in relazione ai file di un user account, e una sezione sugli attributi del file ed i suoi permessi.

Capitolo 7.

Stampanti e stampa

Le stampanti rappresentano una risorsa essenziale per la creazione di una *hard copy* — una descrizione fisica su carta dei dati — una versione di documenti e collaterali idonei per aziende, per uso accademico e privato. Le stampanti sono divenute periferiche indispensabili in tutti i livelli dell'informatica, delle aziende e istituzionali.

Questo capitolo affronta le diverse stampanti disponibili, e mette a confronto i diversi usi negli ambienti informatici. Descrive anche il modo in cui Red Hat Enterprise Linux supporta il processo di stampa.

7.1. Tipi di stampante

Come qualsiasi altra periferica sono disponibili diverse stampanti. Alcune stampanti impiegano delle tecnologie che imitano il tipo di scrittura manuale, mentre altre spruzzano inchiostro, oppure usano un getto laser per generare delle immagini. Sono presenti delle interfacce hardware della stampante con PC o con la rete, che usano protocolli paralleli o seriali per il networking dei dati. Ci sono diversi fattori da tener presente per la valutazione della stampante più idonea per un impiego nel vostro ambiente informatico.

Le seguenti sezioni affrontano le varie tipologie di stampanti ed i protocolli da esse usate per comunicare con i computer.

7.1.1. Considerazioni sulla stampa

Nella valutazione di una stampante sono da tener in considerazione diversi aspetti. Quanto segue riporta alcuni criteri comuni per determinare le vostre esigenze di stampa.

7.1.1.1. Funzione

La valutazione dei requisiti della vostra organizzazione e come una stampante possa far fronte ad essi, è essenziale per determinare il giusto tipo di stampante per il vostro ambiente. La domanda da porsi è "*Cosa c'è da stampare?*" A causa della presenza di stampanti di testo specializzate, immagini ecc., dovete essere sicuri di ottenere il tool più idoneo ai vostri scopi.

Per esempio, se avete bisogno di immagini a colori ad alta qualità, è consigliato l'utilizzo di una stampante con un processo di sublimazione '*dye-sublimation*' oppure del tipo trasferimento termico '*thermal wax transfer*', invece di usarne una di tipo laser o ad impatto.

Al contrario le stampanti laser o a getto sono ideali per la stampa di documenti o di bozze, per una distribuzione interna (come ad esempio stampanti del tipo '*high volume*' generalmente chiamate stampanti *workgroup*). La determinazione dei requisiti di un utente, permette ad un amministratore di scegliere la stampante più idonea.

Altri fattori da considerare sono il *duplexing* — la possibilità di stampare entrambe le facciate di un foglio. Tradizionalmente, le stampanti potevano eseguire il lavoro di stampa solo su di un lato del foglio (generalmente chiamato stampa *simplex*). Molti modelli di tipo '*lower-end*' non hanno come default il metodo '*duplexing*' (tuttavia essi permettono all'utente di eseguire manualmente un lavoro di '*duplexing*' e cioè girando le pagine). Alcuni modelli offrono la possibilità di integrare un hardware aggiuntivo per ottenere un lavoro di duplexing; da considerare però che la suddetta aggiunta di hardware, potrebbe aumentare considerevolmente i costi. Tuttavia, il lavoro di duplexing potrebbe diminuire, a lungo andare, i costi derivati dall'uso di carta, riducendo così i costi riguardanti al *materiale di consumo* —.

Altro fattore da considerare è la misura del foglio. Molte stampanti sono in grado di gestire le misure più comuni:

- *lettera* — (8 1/2" x 11")
- *A4* — (210mm x 297mm)
- *JIS B5* — (182mm x 257mm)
- *legal* — (8 1/2" x 14")

Se alcune sezioni (come ad esempio marketing o design) hanno requisiti specifici come la creazione di poster o banners, possono essere usate stampanti del tipo *large-format* in grado di usare fogli A3 (297mm x 420mm) o tabloid (11" x 17"). In aggiunta, ci sono delle stampanti capaci di usare misure più grandi, anche se quest'ultime sono usate per scopi più specifici, come ad esempio la stampa dei blueprint.

In aggiunta, è da considerare i contenuti di tipo 'high-end', come ad esempio moduli di rete per un workgroup, ed il remote site printing.

7.1.1.2. Costo

Nella scelta della stampante da considerare anche il fattore costo. Tuttavia, determinare il costo associato all'acquisto della stampante, non è sufficiente. Ci sono altri costi da considerare, come ad esempio i beni di consumo, il mantenimento e le parti di ricambio.

Come indicato dal nome, il materiale di consumo rappresenta generalmente il materiale consumato durante il lavoro di stampa. Tale materiale è rappresentato dall'*inchiostro* e da altri *materiali* utilizzati per la stampa.

I suddetti materiali vengono usati per la stampa del testo o dell'immagine. La scelta del materiale dipende dal tipo di informazioni da stampare.

Per esempio, la creazione di una stampa molto accurata di una immagine digitale, necessita di un tipo di foglio speciale in grado di resistere ad una esposizione prolungata a luce naturale o artificiale, e di assicurare una certa accuratezza nella riproduzione del colore, queste qualità sono conosciute come *solidità* del colore. Per documenti del tipo *archival-quality* resistenti nel tempo con un alto livello professionale di leggibilità (come contratti, curriculum e rapporti permanenti), dovrebbe essere usato un foglio di tipo *opaco* (e non lucido). È importante anche lo *stock* (o spessore) del foglio, in quanto alcune stampanti non hanno un percorso diretto. L'utilizzo di carta troppo sottile o troppo spessa potrebbe verificare il blocco della stampante. Alcune stampanti sono in grado di stampare su *fogli lucidi*, in modo da proiettare su schermo le informazioni desiderate.

I suddetti materiali specializzati, come precedentemente riportato, possono avere un certo impatto sui costi del materiale di consumo, e quindi da tener presente nella valutazione dei propri requisiti.

Inchiostro è un termine generale, in quanto non tutte le stampanti usano inchiostro liquido. Per esempio, le stampanti laser usano una polvere conosciuta come *toner*, mentre le stampanti ad impatto usano nastri saturi d'inchiostro. Vi sono stampanti specializzate che riscaldano l'inchiostro durante il processo di stampa, mentre altre spruzzano piccoli quantitativi d'inchiostro sul materiale usato per il suddetto processo. Il costo per la sostituzione dell'inchiostro, una volta terminato, può variare e dipende dal tipo di container che detiene l'inchiostro stesso, e cioè se esso può essere *ricaricato* (riempito) o se necessita la sostituzione della *cartuccia*.

7.2. Stampanti ad impatto

Le stampanti ad impatto rappresentano la tecnologia di stampa più vecchia ancora in uso. Alcuni dei maggiori rivenditori continuano a produrre, vendere e supportare questo tipo di stampante. Le

stampanti ad impatto sono le stampanti più idonee per un ambiente di stampa a basso costo. Le tre forme più comuni di stampanti ad impatto sono *dot-matrix*, *daisy-wheel*, e *line printers*.

7.2.1. Stampanti dot-Matrix

La tecnologia usata per le stampanti del tipo dot-matrix è molto simile. Il foglio viene pressato contro un *tamburo* (un cilindro ricoperto), e viene tirato in avanti in modo intermittente durante la prosecuzione della stampa. La *testina di stampa* guidata elettromagneticamente, si muove attraverso il foglio e colpisce il nastro della stampante situato tra il foglio ed il pin della testina di stampa. L'impatto stampa sul foglio alcuni punti d'inchiostro i quali formano caratteri in grado di essere letti dall'uomo.

Le stampanti dot-matrix variano a seconda della risoluzione e della qualità di stampa, esse possono essere costituite da 9 o 24-pin testine di stampa. Maggiore è il numero di pin, e maggiore sarà la risoluzione. La maggior parte delle stampanti dot-matrix hanno una risoluzione massima di circa 240 *dpi* (punti per pollice). Anche se la risoluzione non è qualitativamente ai livelli delle stampanti laser o a getto, la stampa di tipo dot-matrix (o altro tipo di stampa ad impatto), presenta dei vantaggi. A causa della forza che la testina di stampa deve imprimere sulla superficie del foglio per trasferire l'inchiostro presente nel nastro sulla pagina, questo tipo di stampa è ideale per ambienti che hanno necessità di produrre *copie carbone* attraverso l'uso di documenti con diverse sezioni. Questi documenti hanno la parte inferiore composta di carbone (o altri materiali sensibili alla pressione) e creano un segno sul foglio sottostante quando si applica pressione. Rivenditori o piccole aziende spesso usano queste copie carbone come ricevuta o scontrini fiscali.

7.2.2. Stampanti Daisy-Wheel

Se avete avuto esperienza con una qualsiasi macchina da scrivere, allora sarete in grado di copire il tipo di tecnologia che è alla base delle stampanti di tipo daisy-wheel. Le testine di stampa di queste stampanti sono composte da rulli di plastica o di metallo intagliati in *petali*. Ogni petalo ha la forma di una lettera (maiuscola e minuscola), numeri, o di punteggiatura. Quando uno di questi petali colpisce il nastro della stampante, ne risulta, sul foglio, un carattere composto da inchiostro. Questo tipo di stampante risulta essere lenta e rumorosa. Esse non sono in grado di stampare grafici, e non sono in grado di modificare le fonti a meno che non si sostituisce fisicamente il rullo. Con l'avvento di stampanti laser, le stampanti di tipo daisy-wheel vengono usate sempre meno in ambienti informatici moderni.

7.2.3. Stampanti di linea

Un altro tipo di stampa ad impatto, in alcuni casi simile al tipo daisy-wheel, è la *line printer*. Tuttavia, invece di un rullo di stampa, le line printer hanno un meccanismo che permettono a diversi caratteri di essere stampati simultaneamente sulla stessa riga. Il meccanismo è in grado di usare un *tamburo di stampa* rotante oppure una *catena di stampa* racchiusa da un anello. Durante la rotazione del tamburo o della catena sulla superficie del foglio, alcuni martelletti elettromeccanici posti dietro al foglio, lo spingono (insieme al nastro) sulla superficie del tamburo o della catena, segnando il foglio con la forma del carattere presente su uno di essi.

A causa della natura del meccanismo di stampa, le line printer sono molto più rapide delle stampanti di tipo dot-matrix o daisy-wheel. Tuttavia possono essere molto rumorose, con una capacità 'multi-font', e spesso offrono una bassa qualità di stampa rispetto alle più recenti tecnologie.

A causa della loro rapidità, le line printer usano carta speciale del tipo *tractor-fed*, che presentano, su ambedue i lati, una serie di buchi. Tale caratteristica rende possibile un tipo di stampa molto veloce, con interruzioni necessarie solo quando lo scatolo contenente la carta si esaurisce.

7.2.4. Materiali di consumo per la stampante ad impatto

Su tutte le stampanti, quelle ad impatto hanno un costo, inerente al materiale di consumo, molto basso. I nastri ad inchiostro e la carta, rappresentano i costi primari per le stampanti ad impatto. Alcune di esse (generalmente quelle di tipo dot-matrix e le line printer) hanno bisogno di una carta speciale chiamata 'tractor-fed', che in alcuni casi può aumentare i costi operativi.

7.3. Stampanti a getto 'Inkjet'

Una stampante del tipo *Inkjet* utilizza la tecnologia più diffusa nel campo delle stampanti. Le caratteristiche relativamente a basso costo e idonee per compiti multipli, rendono queste stampanti molto versatili e idonee per un uso personale e per piccole aziende.

Le suddette stampanti utilizzano un inchiostro particolare, ad acqua, in grado di asciugarsi rapidamente, sono composte da una serie di piccole bocche che spruzzano inchiostro sulla superficie del foglio. La testina di stampa viene guidata da una cinghia mossa da un motorino, in grado di muovere la suddetta testina attraverso il foglio.

Le stampanti di tipo Inkjet sono state create originariamente per stampare in modo *monocromatico* (solo bianco e nero). Tuttavia le testine di stampa sono state migliorate includendo oggi diversi colori come il magenta, il giallo, il nero e il cyan. Questa combinazione (chiamata *CMYK*), permette di stampare le immagini con una qualità molto simile a quella riprodotta da un laboratorio di sviluppo di fotografie (quando si usano però determinati tipi di carta politenata.) Quando accoppiati con una stampa di testo di qualità, le stampanti inkjet rappresentano la scelta migliore per una stampa monoscromatica o a colori.

7.3.1. Materiali di consumo per le stampanti Inkjet

Le stampanti Inkjet generalmente sono di basso costo ma tendono ad incrementare in base alla qualità, ai contenuti extra, e alla capacità di stampare su formati più grandi rispetto alle normali stampanti di base. Anche se il costo di una stampante Inkjet è basso rispetto ad altri tipi di stampanti, è da tener in giusta considerazione il costo dovuto al materiale di consumo. A causa dell'alta richiesta per stampanti inkjet e dell'espansione della gamma informatica da un uso privato ad enterprise, l'acquisto di materiale di consumo potrebbe essere costoso.



Nota bene

Se volete acquistare una stampante inkjet, assicuratevi sempre di sapere il tipo di cartuccia da usare. Questo è molto importante soprattutto se scegliete una stampante a colori. Le stampanti inkjet CMYK richiedono inchiostro per ogni colore; tuttavia, è importante sapere se ogni colore viene conservato in una cartuccia separata oppure no.

Alcune stampanti usano una cartuccia con diverse sezioni; questo però potrebbe rappresentare un punto negativo, in quanto a meno che non vi sia un modo per ricaricare un determinato colore, l'intera cartuccia deve essere sostituita. Altre stampanti usano una cartuccia con diverse sezioni per il cyan, magenta, e giallo, ma hanno anche una cartuccia separata per il nero. In un ambiente dove il volume di stampa per il testo è molto elevato, questo tipo di disposizione può risultare idonea. Tuttavia, la soluzione migliore è quella di trovare una stampante con cartucce separate per ogni colore; sarete così in grado di ricaricare facilmente qualsiasi colore ogni qualvolta si esaurisce.

Alcune case produttrici di stampanti inkjet, consigliano l'uso di carta speciale per una stampa d'immagini e documenti qualitativamente molto elevata. Questo tipo di carta usa un rivestimento lucido capace di assorbire inchiostro colorato, il quale previene il fenomeno di *clumping* (la tendenza dell'inchiostro a base di acqua, di raggrupparsi in certe aree, causando delle macchie sul foglio

stampato) o *banding* (dove si verifica la presenza di linee estranee al lavoro di stampa.) Consultate la documentazione della vostra stampante sull'uso dei fogli consigliati.

7.4. Stampanti laser

Le stampanti laser rappresentano un'altra alternativa alla stampa ad impatto. Le suddette stampanti sono conosciute per il loro 'high volume output' e per un basso costo per pagina. Esse vengono spesso impiegate nelle aziende sotto forma di un workgroup o centro di stampa dipartimentale, dove la prestazione, resistenza, e le esigenze output sono una priorità. A causa della capacità delle stampanti laser ad ottemperare a questi requisiti (con un costo per pagina molto ragionevole), la tecnologia viene ampiamente riconosciuta come il cavallo di battaglia del metodo di stampa usato dalle suddette aziende.

Le stampanti laser condividono la stessa tecnologia delle fotocopiatrici. Alcuni rulli tirano il foglio da un vassoio attraverso un *rullo di carica*, il quale conferisce una carica elettrostatica. Contemporaneamente viene data una carica opposta ad un tamburo di stampa. La superficie del tamburo viene analizzata da un raggio laser, che prende in considerazione solo i punti corrispondenti al testo o ad una immagine desiderati. La carica conferita al testo o all'immagine viene usata successivamente per far aderire il toner alla superficie del tamburo.

Il foglio e il tamburo vengono portati a contatto; la diversa carica a loro conferita fa in modo che il toner aderisca alla carta. E per finire, il foglio passa tra alcuni *fusing roller*, che riscaldano il foglio e sciogliono il toner, fondendolo così sulla superficie del foglio.

7.4.1. Stampanti laser a colori

Le stampanti laser a colori cercano di combinare le caratteristiche migliori della tecnologia laser e quella inkjet, raggruppandole in un pacchetto di tipo multi-uso. Questa tecnologia è basata su di una stampa laser monocromatica tradizionale, utilizzando però componenti aggiuntivi per creare immagini e documenti a colori. Invece di usare solo un toner nero, le stampanti laser a colori usano una combinazione toner di tipo CMYK. Il tamburo di stampa è in grado sia di ruotare ogni singolo colore o di usare un colore per volta, oppure di usarne tutti e quattro contemporaneamente su di una 'base', per poi passare il foglio attraverso il tamburo, trasferendo l'immagine sul foglio stesso. Le stampanti laser a colori usano anche un *olio di fusione* insieme ai rulli di fusione 'heated fusing roll', i quali legano maggiormente il toner al foglio, dando così diversi tipi di lucentezza all'immagine finita.

A causa delle caratteristiche migliorate, le stampanti laser a colori sono generalmente più costose di quelle di tipo laser monocromatico. Considerando i costi implicati con queste tecnologie di stampa, alcuni amministratori potrebbero desiderare di separare la funzionalità monocromatica (testo) ed il colore (immagine), rispettivamente per una apposita stampante laser monocromatica ed una stampante laser (o inkjet) a colori.

7.4.2. Materiali di consumo per stampanti a colori

A seconda della stampante laser impiegata, i costi sono proporzionali al volume di lavoro. Il toner è contenuto in cartucce che generalmente possono essere sostituite; tuttavia alcuni modelli possono avere cartucce ricaricabili. Le stampanti laser a colori necessitano una cartuccia per ogni colore. In aggiunta, esse hanno bisogno di olii speciali per far amalgamare nel miglior modo possibile il toner con i fogli, limitando così ogni possibilità di perdite. Tutto ciò non fa altro che aumentare il costo dei materiali di consumo delle stampanti laser a colori; tuttavia, è da puntualizzare che i suddetti materiali in linea di massima, sono sufficienti per 6000 pagine, molto di più di una normale stampante ad impatto o inkjet. I fogli da utilizzare con le stampanti laser non rappresentano un problema particolare, in quanto è possibile utilizzare in molti casi, fogli normali da fotocopia o xerografici. Se desiderate invece eseguire una stampa di immagini ad alta qualità, si consiglia di utilizzare carta lucida.

7.5. Altri tipi di stampante

Sono disponibili altri tipi di stampanti, molte delle quali per scopi particolari tipo grafiche professionali o editoriali. Tuttavia queste stampanti non sono ad uso generale. A causa della loro particolare classificazione, il loro costo (per quanto riguarda i materiali di consumo), tende ad essere più elevato rispetto alla struttura principale.

Stampanti a cera termica

Queste stampanti vengono usate maggiormente per diapositive e per *provini* a colori (creando documenti di testo ed immagini prima di inviare i documenti master per la stampa su stampanti in offset a quattro colori). Le stampanti a cera termica usano nastri CMYK guidati da una cinghia e fogli particolari o trasparenti. La testina di stampa contiene elementi capaci di riscaldare ogni colore sul foglio man mano che lo stesso attraversa la stampante.

Stampanti a sublimazione

Usate in organizzazioni come i service bureau — dove la qualità dei documenti, opuscoli e presentazioni è prioritaria rispetto ai beni di consumo — Il tipo di stampante a sublimazione 'dye-sublimation (o dye-sub)' rappresenta il cavallo di battaglia nella qualità di stampa CMYK. I concetti dietro alle stampanti dye-sub sono molto simili alle stampanti a cera termica ad eccezione dell'uso di pellicole dye plastiche invece di cera colorata. La testina di stampa riscalda la pellicola colorata e vaporizza l'immagine su di un foglio ricoperto in modo speciale.

Dye-sub è molto diffuso nel campo del design, editoriale ed anche della ricerca scientifica, dove è necessaria un tipo di stampa molto precisa e dettagliata. Tali caratteristiche però comportano anche alcuni costi, in quanto queste stampanti sono conosciute anche per il loro alto costo per pagina.

Stampanti ad inchiostro solido

Usate maggiormente nel campo del design industriale e nell'imballaggio, il costo per le stampanti ad inchiostro solido variano a seconda della loro abilità di stampare su diversi tipi di fogli. Le suddette stampanti, come implica anche il loro nome, usano inchiostro solido che viene sciolto a spruzzato attraverso delle piccolissime bocche presenti sulla testina di stampa. Il foglio viene poi inviato attraverso un rullo di fusione in modo da amalgamare maggiormente l'inchiostro con il foglio stesso.

La stampante ad inchiostro solido è ideale per la produzione di provini per nuovi design per le confezioni di prodotti, per questo motivo numerose aziende non hanno la necessità di usare questo tipo di stampante.

7.6. Tecnologie e linguaggi di stampa

Prima dell'avvento della tecnologia laser e inkjet, le stampanti ad impatto erano in grado di stampare in maniera standard, con nessuna variazione nella dimensione delle lettere o stile. Oggi, le stampanti sono in grado di processare documenti complessi con immagini di tipo embedded, carte e tabelle in diversi linguaggi e immagini multiple, tutto in una pagina. Tale complessità deve aderire ad alcune convenzioni riguardanti il formato. Per questo motivo si è arrivati al *page description language* (o PDL) —, un linguaggio specializzato di formattazione del documento creato in particolare per la comunicazione tra i computer e le stampanti.

Attraverso gli anni, i produttori di stampanti hanno sviluppato i propri linguaggi proprietari per descrivere i diversi formati. Tuttavia, tali linguaggi sono applicati solo alle stampanti che i produttori hanno loro stesso creato. Se, per esempio, avete inviato un file ad una tipografia professionale pronto per la stampa, usando un PDL proprietario, non è sicuro che il file inviato sia compatibile con le macchine della stampante. In questo caso verrà evidenziato il problema della portabilità.

Xerox® ha sviluppato il protocollo Interpress™ per le proprie stampanti, in contrasto non è stato mai adottato tale linguaggio dal resto dell'industria. Due sviluppatori che hanno contribuito allo sviluppo di Interpress hanno abbandonato Xerox e formato Adobe®, una compagnia software che concentra le proprie risorse principalmente nelle grafiche elettroniche e in documenti professionali. Adobe ha sviluppato un PDL chiamato *PostScript*™, il quale è in grado di utilizzare un linguaggio per descrivere una formattazione di testo e delle informazioni d'immagine che possono essere processate dalle stampanti. Allo stesso tempo, la Compagnia Hewlett-Packard® ha sviluppato il *Printer Control Language*™ (o PCL) per un utilizzo nella propria gamma di stampanti laser e inkjet. PostScript e PCL sono PDL molto comuni e sono supportati da molte case produttrici.

I PDL funzionano con lo stesso principio dei linguaggi di programmazione del computer. Quando un documento è pronto per essere stampato, il PC o la workstation prende le immagini, le informazioni tipografiche, e l'impostazione del documento, usandoli come oggetti che formano le istruzioni per la stampante da processare. La stampante traduce i suddetti oggetti in *rasters*, una serie di righe capaci di formare una immagine del documento (chiamata *Raster Image Processing* o RIP), e stampa l'output sulla pagina come una immagine, completa di testo e di grafiche. Questo processo rende il lavoro di stampa più costante, risultando talvolta in una leggera variazione nella stampa tra una stampante e l'altra. I PDL sono stati creati per essere idonei ad ogni formato, e scalabili per essere adatti su fogli con misure differenti.

La scelta della stampante giusta dipende esclusivamente dalla qualità che il vostro ufficio, o sezione all'interno di una organizzazione, ha adottato per far fronte ai propri requisiti. Molti dipartimenti usano word processing o altri software di produttività che utilizzano il linguaggio PostScript. Tuttavia, è da prendere anche in considerazione la necessità di utilizzare PDL o altra forma di stampa proprietaria,

7.7. Networked contro le stampanti locali

A seconda delle necessità organizzative potrebbe essere non necessario assegnare una stampante ad ogni membro appartenente alla vostra organizzazione. I suddetti costi possono rappresentare una grossa fetta del budget, lasciando disponibili meno risorse. Mentre le stampanti locali collegate tramite un cavo parallelo o USB ad ogni workstation, rappresentano una soluzione ideale per l'utente, generalmente non lo è sotto il profilo economico.

I produttori di stampanti hanno affrontato la necessità di sviluppare delle stampanti *departmentali* (o workgroup). Queste macchine sono generalmente resistenti, veloci, con materiali di consumo a lunga durata. Le stampanti di tipo workgroup generalmente sono collegate ad un server di stampa, un dispositivo di tipo standalone (come ad esempio una workstation riconfigurata) che gestisce i lavori di stampa e gli itinerari output alle stampanti corrette quando queste sono disponibili. Le stampanti departmentali più recenti includono delle interfacce di rete aggiunte o interne che eliminano la necessità di un server di stampa.

7.8. Informazioni specifiche su Red Hat Enterprise Linux

Quanto segue descrive le varie caratteristiche specifiche di Red Hat Enterprise Linux per quanto riguarda stampanti e lavoro di stampa.

Lo **Strumento di configurazione della stampante** permette agli utenti di configurare una stampante. Questo tool aiuta a gestire il file di configurazione della stampante, le directory spool e i filtri di stampa.

Red Hat Enterprise Linux 4 usa il sistema di stampa CUPS. Se un sistema è stato migliorato da una versione precedente di Red Hat Enterprise Linux che utilizzava CUPS, il processo di miglioramento non avrà intaccato le code precedentemente configurate, e il sistema continuerà ad usare CUPS.

Per usare **Strumento di configurazione della stampante** dovete avere i privilegi root. Per iniziare un'applicazione, selezionare **Pulsante del menu principale** (sul pannello) => **Impostazioni del sistema** => **stampa**, o inserire il comando `system-config-printer`. Questo comando determina

automaticamente se eseguire la versione grafica o di testo, a seconda se il comando viene eseguito in un ambiente desktop grafico oppure da una console di testo.

Per forzare lo **Strumento di configurazione della stampante** in modo da essere eseguito come se fosse un'applicazione di testo, usando il comando `system-config-printer-tui` dal prompt della shell.



Importante

Non modificate il file `/etc/printcap` oppure i file presenti nella directory `/etc/cups/`. Ogni volta che il demone della stampante (`cups`) viene avviato o riavviato, vengono creati dinamicamente, nuovi file di configurazione. I suddetti file vengono altresì creati dinamicamente quando vengono effettuati i cambiamenti con lo **Strumento di configurazione della stampante**.

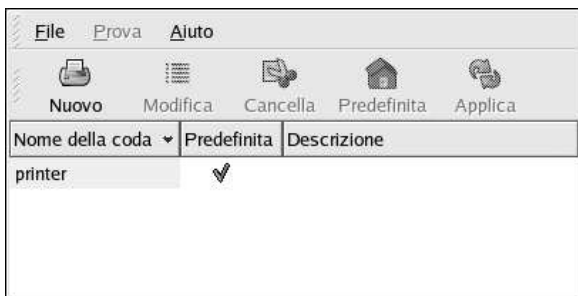


Figura 7-1. Strumento di configurazione della stampante

Le seguenti code di stampa possono essere configurate:

- **Collegata-localmente** — Una stampante collegata direttamente al computer attraverso una porta parallela o USB.
- **Networked CUPS (IPP)** — una stampante che può essere accessa tramite una rete TCP/IP attraverso il protocollo di stampa di Internet, conosciuto anche come IPP (per esempio, una stampante collegata ad un altro sistema Red Hat Enterprise Linux che esegue CUPS sulla rete).
- **Networked UNIX (LPD)** — Una stampante collegata ad un sistema UNIX diverso che può essere accesso tramite una rete TCP/IP (per esempio, una stampante collegata ad un altro sistema Red Hat Enterprise Linux che esegue LPD sulla rete).
- **Windows Rete (SMB)** — Una stampante collegata ad un sistema diverso il quale condivide una stampante attraverso una rete SMB (per esempio, una stampante collegata ad una macchina Microsoft Windows).
- **Novell Rete (NCP)** — Una stampante collegata ad un sistema diverso il quale usa una tecnologia di rete NetWare di Novell.
- **JetDirect Rete** — Una stampante collegata direttamente alla rete attraverso HP JetDirect invece di un computer.

**Importante**

Se aggiungete una nuova coda di stampa o modificate una già esistente, è necessario confermare i cambiamenti.

Facendo clic sul pulsante **Applica**, salvate qualsiasi cambiamento che avete apportato riavviando il demone della stampante. I cambiamenti non verranno scritti sul file di configurazione fino a quando il demone della stampante non viene riavviato. Alternativamente, potete scegliere **File => Salva cambiamenti** e poi scegliere **Action => Applica**.

Per maggiori informazioni sulla configurazione delle stampanti con Red Hat Enterprise Linux consultate *Red Hat Enterprise Linux System Administration Guide*.

7.9. Risorse aggiuntive

Le configurazioni di stampa e la stampa di rete, sono argomenti molto vasti, che necessitano una certa conoscenza dell'hardware, networking e della gestione del sistema. Per informazioni più dettagliate sull'impiego dei servizi della stampante nei vostri ambienti, consultate le seguenti risorse.

7.9.1. Documentazione installata

- Pagina man di `lpr(1)` — Imparate ad eseguire una stampa delle informazioni di file selezionati sulla stampante da voi scelta.
- Pagina man di `lprm(1)` — Imparate a rimuovere i lavori di stampa dalla coda di una stampante.
- Pagina man di `cupsd(8)` — Per saperne di più sul demone della stampante CUPS.
- Pagina man di `cupsd.conf(5)` — Per saperne di più sul formato del file, per il file di configurazione del demone della stampante CUPS.
- Pagina man di `classes.conf(5)` — Per saperne di più sul formato del file, per il file di configurazione classe CUPS.
- File in `/usr/share/doc/cups-<version>` — Per saperne di più sul sistema di stampa CUPS.

7.9.2. Siti web utili

- <http://www.webopedia.com/TERM/p/printer.html> — Definizioni generali delle stampanti e descrizioni delle varie tipologie.
- <http://www.linuxprinting.org/> — Un database di documenti sulla stampa, insieme con un database di quasi 1000 stampanti compatibili con i mezzi di stampa fornite dalla Linux.
- <http://www.cups.org/> — Documentazione, Domande frequenti (FAQ), e newsgroup su CUPS.

7.9.3. Libri correlati

- *Stampa di rete* di Matthew Gast e Todd Radermacher; O'Reilly & Associati, Inc. — Informazioni complete su come usare Linux come server di stampa in ambienti eterogenei.
- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — Include un capitolo sulla configurazione della stampante.

Capitolo 8.

Pianificazione di un system disaster

La pianificazione di un system disaster rappresenta una problematica che viene facilmente dimenticata dagli amministratori di sistema — non è piacevole, ma purtroppo sembra che ci sia sempre qualcos'altro a cui pensare. Tuttavia, non pianificare un disaster rappresenta l'errore più grande che un amministratore possa fare.

Anche se spesso si è portati a pensare ai diversi disastri naturali (come incendi, inondazioni o temporali), possono rientrare in questa categoria i problemi di ogni giorno che influiscono nella nostra vita quotidiana. Per questo motivo la definizione di disaster per un amministratore, rappresenta un evento non pianificato che influenza la normale funzione dell'organizzazione.

Anche se risultasse impossibile elencare tutti i tipi di disastri che si possono verificare, questa sezione esamina i fattori principali che fanno parte di ogni disaster, in modo da poter esaminare non le probabilità, ma le reali condizioni che si possa verificare tale evento.

8.1. Tipi di disaster

In generale, vi sono quattro fattori che possono innescare un disaster. Questi fattori sono:

- Problemi hardware
- Problemi software
- Problemi riguardanti l'ambiente
- Errori umani

8.1.1. Problemi hardware

I problemi hardware sono più facili da capire — se si verifica tale errore, non è più possibile lavorare. Quello che risulta più difficile capire è la natura dei problemi e come subirne il minor impatto possibile. Ecco alcuni approcci:

8.1.1.1. Avere un hardware di riserva

L'esposizione ai problemi riguardanti l'hardware può essere ridotta conservando unità hardware di riserva. Naturalmente questo approccio presuppone due cose:

- Avere qualcuno in grado di diagnosticare il problema, identificare l'hardware incriminato e sostituirlo.
- Disponibilità di una unità hardware per la sostituzione.

Queste problematiche vengono affrontate in dettaglio nelle seguenti sezioni.

8.1.1.1.1. Essere competenti

A seconda della vostra esperienza e dell'hardware incriminato, avere la necessaria competenza non rappresenterebbe una problematica. Tuttavia se non avete mai lavorato in passato con unità hardware, vi consigliamo di consultare qualche collega per seguire un corso introduttivo su come riparare i PC. Poichè tale corso potrebbe risultare non sufficiente per permettervi di affrontare i diversi problemi con un server di livello enterprise, esso rappresenta il modo migliore per ottenere quelle conoscenze

di basi necessarie (gestione corretta dei tool e dei componenti, procedure per la diagnosi di base, e così via).



Suggerimento

Prima di accingervi a porre rimedio a tale problema, assicuratevi che l'hardware in questione:

- Non sia sotto garanzia
- Non sia vincolato da alcun contratto di servizio/mantenimento di ogni tipo

Se cercate di far fronte al problema che si è verificato ma l'hardware è ancora sotto la garanzia e/o vincolato da un contratto particolare, probabilmente sarete in procinto di violare i termini di questi agreement, annullando così la copertura della garanzia o di altri accordi.

Tuttavia, anche se avete una conoscenza di base, potreste essere in grado di diagnosticare effettivamente e sostituire le unità hardware in questione — se avete precedentemente scelto una gamma idonea di unità di ricambio.

8.1.1.1.2. Cosa selezionare?

Questa domanda evidenzia le diverse sfaccettature della natura di qualsiasi elemento appartenente al disaster recovery. Quando considerate il tipo di hardware da conservare, ecco alcune delle problematiche da ricordare:

- Downtime massimo permesso
- Le conoscenze necessarie per eseguire una riparazione
- Il budget disponibile per le unità di ricambio
- Il loro spazio di storage necessario
- Altro hardware in grado di utilizzare le stesse unità di ricambio

Ciascuno di questi punti è in relazione al tipo di unità di ricambio da conservare. Per esempio, conservare un sistema completo tenderebbe a minimizzare il downtime e richiederebbe una conoscenza minima per l'installazione, ma risulterebbe più costoso rispetto ad una CPU o modulo RAM di ricambio. Tuttavia, tale spesa potrebbe essere giustificata, se la vostra organizzazione possiede diversi server tutti identici tra loro, che possono trarre beneficio da un singolo sistema di riserva.

Indipendentemente dalla decisione finale, bisogna affrontare la seguente ed inevitabile domanda.

8.1.1.1.2.1. Quanto ne devo conservare?

La questione dei livelli delle unità di ricambio possiede anch'essa diverse sfaccettature. Ecco le problematiche principali:

- Downtime massimo permesso
- Percentuale d'errore prevista
- Tempo stimato per rifornire lo stock
- Il budget disponibile per le unità di ricambio
- Il loro spazio di storage necessario
- Altro hardware in grado di utilizzare le stesse unità di ricambio

Per un sistema che è in grado di affrontare un downtime di due giorni, e che necessità di una unità di riserva che viene usata una volta l'anno, la quale può essere reperita in 24 ore, potrebbe essere consigliato di avere al massimo una unità di riserva (in certi casi è possibile anche essere sprovvisti di tale unità se siete in grado di assicurare la suddetta entro e non oltre le 24 ore di tempo).

Viceversa, un sistema che non è in grado di permettersi un periodo di downtime maggiore a due minuti, e che possiede una unità di riserva la quale viene usata una volta al mese (e che la stessa impiega diverse settimane per essere reperita), potrebbe significare la necessità di conservare non una ma diverse unità di riserva.

8.1.1.1.3. Unità di riserva che non sono di 'riserva'

Quando una unità di riserva risulta non essere di riserva? Quando essa rappresenta un hardware idoneo ad un suo uso giornaliero, ma che risulta essere anche idoneo come unità di riserva per sistemi con una priorità più elevata. Questo approccio presenta alcuni benefici:

- Meno risorse economiche per le unità di riserva "non produttive"
- L'hardware viene annoverato come unità operativa

Sono presenti, tuttavia, alcuni lati negativi per questo tipo di approccio:

- La produzione normale dei compiti a bassa priorità viene interrotta
- Presenza di una potenziale esposizione se si verifica un errore hardware con bassa priorità (lasciando nessuna unità di riserva per l'hardware con una priorità elevata)

A causa di queste limitazioni, l'utilizzo come riserva di un altro sistema di produzione potrebbe risultare utile, ma il successo di questo approccio dipende dal carico specifico di lavoro, e dall'impatto che ha l'assenza del sistema nelle operazioni generali del centro dati.

8.1.1.2. Contratti sul servizio

I contratti relativi ai servizi fanno sì che i problemi di natura hardware vengano affrontati e possibilmente risolti da altre persone. Tutto ciò che viene a voi richiesto è quello di confermare che si sia verificato un problema all'hardware e che lo stesso non sia dovuto al software. Potete successivamente contattare il personale preposto per risolvere tale problema.

Sembra facile. Ma come molte altre cose nella vita, c'è di più. Ecco alcune considerazioni da fare nella selezione di un contratto relativo ad un servizio:

- Orari di copertura
- Tempo di risposta
- Disponibilità delle parti di ricambio
- Budget disponibile
- Hardware ricoperto dal contratto

Affronteremo questi dettagli più da vicino nelle seguenti sezioni.

8.1.1.2.1. Orari di copertura

Sono disponibili diversi contratti a seconda delle diverse necessità; una delle variabili inerenti ai diversi contratti è relativa agli orari di copertura. Se non desiderate pagare una copertura di tipo premium, purtroppo non è sempre possibile chiamare in ogni momento ed aspettarsi un tecnico pronto ad assistervi.

In base al vostro contratto, potreste trovarvi nella situazione dove non sarete in grado di usufruire del servizio se non per un determinato giorno/orario, oppure non sarà possibile ottenere l'assistenza di un tecnico se non per un determinato giorno/orario specificato nel vostro contratto.

In molti casi gli orari di copertura vengono definiti in termini di orari e giorni durante i quali i tecnici sono disponibili. Alcune degli orari di copertura più comuni sono:

- Da lunedì al venerdì, dalle 09:00 alle 17:00
- Da Lunedì al venerdì, 12/18/24 ore ogni giorno (con orari prestabiliti)
- Da Lunedì al venerdì, (o dal lunedì alla domenica), stessi orari come sopra specificati

Come prevedibile, il costo di un contratto aumenta con l'aumentare della copertura del servizio. In generale, avere una copertura che va dal lunedì al venerdì tende ad essere meno costosa se invece si aggiungesse anche il sabato e la domenica.

Ma anche in questa circostanza è possibile ridurre i costi del servizio se siete in grado di svolgere, in prima persona, parte del lavoro.

8.1.1.2.1.1. Servizio Depot

Se avete bisogno solo della disponibilità di un tecnico durante l'orario lavorativo normale, e siete in possesso di una sufficiente esperienza che vi permette di determinare cosa non va, potreste considerare di utilizzare un *servizio depot*. Conosciuto sotto diversi nomi (incluso anche *walk-in service* e *drop-off service*), i rivenditori potrebbero avere dei depositi dove i tecnici possono svolgere il loro lavoro direttamente sull'hardware fornito dagli utenti.

Il suddetto servizio presenta il beneficio di essere molto rapido tanto quanto lo siete voi. Non è necessario aspettare la disponibilità di un tecnico il quale viene di persona nel vostro ufficio, casa ecc. I suddetti tecnici, non sono disponibili su chiamata, essi infatti svolgono il loro lavoro in depositi o laboratori, e appena voi stessi portate il vostro hardware in uno di questi centri, essi esplicano le loro funzioni.

Poichè questo tipo di servizio viene eseguito da un ufficio centrale, è molto probabile che la parte da sostituire sia disponibile. Questo riduce al massimo la probabilità di dover attendere l'arrivo di un nuovo pezzo di ricambio proveniente da un altro deposito distante centinaia di chilometri da voi.

Ci sono comunque alcuni lati negativi. Quello più ovvio è quello di non poter scegliere gli orari di servizio — potete usufruire di questo servizio solo quando il deposito è aperto. Un altro aspetto negativo è che i tecnici non vanno oltre il loro normale orario lavorativo, ciò significa che se il vostro problema si è verificato alle 16:30 di venerdì, e voi portate il vostro sistema al deposito per le 17:00, il lavoro verrà eseguito non prima di lunedì mattina.

Un altro fattore che può essere determinante nella scelta di questo tipo di servizio è quello di valutare se il deposito si trovi in zone limitrofe. Se la vostra organizzazione è situata in un'area metropolitana non dovrebbe essere un problema. Tuttavia, le organizzazioni che si trovano in zone rurali potrebbero avere qualche problema in più.



Suggerimento

Se desiderate scegliere tale servizio, prendete in considerazione anche il modo con il quale trasportare il vostro hardware al deposito più vicino. A tale scopo siete in grado di usare una vettura della vostra compagnia oppure la vostra? Se è la vostra, avete una capacità di carico accettabile

con il relativo spazio a voi necessario? La vostra vettura è assicurata? Per caricare l'hardware nella vettura è necessario l'aiuto di un'altra persona?

Anche se queste potrebbero essere considerazioni di secondo piano, esse devono essere risolte prima di scegliere definitivamente l'uso di questo servizio.

8.1.1.2.2. Tempo di risposta

In aggiunta agli orari di copertura, molti agreement hanno un livello specifico inerente il tempo di risposta. In altre parole, quando telefonate per chiedere assistenza, quanto tempo trascorre dalla vostra telefonata prima che un tecnico sia in grado di contattarvi? Come potete immaginare, ad un tempo di risposta molto rapido ne consegue un agreement più costoso.

Vi sono dei limiti ai tempi di risposta disponibili. Per esempio, il tempo necessario per arrivare dall'ufficio del rivenditore al vostro, risulta influenzare molto sui tempi di risposta disponibili¹. Il tempo medio di risposta che si aggira intorno alle quattro ore, viene generalmente considerato molto buono. I tempi di risposta più lenti possono andare dalle otto (il che può essere considerato come assistenza per il "giorno successivo" per un agreement standar), alle 24. Come molti agreement il costo è — negoziabile.



Nota Bene

Anche se non è una cosa comune, dovete essere a conoscenza che gli agreement che presentano alcune clausole sul tempo di risposta, possono portare il servizio di una organizzazione oltre ai limiti di adempimento della stessa. Può capitare che alcune organizzazioni molto occupate in tale servizio, possano inviare il loro personale — *chiunque* — per rispondere ad una chiamata con un tempo di risposta brevissimo, solo per rientrare nei limiti prefissati dal tempo di risposta stesso. Questa persona trova il difetto, chiama "l'ufficio" in modo che un'altra persona possa portare il "pezzo di ricambio idoneo."

Ma in effetti essi stanno aspettando la disponibilità di un tecnico in grado di far fronte alla situazione.

Questa situazione potrebbe essere accettabile in circostanze straordinarie (come ad esempio la presenza di problemi di alimentazione che hanno danneggiato la maggior parte dei sistemi presenti nell'area di servizio), ma se questo metodo viene impiegato regolarmente, è consigliabile contattare il responsabile del servizio e chiedere spiegazioni.

Se le vostre necessità richiedono un tempo di risposta molto rapido (e il vostro budget risulta essere molto ampio), vi è un approccio che può ridurre maggiormente il vostro tempo di attesa — fino a zero.

8.1.1.2.2.1. Tempo di risposta pari a zero — Significa avere un tecnico interno

Data la situazione appropriata (voi rappresentate uno dei clienti più importanti presenti nell'area), date determinate necessità (*qualsiasi* downtime non è accettabile), e risorse finanziarie (se domandate il prezzo forse significa che non potete permettervelo), potreste rappresentare il cliente giusto per avere un tecnico permanente interno. I benefici di un tecnico interno sono ovvi:

- Risposta istantanea ad un problema

1. Generalmente viene considerato come tempo di risposta ideale, in quanto i tecnici spesso sono responsabili di parte di territorio che va ben oltre quello del loro ufficio. Se vi trovate a due estremità diverse di una zona assegnata ad un particolare tecnico, il tempo di risposta potrebbe essere più lungo.

- Un approccio più attivo per la gestione di un sistema

Come potete immaginare questa opzione potrebbe risultare *molto* costosa, particolarmente se avete la necessità di avere a disposizione un tecnico 24 ore su 24, sette giorni su sette. Ma se questo approccio risulta essere il più idoneo per la vostra organizzazione, ricordatevi alcuni punti in modo da ottenere il massimo dei benefici.

Per prima cosa, i tecnici interni necessitano di più risorse rispetto ad un impiegato normale, ad esempio spazio sufficiente, un telefono, schede d'accesso appropriate e/o chiavi, e così via.

Essi possono risultare di poco aiuto se non possiedono gli strumenti giusti. Per questo, assicuratevi di avere uno storage sicuro per le parti di ricambio utili al tecnico. In aggiunta, assicuratevi che il tecnico conservi uno stock di parti idonee alla vostra configurazione, e che le stesse non vengano "cannibalizzate" da altri tecnici.

8.1.1.2.3. *Disponibilità delle parti di ricambio*

Ovviamente, la disponibilità delle parti di ricambio ricopre un ruolo molto importante nel limitare l'esposizione ad eventuali problemi del vostro hardware. All'interno di un contesto di un agreement, la disponibilità delle suddette parti ricopre una certa importanza, in quanto tale disponibilità non è fondamentale solo per il vostro ufficio, ma anche per i clienti presenti nel territorio che ne fanno uso. Il problema della disponibilità potrebbe intaccare negativamente la vostra organizzazione, in quanto si potrebbe verificare che una organizzazione a voi concorrente possa aver acquistato un maggior numero di parti hardware, e potrebbe avere un trattamento migliore quando si ripresenta la necessità di acquistare parti di ricambio (oppure in termini di tecnici).

Sfortunatamente, non c'è molto che si possa fare in queste circostanze, l'unica alternativa è di risolvere il problema con il service manager.

8.1.1.2.4. *Budget disponibile*

Come sopra riportato, i contratti possono variare nel prezzo a seconda della natura dei servizi forniti. Tenete presente che i costi associati con un contratto rappresentano una spesa periodica; ogni talvolta un contratto stà per scadere, è necessario negoziare un nuovo contratto e quindi pagarlo.

8.1.1.2.5. *Hardware ricoperto dal contratto*

Ecco un'area dove i costi possono essere mantenuti al minimo. Considerate per un istante di aver appena raggiunto un accordo che vi permette di avere un tecnico interno 24 ore su 24, sette giorni su sette, con la presenza di pezzi di ricambio interni — e voi lo nominate. Ogni singola parte hardware che avete acquistato è coperta dal servizio, incluso il PC che la ricezione utilizza per i suoi compiti non critici.

In questo caso potreste domandarvi se il suddetto PC abbia *effettivamente* bisogno di un tecnico interno 24 ore su 24 sette giorni su sette. Anche se il PC risulta essere vitale per il lavoro di ricezione, l'impiegato preposto a tale compito lavora solo dalle 09:00 alle 17:00; quindi sarà molto difficile che:

- Il PC verrà utilizzato dalle 17:00 alle 09:00 del giorno seguente (senza considerare i fine settimana)
- Verrà notata la presenza di un problema su questo PC, tra le 09:00 e le 17:00

Per questo motivo risulta inutile l'assistenza del suddetto PC durante il fine settimana, risulterebbe essere una perdita di denaro.

La cosa migliore da fare risulta quella di dividere in due l'agreement sul servizio, in modo tale che un hardware adibito per compiti non critici, venga raggruppato separatamente da quello adibito a compiti più importanti. In questo modo è possibile contenere il più possibile i costi.

**Nota Bene**

Se avete venti server configurati in modo simile e che risultano essere critici per la vostra organizzazione, potrebbe risultare conveniente avere un agreement solo per uno o due di essi, con il resto coperto da un agreement meno oneroso. In questo modo non importa quale server presenta un problema durante il fine settimana, voi sarete sempre in grado di dire che il server in questione è *quello* ricoperto dal servizio più costoso.

Non fate quanto sopra riportato. Non solo perchè risulta essere disonesto, ma anche perchè molti rivenditori mantengono traccia dell'agreement stipulato, implementando numeri seriali. Anche se riuscite a trovare un modo per aggirare questo ostacolo, risulta essere molto più oneroso per voi affrontarne le spese una volta scoperti, che invece pagare regolarmente il servizio desiderato.

8.1.2. Problemi software

I problemi software possono risultare in un downtime esteso. Per esempio, i possessori di un determinato tipo di sistemi computerizzati conosciuti per le loro caratteristiche di alta disponibilità, hanno recentemente riscontrato un problema. Un bug presente nel codice di gestione dell'orario del sistema operativo del computer, ha come risultato il crashing del sistema stesso in un determinato istante, e in un determinato giorno. Anche se questa situazione in particolare potrebbe risultare un esempio un pò estremo, altre problematiche relative al software potrebbero essere meno drammatiche ma allo stesso tempo devastanti.

Le problematiche relative al software possono colpire una delle due aree:

- Sistemi operativi
- Applicazioni

Ogni tipo di problematica presenta uno specifico impatto e viene affrontata in dettaglio nelle seguenti sezioni:

8.1.2.1. Problematiche relative al sistema operativo

Con questo tipo di problema, il sistema operativo è responsabile per l'interruzione del sistema. Le problematiche del sistema operativo hanno le loro origini da due diverse aree:

- Crash del sistema
- Sospensione del sistema

La cosa principale da ricordare delle problematiche riguardanti il sistema operativo è che le suddette problematiche possono far perdere tutto ciò che era in esecuzione durante il verificarsi di un problema. Per questo motivo essi possono avere un impatto devastante nella catena produttiva di una organizzazione.

8.1.2.1.1. Crash del sistema

Il crash del sistema operativo si verifica quando si è in presenza di un errore attraverso il quale il sistema non è in grado di ripristinare le sue funzioni. I motivi di questi crash variano e vanno dalla incapacità di gestire un problema hardware, alla presenza di un bug nel codice al livello del kernel comprendente il sistema operativo. Quando si verifica un crash del sistema operativo, il sistema deve essere riavviato per continuare la sua funzione.

8.1.2.1.2. Sospensione del sistema

Quando il sistema operativo non gestisce più gli eventi del sistema, il sistema stesso si arresta. Questa procedura è conosciuta come *hang*. Tali sospensioni possono essere causate da *deadlocks* (due utenti che si contendono reciprocamente le loro risorse), e *livelocks* (due o più processi che rispondono alle rispettive attività, senza eseguire alcun lavoro utile), ma il risultato finale è lo stesso — carenza di produttività.

8.1.2.2. Problemi con le applicazioni

Diversamente dai problemi relativi ai sistemi operativi, quelli inerenti alle applicazioni possono essere più limitate. A seconda dell'applicazione specifica, il verificarsi di un problema ad una singola applicazione potrebbe influenzare solo una persona. D'altra canto se il suddetto problema riguarda un'applicazione server capace di servire un gran numero di applicazioni client, le conseguenze saranno più di larga scala.

Il problema riguardante l'applicazione, come quello riguardante il sistema operativo, è che essa può essere soggetta a sospensioni o ad un crash; la sola differenza è che ad essere sospesa o essere soggetta ad un crash è in questo caso è solo l'applicazione.

8.1.2.3. Come ottenere il supporto — Supporto software

Proprio come i rivenditori hardware forniscono un supporto per i loro prodotti, anche i rivenditori software rendono disponibili per i loro clienti dei pacchetti di supporto. Eccetto le ovvie differenze (nessun ricambio hardware è necessario, e la maggior parte del lavoro può essere svolto attraverso il telefono), i contratti di supporto software possono essere molto simili a quelli hardware.

Il livello di supporto fornito da un rivenditore software può variare. Ecco alcune delle strategie più comuni usate al giorno d'oggi:

- Documentazione
- Supporto autonomo
- Supporto email o web
- Supporto telefonico
- Supporto On-site

Ogni tipo di supporto viene analizzato nelle seguenti sezioni.

8.1.2.3.1. Documentazione

Anche se spesso trascurata, la documentazione software rappresenta il primo tool di supporto. Sia online che stampata, la documentazione spesso contiene le informazioni utili necessarie per risolvere i problemi.

8.1.2.3.2. Supporto autonomo

Il supporto autonomo rappresenta quel tipo di supporto dove il cliente, usando le risorse online, risolve i problemi relativi al proprio software. Molto spesso queste risorse prendono la forma di FAQ basate sul web (domande frequenti), o basate sulla conoscenza personale.

Le FAQ molto spesso non hanno alcuna capacità di selezione, condizionando il cliente a controllare ogni singola domanda fino a quando non si trova quella giusta. Le risorse basate sulla conoscenza

sono spesso più complesse, permettendo la ricerca di determinati termini. Le suddette risorse possono essere molto varie, rendendole uno dei migliori tool per la risoluzione dei problemi.

8.1.2.3.3. Supporto email o web

Molto spesso ciò che sembra un sito web di supporto autonomo, include anche indirizzi email e forme di supporto web i quali rendono possibile l'invio di domande al personale di supporto. Anche se il suddetto metodo potrebbe sembrare un miglioramento di tale supporto, tutto dipende però dall'efficienza del personale preposto alle risposte.

Se il suddetto personale risulta superficiale nelle risposte, questo potrebbe essere causato dalle esigenze di rispondere ad ogni singola richiesta molto velocemente in modo da rispondere al maggior numero possibile di email. La ragione sta nel fatto che il personale facente parte di questi gruppi di supporto, vengono valutati in base al numero di email che hanno risposto. Risulta essere quindi molto difficile anche il seguire in modo più approfondito il problema di una persona, in quanto — la persona preposta alla risposta della vostra email non farà altro che pensare a rispondere il più velocemente possibile in modo da concentrarsi sull'email successiva.

Il miglior modo per far fronte a tale problema è quello di accertarsi che la vostra email contenga domande specifiche, come ad esempio:

- Descrivere chiaramente la natura del problema
- Includere tutti i possibili numeri della versione
- Descrivere le vostre azioni atte alla risoluzione del problema (applicazione delle ultimissime patch, riavvio con una configurazione minima, ecc.)

Dando al tecnico più informazioni possibili avrete una maggiore possibilità di ottenere una risposta esauriente.

8.1.2.3.4. Supporto telefonico

Come indicato dal nome, il supporto telefonico vi dà la possibilità di parlare direttamente con un tecnico. Questo tipo di supporto è molto simile al supporto hardware, disponibile anche con diversi livelli di supporto (con diversi orari di copertura, di tempi di risposta, ecc.)

8.1.2.3.5. Supporto On-site

Conosciuto anche come consultazione on-site, tale supporto viene riservato per la risoluzione di problemi specifici o l'effettuazione di modifiche critiche, come ad esempio l'installazione iniziale del software, la configurazione, e così via. Come previsto, esso rappresenta il supporto software più costoso.

Vi sono alcuni casi dove il supporto on-site risulta essere quello più idoneo. Per esempio, considerate una piccola organizzazione che possiede solo un amministratore. La suddetta organizzazione sta per impiegare il primo server database, ma l'utilizzo (e l'organizzazione) non sono sufficientemente grandi da poter giustificare l'assunzione di un amministratore per la gestione del database stesso. In questa situazione, risulterebbe essere meno costoso impiegare uno specialista proveniente direttamente dal rivenditore, e quindi gestire tutte le problematiche dovute all'impiego iniziale (e occasionalmente in una fase futura, se necessario), che istruire un amministratore che userà le nozioni appena apprese in modo molto saltuario.

8.1.3. Problematiche riguardanti l'ambiente

Anche se l'hardware funziona correttamente e il software è stato configurato in modo corretto, è sempre possibile riscontrare alcuni problemi. I problemi più comuni che si possono verificare esternamente al sistema, sono quelli che implicano l'ambiente fisico nel quale risiede il sistema stesso.

I problemi riguardanti l'ambiente possono essere suddivisi in quattro categorie:

- Integrità dell'edificio
- Elettricità
- Aria condizionata
- Condizioni atmosferiche ed ambiente esterno

8.1.3.1. Integrità dell'edificio

Per una struttura così semplice, un edificio è in grado di fornire diverse funzioni. Esso fornisce una copertura, fornisce il micro-clima ideale per gli oggetti e le persone presenti al suo interno, possiede meccanismi in grado di fornire alimentazione e protezione contro il fuoco, il furto e contro atti vandalici. Garantendo tutto ciò, non deve destare sorpresa se la scelta dell'edificio più idoneo ricopre una certa importanza. Ecco alcune considerazioni:

- Ci possono essere delle infiltrazioni d'acqua, permettendo alla stessa di arrivare fino al centro dati.
- Alcuni edifici possono risultare non idonei (problemi inerenti l'acqua, il sistema fognario o il riciclo di aria), rendendo così l'edificio non utilizzabile.
- I pavimenti possono avere un limite di sopportazione del carico non sufficiente per poter ospitare l'equipaggiamento da voi scelto per far parte del centro dati.

È importante avere una mentalità molto elastica nel riconoscere i diversi motivi per i quali un edificio potrebbe risultare non idoneo. Ecco un elenco che vi aiuterà a determinare le diverse cause.

8.1.3.2. Elettricità

Poichè l'elettricità rappresenta la linfa vitale per ogni sistema operativo, le problematiche relative all'alimentazione sono di vitale importanza. Ci sono diversi aspetti da considerare relativi all'alimentazione; essi vengono affrontati in modo più dettagliato nelle seguenti sezioni.

8.1.3.2.1. La sicurezza della vostra alimentazione

Come prima cosa, è necessario determinare quanto possa essere sicura la vostra fonte normale di elettricità. Proprio come ogni centro dati, la vostra fonte di elettricità viene rappresentata da una compagnia locale in grado di fornire elettricità tramite cavi elettrici. Per questo motivo, ci sono alcuni limiti su quello che potete fare per assicurarvi che la vostra fonte di energia sia sicura.



Suggerimento

Le organizzazioni situate vicino ai confini di un'azienda in grado di fornire elettricità, potrebbero essere in grado di negoziare i propri collegamenti elettrici in modo seguente:

- Con l'azienda in grado di servire la vostra area
- E con l'azienda in grado di fornire elettricità in un'area diversa dalla vostra ma confinante

I costi necessari per l'uso di cavi elettrici provenienti da un'azienda fornitrice di energia elettrica in una area diversa ma confinante, sono molto elevati rendendo disponibile questa opzione solo per le

organizzazioni più grandi. Tuttavia, alcune organizzazioni in molti casi, possono trovare che i benefici per una tale opzione possano giustificare i costi.

È consigliabile sempre controllare i metodi attraverso i quali l'alimentazione viene resa disponibile alla vostra organizzazione e al vostro edificio in generale. Per questo motivo controllate se i cavi sono sotterranei oppure utilizzano tralicci elettrici. Quest'ultimi sono soggetti a:

- Danni causati per colpa di condizioni atmosferiche avverse (ghiaccio, vento, temporali)
- Incidenti stradali che possono danneggiare i tralicci e/o i trasformatori
- Presenza di animali che possono compromettere il normale lavoro dei cavi

Tuttavia, anche i cavi sotterranei presentano i loro lati negativi:

- Danni dovuti ai lavori di scavatura
- Inondazioni
- Temporali (anche se in maniera più lieve)

Continuando a seguire il percorso dei cavi, controllate se questi passano, prima di arrivare all'interno del vostro edificio, attraverso un trasformatore esterno. Controllate anche se il trasformatore sia protetto contro incidenti dovuti a veicoli o alla caduta di alberi. Inoltre controllate se gli interruttori siano protetti per evitare un uso non autorizzato.

Una volta all'interno dell'edificio, controllate se i cavi elettrici (o i pannelli ai quali essi sono collegati), possano essere soggetti ad altri problemi. Per esempio, controllate se un problema di natura idraulica possa presentare un problema per i cavi.

Continuando a seguire i cavi fino a raggiungere il centro dati, controllate se vi sono altri ostacoli che possano intaccare il normale approvvigionamento di energia elettrica. Per esempio, controllare se il centro dati condivide uno o più circuiti con i carichi che non fanno parte del centro dati stesso, se così fosse, è possibile che si possa verificare un sovraccarico di corrente, azionando così i circuiti di protezione e intaccando così il normale funzionamento del centro dati.

8.1.3.2.2. Qualità di alimentazione

Non è sufficiente assicurarsi che la fonte di alimentazione del centro dati sia la più sicura possibile. È importante anche assicurarsi della qualità di alimentazione distribuita all'intero del centro dati. Per questo motivo ci sono diversi fattori da considerare:

Voltaggio

Il voltaggio della corrente elettrica deve essere stabile, con nessuna riduzione (spesso riferita come *perdite*, *cadute*, o *sottotensioni*) o aumenti (spesso conosciuti come *picchi* e *sovratensioni*).

Forma dell'onda elettrica

La forma dell'onda elettrica deve essere sinusoidale, con un valore minimo di *THD* (Total Harmonic Distortion).

Frequenza

La frequenza deve essere stabile (molte nazioni usano una frequenza di alimentazione di 50Hz o 60Hz).

Rumore

L'alimentazione non deve includere qualsiasi rumore di tipo *RFI* (Radio Frequency Interference) o *EMI* (Electro-Magnetic Interference)

Corrente

L'alimentazione deve essere fornita con una corrente sufficiente per far funzionare il centro dati.

L'alimentazione generalmente fornita dall'azienda produttrice di energia elettrica, normalmente non è idonea agli standard necessari per un centro dati. Per questo motivo, sono necessari alcuni tipi di adattatori di alimentazione. Sono presenti diversi approcci:

Scaricatori modulari

Gli scaricatori modulari — filtrano i flussi provenienti dalla sorgente di alimentazione. La maggior parte di essi non fanno altro, lasciando così l'apparecchiatura vulnerabile ad altri problemi.

Adattatori di alimentazione

Gli adattatori di alimentazione hanno un approccio più completo; a seconda della natura dell'unità, essi possono far fronte alla maggior parte dei problemi sopra indicati.

Unità Motore-Generatore

Una unità motore-generatore è composta da un grande motore elettrico alimentato dalla vostra fonte di alimentazione normale. Il motore è collegato ad una ruota rotante molto grande, la quale è collegata ad un generatore. In questo modo, il motore è in grado di far girare la ruota, facendo produrre al generatore una quantità di elettricità sufficiente per alimentare il centro dati. In questo modo, l'alimentazione del centro dati è isolata da quella esterna, ciò comporta quindi una riduzione di numerosi problemi. La ruota rotante fornisce anche la possibilità di mantenere costante l'alimentazione anche in presenza di cadute di tensione in quanto la ruota rotante ha bisogno di alcuni secondi prima di fermarsi.

Alimentazione costante 'Uninterruptible Power Supplies'

Alcuni tipi di Uninterruptible Power Supplies (più comunemente conosciuti come *UPS*), includono alcune (se non tutte) protezioni di un adattatore di alimentazione ².

Con le ultime due tecnologie sopra elencate, iniziamo ad affrontare un argomento molto importante che coinvolge un gran numero di persone — l'alimentazione di backup. Nella sezione successiva, vengono affrontati i diversi approcci per fornire un'alimentazione di backup.

8.1.3.2.3. Alimentazione di backup

Un termine relativo all'alimentazione che quasi tutti conoscono è *blackout*. Un *blackout* non è altro che una perdita di corrente che può andare da una frazione di secondo a delle settimane.

Poiché la durata di questi blackout è variabile, è necessario implementare una forma di alimentazione di backup, in modo da far fronte a qualsiasi imprevisto.



Suggerimento

Il blackout medio dura circa pochi secondi; quelli di durata superiore sono più rari. Per questo motivo, proteggetevi contro blackout che hanno una durata breve, e successivamente provate a risolvere il problema di come far fronte ad un blackout più lungo.

2. La tecnologia riguardante gli UPS viene affrontata in modo più dettagliato nella Sezione 8.1.3.2.3.2.

8.1.3.2.3.1. Fornire alimentazione per pochi secondi

Poichè la maggior parte dei blackout durano solo pochi secondi, la soluzione adottata per fornire alimentazione di backup deve comprendere due caratteristiche principali:

- Un tempo relativamente breve per smistarsi sull'alimentazione di backup (conosciuto come *tempo di trasferimento*)
- Un *periodo di esecuzione* (il periodo di durata del backup), misurato in secondi e minuti

Le soluzioni per l'alimentazione di backup che corrispondono a queste caratteristiche sono gli insiemi motore-generatore e gli UPS. La ruota rotante presente nell'insieme motore-generatore, permette di far fronte a blackout di breve durata, circa un secondo. Essi generalmente sono molto grandi e costosi, e ideali per centro dati medio-grandi.

Tuttavia, un altro tipo di tecnologia — chiamata UPS — può sostituire la soluzione rappresentata da un insieme motore-generatore in quanto troppo costosa. L'UPS è in grado di far fronte a blackout più lunghi.

8.1.3.2.3.2. Fornire alimentazione per pochi minuti

Gli UPS possono essere acquistati in diverse dimensioni — abbastanza piccoli per eseguire un PC low-end piccolo per cinque minuti, o sufficientemente grande per alimentare un intero centro dati per più di una ora.

Gli UPS sono composti dalle seguenti parti:

- Un *interruttore di smistamento*, per smistarsi dalla fonte di alimentazione principale a quella di backup
- Una batteria, per fornire alimentazione di backup
- Un *invertitore*, il quale converte la corrente DC continua della batteria in corrente AC alternata necessaria per l'hardware del centro dati

Non tenendo in considerazione la misura e la capacità della batteria dell'unità, gli UPS sono principalmente di due tipi

- L'UPS *offline* utilizza il proprio invertitore per generare alimentazione solo quando la fonte primaria non è in grado di farlo
- L'UPS *online* utilizza il proprio invertitore per generare continuamente alimentazione, alimentando l'invertitore tramite la sua batteria solo quando la fonte di alimentazione principale non è in grado di farlo.

Ogni tipo presenta i propri vantaggi e svantaggi. L'UPS online è generalmente meno costoso, in quanto l'invertitore non deve essere impostato per un lavoro costante. Tuttavia, se si verifica un problema in un invertitore di un UPS offline, esso non verrà notato (almeno fino a quando non si verifica un problema all'alimentazione).

Gli UPS online tendono ad essere migliori nel fornire un'alimentazione pulita al vostro centro dati; esso fornisce essenzialmente alimentazione in modo costante.

Non ha importanza quale tipo di UPS selezionate, è necessario configurarlo in modo da essere idoneo al carico (e quindi assicuratevi che l'UPS abbia una capacità sufficiente per produrre elettricità come richiesto, e cioè con un corretto voltaggio e corrente), e determinare la durata entro la quale deve operare per essere in grado di fornire alimentazione tramite la batteria.

Per ottenere tutte queste informazioni, è necessario determinare prima il carico che l'UPS deve fornire. Controllate ogni singolo pezzo e determinate la quantità di alimentazione che è possibile fornire (queste informazioni sono generalmente riportate su di una etichetta). Annotate il voltaggio, i watt,

e/o ampere. Una volta ottenute tutte queste informazioni, convertitele in VA (Volt-Ampere). Se avete un numero di watt, potete usare i watt elencati come VA; se avete gli ampere, moltiplicateli per i volt per ottenere i VA. Aggiungendo i risultati di VA, potete ottenere approssimativamente il numero di VA necessari per il vostro UPS.



Nota Bene

Questo tipo di calcolo per la determinazione di VA non è proprio corretto; tuttavia, per ottenere il vero valore di VA dovrete sapere il fattore di alimentazione per ogni unità, ma purtroppo questa informazione viene raramente fornita. In ogni modo, i numeri riguardanti il VA ottenuti in questo modo, riflettono valori molto approssimativi, lasciando così un ampio margine d'errore.

La determinazione del periodo di esecuzione si traduce più in un aspetto aziendale che in un aspetto tecnico — contro che tipo di disfunzioni desiderate proteggervi, e quanto desiderate spendere? Molti scelgono periodi di esecuzione che si aggirano ad una ora, in alcuni casi anche a due ore, in quanto oltre a questi limiti, le batterie che forniscono questa copertura sono molto costose.

8.1.3.2.3.3. Fornire alimentazione per alcune ore (ed oltre)

Quando si hanno delle disfunzioni nell'alimentazione che si protraggono per giorni, le scelte da fare sono molto costose. Le tecnologie in grado di far fronte ad un simile problema sono limitate a generatori alimentati da motori — principalmente turbine a gas o diesel.



Nota Bene

Ricordatevi che i generatori alimentati da motori necessitano di carburante durante la loro esecuzione. È necessario quindi conoscere il "consumo" di carburante del vostro generatore con carico massimo, e conseguentemente organizzare la consegna di carburante.

A questo punto le vostre opzioni sono svariate, pur assumendo che la vostra organizzazione abbia a disposizione sufficienti risorse finanziarie. In questo caso è consigliabile la presenza di esperti, in modo da potervi aiutare nella determinazione della soluzione migliore per la vostra organizzazione. Pochissimi amministratori di sistema possiedono la conoscenza necessaria per pianificare l'acquisizione e l'impiego di questi sistemi.



Suggerimento

È possibile affittare generatori portatili di qualsiasi dimensione, sfruttando i loro benefici senza far fronte alle spese iniziali necessarie per il loro acquisto. Tuttavia, ricordatevi che a causa dei numerosissimi problemi, spesso vi è una carenza di generatori, e quindi il loro affitto potrebbe risultare molto oneroso.

8.1.3.2.4. Pianificazione di un blackout esteso

Un blackout di cinque minuti può sembrare un piccolo inconveniente durante una normale giornata lavorativa, ma quali conseguenze può apportare un blackout di una ora? Cinque ore? Un giorno? Una settimana?

Anche se un centro dati funziona regolarmente, un blackout esteso potrebbe in qualche modo influenzare la vostra organizzazione. Considerate i seguenti punti:

- Cosa può succedere se non vi è più alimentazione per gestire il controllo dell'ambiente in un centro dati?
- Cosa può succedere se non vi è più alimentazione per gestire il controllo dell'ambiente di tutto l'edificio?
- Cosa può succedere se non vi è più alimentazione per far funzionare le workstation personali, il sistema telefonico e le luci?

Il punto da fare in questo caso è il seguente, fino a quando la vostra organizzazione è in grado di tollerare un blackout esteso. Oppure, se questo punto non rappresenta una opzione, la vostra organizzazione deve considerare la possibilità di continuare ad operare in ogni circostanza, adottando generatori molto grandi per alimentare tutto l'edificio.

Naturalmente è possibile che la causa che abbia determinato un blackout esteso nei confronti della vostra organizzazione, possa aver influenzato anche il mondo esterno, causando un serio problema nell'abilità della vostra organizzazione a continuare le normali funzioni, anche se la stessa possiede una capacità illimitata per quanto riguarda i generatori di alimentazione.

8.1.3.3. Riscaldamento, ventilazione, e aria condizionata

I sistemi di riscaldamento, di ventilazione e di aria condizionata (*HVAC*) utilizzati negli edifici moderni sono molto sofisticati. Spesso sono controllati da computer, il sistema HVAC è indispensabile per fornire un ambiente lavorativo molto confortevole.

Generalmente un centro dati possiede un sistema di aria condizionata aggiuntivo, in questo modo è in grado di affrontare il problema dovuto alle alte temperature generate dai computer e dalle apparecchiature associate. La presenza di problemi in un sistema HVAC può risultare devastante per il normale funzionamento di un centro dati. Data la loro complessità e la natura elettro-meccanica, le possibilità che questo avvenga sono molto elevate ed al tempo stesso molto varie. Ecco alcuni esempi:

- Le unità che gestiscono l'aria condizionata (generalmente composte da ventilatori molto grandi con motori elettrici) possono presentare alcuni problemi di natura elettrica, di portante, di cinghia/puleggia ecc.
- Le unità di ventilazione (chiamate spesso *chillers*), possono perdere le loro capacità a causa di perdite, o a causa del bloccaggio del loro compressore e/o motore.

La riparazione e la gestione dei sistemi HVAC rappresenta un campo molto specializzato — un campo che un normale amministratore di sistema dovrebbe lasciare agli esperti. L'unica cosa che egli può garantire, è quella di assicurare l'esecuzione di un controllo giornaliero ai suddetti sistemi (considerate anche un controllo più frequente), e che gli stessi siano mantenuti in accordo alle indicazioni del rivenditore.

8.1.3.4. Condizioni atmosferiche e fattori esterni

Talvolta si possono verificare particolari condizioni atmosferiche che possono causare problemi ad un amministratore:

- Abbondanti nevicate e formazione di ghiaccio possono impedire al personale di arrivare al centro dati, è possibile che queste condizioni siano in grado di ostruire i condensatori dell'aria condizionata, causando un aumento delle temperature proprio quando il personale non è presente e quindi non è in grado di apportare le correzioni necessarie.
- Forti raffiche di vento possono danneggiare le comunicazioni e l'alimentazione, in alcuni casi danneggiando anche le strutture dell'edificio.

Sono presenti altri tipi di condizioni atmosferiche in grado di causare diversi problemi, anche se questi non sono molto conosciuti. Per esempio, temperature troppo elevate possono causare problemi ai sistemi di raffreddamento, creando problemi di alimentazione a causa del sovraccarico della rete elettrica.

Anche se non è possibile influenzare le condizioni atmosferiche, sapere il modo con il quale queste condizioni agiscono sul normale funzionamento del vostro centro dati, può risultare utile per mantenere le vostre strutture sempre efficienti.

8.1.4. Errori umani

È stato detto che i computer *sono* perfetti. Il motivo di questa affermazione risiede nel fatto che se si effettua un controllo approfondito, ne risulterà che la causa principale sia un errore umano. In questa sezione, vengono esaminate le cause e gli impatti causati dai diversi errori umani.

8.1.4.1. Errori causati da un utente finale

Gli utenti di computer possono commettere errori gravi in grado di avere un serio impatto sull'intero sistema. Tuttavia, a causa del loro ambiente non privilegiato, i suddetti errori sono generalmente di natura locale. Poiché molti utenti interagiscono con un computer attraverso una o più applicazioni, è proprio in esse che si possono ricercare gli errori dovuti ad utenti finali.

8.1.4.1.1. *Uso improprio delle applicazioni*

Quando le applicazioni vengono usate in modo improprio, si possono verificare diversi problemi:

- Sovrascrittura dei file
- Utilizzo di dati non corretti come input per un'applicazione
- File organizzati e chiamati in modo non corretto
- Cancellazione accidentale di file

L'elenco potrebbe protrarsi ancora, ma quanto stilato è sufficiente per farsi un'idea. Poiché gli utenti non hanno i privilegi del super utente, gli errori che essi commettono sono limitati ai propri file. Per questo motivo, ecco riportato il miglior approccio:

- Educate gli utenti ad un uso corretto delle proprie applicazioni e delle tecniche corrette di gestione dei file
- Assicuratevi di eseguire periodicamente un backup dei file degli utenti e che il processo di ripristino sia semplice e veloce

Non vi è molto altro che possiate fare per limitare al massimo gli errori degli utenti.

8.1.4.2. Errori dovuti al personale addetto alle operazioni

Gli operatori hanno un rapporto più stretto con l'organizzazione rispetto agli utenti finali. Generalmente gli errori dei suddetti utenti si traducono in errori nelle applicazioni, al contrario gli operatori tendono ad eseguire compiti molti più ampi e complessi. Anche se la natura dei compiti viene decisa da altri, alcuni dei compiti possono includere l'utilizzo di utility del tipo system-level, dov'è maggiore la possibilità di recare danni più complessi. Per questo motivo, il tipo di errore che un operatore può fare è quello di non seguire scrupolosamente le procedure.

8.1.4.2.1. Inosservanza delle procedure

Tutti gli operatori dovrebbero avere una serie di procedure documentate e disponibili per ogni azione da eseguire³. È possibile anche che le procedure non siano state aggiornate. Ecco elencati i vari motivi:

- L'ambiente lavorativo ha subito un cambiamento e le procedure non sono mai state aggiornate. Per questo motivo le procedure memorizzate precedentemente dall'operatore sono invalide. A questo punto, anche se le procedure sono state aggiornate, (molto improbabile, dato che le stesse non sono state aggiornate prima), l'operatore non ne risulterà a conoscenza.
- L'ambiente lavorativo è cambiato, ma non risulta alcuna presenza di procedure. Questo è un esempio ancora più critico di quello precedente.
- Le procedure esistono e sono corrette, ma l'operatore non le segue (o non è in grado di seguirle).

A seconda della struttura di direzione della vostra organizzazione, non vi resta altro da fare che informare il vostro manager. In ogni modo, rendervi sempre disponibili per risolvere un problema risulta essere sempre il miglior approccio.

8.1.4.2.2. Errori durante l'esecuzione delle procedure

Anche se l'operatore segue le procedure corrette è sempre possibile commettere alcuni errori. Se si verificasse un errore, ciò potrebbe significare che l'operatore non ha prestato molta attenzione, (in questo caso è necessario il coinvolgimento del manager dell'operatore in questione).

Potrebbe anche trattarsi di un semplice errore. In questi casi, gli operatori migliori sono in grado di accorgersi che il funzionamento non è corretto, chiedendo così assistenza al personale specializzato. Incoraggiate gli operatori con i quali lavorate a tale prassi, in modo tale che gli stessi contattino il personale addetto nel caso in cui essi siano in grado di accorgersi di un funzionamento non corretto. Anche se molti operatori sono in grado di risolvere la maggior parte dei problemi in modo indipendente, è da tenere in considerazione il fatto che questo compito non rientra nelle loro mansioni. Ricordate che se il problema si aggrava a causa di un intervento di un operatore, tale aggravarsi potrebbe influire in modo negativo sia sulla carriera dell'operatore stesso, e sia sulla vostra possibilità di risolvere il problema in modo rapido.

8.1.4.3. Errori dell'amministratore del sistema

Diversamente dagli operatori, gli amministratori eseguono una gamma molto vasta di compiti utilizzando i computer dell'organizzazione. I loro compiti non si basano spesso su procedure documentate.

Per questo motivo, gli amministratori non fanno altro che aumentare il loro carico di lavoro. Durante l'adempimento dei compiti giornalieri di un amministratore, essi hanno un sufficiente accesso ai sistemi (senza indicare i loro privilegi di super utente), per creare inavvertitamente dei danni.

3. Se gli operatori nella vostra organizzazione non sono in possesso delle suddette procedure, cercate di creare un documento lavorando insieme con loro, con il vostro manager e con i vostri utenti. Senza tali procedure un centro dati ha altissime probabilità di incontrare problemi seri durante le operazioni giornaliere.

Gli amministratori generalmente commettono errori dovuti a configurazioni errate o di gestione

8.1.4.3.1. Errori nella configurazione

Gli amministratori devono sempre configurare tutti i diversi aspetti di un sistema. Questa configurazione potrebbe includere:

- Email
- Account utente
- Rete
- Applicazioni

L'elenco potrebbe andare avanti. Il compito vero e proprio della configurazione varia; alcuni compiti richiedono la modifica di un file di testo (utilizzando una delle tantissime sintassi dei file di configurazione disponibili), mentre altri file richiedono l'esecuzione di una utility di configurazione.

Il fatto che tutti questi compiti vengono gestiti in modo diverso non rappresenta una grossa aggiunta, infatti ogni compito riguardante la configurazione richiede conoscenze diverse. Per esempio, la conoscenza necessaria per configurare un mail transport agent è diversa da quella necessaria per configurare un nuovo collegamento di rete.

Considerando tutto questo, ci si potrebbe meravigliare del fatto che vengono commessi così *pochi* errori. In ogni modo, la configurazione è, e continuerà ad essere uno dei compiti più impegnativi di un amministratore. Ci si può domandare dunque se esista un metodo per ottenere un processo meno propenso ad errori.

8.1.4.3.1.1. Modifica del controllo

La minaccia comune per ogni modifica della configurazione è rappresentata dall'esistenza di una modifica. Il cambiamento potrebbe essere consistente, oppure piccolo, ma esso pur sempre rappresenta un cambiamento, e bisogna affrontarlo in modo particolare.

Molte organizzazioni implementano un tipo di processo per il controllo delle modifiche. L'intento è quello di aiutare gli amministratori (e tutti i gruppi influenzati a tali modifiche), in modo da gestire il processo e di ridurre l'esposizione dell'organizzazione stessa a possibili errori.

Tale processo di controllo sulle modifiche generalmente divide il processo di modifica in quattro fasi. Ecco un esempio:

Ricerca preliminare

La ricerca preliminare cerca di definire in modo chiaro quanto segue:

- La natura della modifica
- Se tale modifica avviene, definire il suo impatto
- Una posizione di ripiego se la modifica non ha avuto successo
- Una valutazione dei possibili errori

Una ricerca preliminare potrebbe includere la prova della modifica durante un periodo previsto di downtime, oppure potrebbe prima implementare la modifica su di un ambiente di prova particolare eseguito su di un hardware di prova preposto.

Programmazione

Il cambiamento viene esaminato con una attenzione particolare sul meccanismo di implementazione. L'implementazione fatta include la sequenza ed il tempo necessario per la modifica (insieme con la sequenza ed il tempo necessario ad ogni fase per eseguire il back out della modifica

nel caso in cui si verificasse un problema), ed assicurarsi che il tempo assegnato alla modifica sia sufficiente, e che non vi sia alcun conflitto con qualsiasi altra attività al livello di sistema.

Il prodotto di questo processo è rappresentato da un elenco di controllo di tutte le fasi da seguire da parte dell'amministratore durante l'esecuzione delle modifiche. Sono incluse con ogni fase le istruzioni da seguire nel caso in cui si verifichi un errore durante una delle fasi. Spesso viene incluso il tempo stimato per facilitare il controllo da parte di un amministratore.

Esecuzione

A questo punto l'esecuzione delle fasi necessarie per implementare il cambiamento dovrebbero essere semplici e chiare. Il suddetto cambiamento può essere o meno implementato, (a seconda se si fossero riscontrati dei problemi).

Controllo

Anche se la modifica non viene implementata, l'ambiente viene controllato in modo da assicurare il suo normale funzionamento.

Documentazione

Se la modifica è stata implementata, tutta la documentazione viene aggiornata in modo da riflettere la diversa configurazione.

Ovviamente, non tutte le modifiche della configurazione richiedono questo livello di dettagli. La creazione di un account utente non dovrebbe richiedere alcuna ricerca preliminare, e la programmazione consisterà nel determinare se l'amministratore sia in possesso di tempo sufficiente per creare un account. L'esecuzione sarà allo stesso modo molto veloce, il controllo potrebbe consistere nell'assicurare che l'account è utilizzabile, e la documentazione consisterà probabilmente nell'invio di una email al nuovo manager dell'utente.

Poichè le modifiche riguardanti la configurazione diventano sempre più complesse, sarà necessario l'adozione di un processo di controllo delle modifiche più formale.

8.1.4.3.2. Errori commessi durante la manutenzione

Questi tipi di errori possono essere molto insidiosi in quanto la pianificazione ed il controllo durante il mantenimento giornaliero sono ridotte al minimo.

Sfortunatamente gli amministratori sono testimoni ogni giorno di questo tipo di errore, anche se alcuni utenti affermano di non aver modificato niente — e che potrebbe essere colpa del computer. Generalmente l'utente che afferma questo, non ricorda quello che ha fatto in passato, e allo stesso modo quando un qualcosa di simile accade anche a voi, vi troverete nella stessa situazione.

Il segreto è quello di poter ricordare le modifiche eseguite durante il mantenimento, in modo tale da essere in grado di risolvere qualsiasi problema rapidamente. Un processo di controllo per le modifiche molto articolato non risulterebbe idoneo in quanto sarà difficile ricordarsi le centinaia di cose fatte durante il giorno. Che cosa si potrebbe fare quindi per ricordarsi tutto quello che un amministratore fa durante una giornata lavorativa?

La risposta è semplice — prendere nota di ogni cosa. Sia che venga eseguito utilizzando un taccuino, un PDA o come commenti nei file interessati, prendete sempre nota. Controllando ciò che avete fatto, avrete una maggiore possibilità di risalire alla causa del problema.

8.1.4.4. Errori causati dal tecnico

Talvolta le persone che dovrebbero aiutarvi a risolvere il problema e a mantenere così il vostro sistema sempre efficiente, possono rappresentare la causa per un aggravarsi del problema stesso. Tutto questo

non è premeditato; ma potrebbe essere frutto di alcune circostanze molto particolari. Lo stesso effetto lo si può avere al lavoro quando i programmatori risolvono un bug, creando a loro insaputa un altro bug.

8.1.4.4.1. Hardware riparato incorrettamente

In questo caso il tecnico potrebbe aver eseguito una diagnosi incorretta del problema e quindi aver eseguito una riparazione non necessaria, oppure si potrebbe verificare il caso di una corretta diagnosi, ma una riparazione non corretta. Potrebbe anch'essere probabile che l'unità usata per la sostituzione era difettosa, o che la procedura di riparazione non è stata seguita in modo corretto.

Per questo motivo è importante sapere sempre tutto quello che il tecnico è in procinto di fare. Facendo questo potreste risalire alla causa di un problema in modo più rapido e semplice. In questo modo il tecnico potrà seguire la cronologia di un problema con una certa logica, invece di affrontare un problema come se fosse uno nuovo senza connessione con il precedente. In questo modo si evita la perdita di tempo che potrebbe essere utilizzato per risolvere il problema sbagliato.

8.1.4.4.2. Risoluzione di un problema creandone un altro

Talvolta anche se un problema è stato diagnosticato e risolto in modo corretto, è possibile che si verifichi un secondo problema. Il modulo della CPU è stato sostituito, ma il sacchetto anti-statico entro il quale risiedeva, ostruisce la ventola causando uno shutdown a causa di temperature elevate. Oppure l'unità disco presente nel RAID array che presentava alcuni problemi è stata sostituita, poiché il connettore su di un'altra unità è stato accidentalmente scollegato, l'array non è ancora in funzione.

Questi risultati sono causati dalla poca attenzione oppure da un errore. Questo purtroppo non ha importanza. Quello che dovete fare è di prestare sempre molta attenzione nel rivisionare ciò che il tecnico ha riparato, e assicurare che il sistema funzioni in modo corretto prima di congelarlo.

8.2. Backup

I backup svolgono due compiti principali:

- Permettere il ripristino di file individuali
- Permettere il ripristino completo dei file system

Il primo compito rappresenta la base per la richiesta tipica di ripristino dei file: un utente cancella accidentalmente un file e successivamente richiede di ripristinarlo direttamente dall'ultimo backup. Le esatte circostanze potrebbero variare, ma generalmente questo rappresenta un uso giornaliero dei backup.

La seconda situazione rappresenta l'incubo di tutti gli amministratori: l'hardware in questione una volta utilizzato come parte produttiva del centro dati, ora non è altro che un corpo estraneo. La parte non reperibile è composta dal software e dai dati che voi e i vostri utenti avete messo insieme durante gli anni. Teoricamente è stato eseguito il backup di tutto il materiale perduto. La domanda da fare è la seguente: effettivamente questo backup è stato eseguito?

Se sì, potete ripristinare il tutto?

8.2.1. Dati diversi: Diversi requisiti di backup

Notate il tipo di dati⁴ processati e conservati da un computer generico. Notate che alcuni dati non sono stati quasi modificati, mentre altri variano costantemente.

La velocità con la quale vengono modificati i dati risulta essere importante per il design di una procedura di backup. Ecco elencati i motivi:

- Un backup non è altro che una rappresentazione di dati ai quali è stato eseguito il backup. Rappresenta i dati in un particolare istante.
- È possibile eseguire raramente il backup dei dati che non cambiano frequentemente, mentre è possibile eseguire un backup più frequente, per i dati che cambiano più spesso.

Gli amministratori che possiedono una buona conoscenza dei propri sistemi, dei propri utenti e delle applicazioni, dovrebbero essere in grado di raggruppare velocemente i diversi dati in categorie. Ecco alcuni esempi:

Sistema operativo

Questi dati cambiano normalmente solo durante i processi di aggiornamento, di installazione di bug fix, e qualsiasi modifiche specifiche al sito.



Suggerimento

È necessario eseguire un backup del sistema operativo? Questa è una delle domande che ha perseguitato un amministratore durante gli anni. D'altro canto se un processo di installazione è relativamente semplice, e se le applicazioni del bug fix e le personalizzazioni sono documentate correttamente e facilmente riproducibili, l'opzione migliore potrebbe essere quella di installare nuovamente il sistema operativo.

Se invece non vi sono dubbi che una nuova installazione è in grado di riprodurre l'ambiente originale del sistema, allora è meglio eseguire un backup del sistema operativo, anche se tale procedimento viene eseguito meno frequentemente dei backup dei dati di produzione. I backup occasionali del sistema operativo possono essere d'aiuto solo quando vi è la necessità di ripristinare alcuni file system (per esempio, a causa di cancellazioni accidentali).

Application Software

Questi dati cambiano ogni qualvolta vengono installate, rimosse o migliorate delle applicazioni.

Application Data

Questi dati cambiano ogni qualvolta vengono eseguite le applicazioni associate. A seconda dell'applicazione specifica e della vostra organizzazione, questo può significare che i cambiamenti avvengono secondo secondo oppure dopo la fine di ogni anno fiscale.

User Data

Questi dati cambiano a seconda dei modelli in uso nella vostra community. In molte organizzazioni questo significa che i cambiamenti avvengono in ogni momento.

In base a queste categorie (oppure ad altre categorie aggiuntive che sono specifiche alla vostra organizzazione), dovrete avere una buona idea riguardo la natura dei backup necessari per proteggere i vostri dati.

4. Usiamo in questa sezione il termine *dati* per descrivere qualsiasi cosa che sia stata processata tramite il software di backup. Ciò include il software del sistema operativo, application software, ed anche i dati effettivi. Non ha importanza cosa sia, per quanto riguarda il software, sono tutti dei dati.

**Nota Bene**

Ricordate che la maggior parte dei software di backup interagiscono con i dati, ad un livello di file system o di directory. In altre parole, la struttura della directory del vostro sistema risulta essere importante su come vengono eseguiti i backup. Questo è un altro motivo per affermare che è sempre una buona idea considerare attentamente la migliore struttura della directory per un nuovo sistema, group file e directory a seconda del loro utilizzo previsto.

8.2.2. Software di backup: Acquistare o Creare

Per poter eseguire i backup è innanzitutto necessario avere il software corretto. Tale software non deve solo essere in grado di eseguire i compiti di base come per esempio eseguire copie di bit su media di backup, ma deve rappresentare una interfaccia idonea per il personale e per i requisiti della vostra organizzazione. Alcune delle caratteristiche da considerare sono:

- Programmazione dei backup da eseguire nel momento appropriato
- Gestione della posizione, della rotazione e dell'utilizzo del media di backup
- Coadiuvare gli operatori (e/o con i modificatori robotizzati dei media) per assicurare la disponibilità del media idoneo.
- Assistenza degli operatori nell'individuazione del media contenente un backup specifico di un determinato file

Come potete notare, una soluzione effettiva di backup richiede molto di più.

A questo punto molti amministratori vanno alla ricerca di una delle seguenti soluzioni:

- Acquisto di una soluzione già sviluppata
- Creazione di un sistema di backup (possibilmente integrando una o più tecnologie open source)

Ogni singolo approccio possiede un punto a favore ed uno a sfavore. Data la complessità dei compiti, una soluzione interna potrebbe non essere in grado di gestire al meglio alcuni aspetti (come ad esempio la gestione dei media, oppure avere una documentazione ed un supporto tecnico non completi). Tuttavia, per alcune organizzazioni quanto sopra descritto potrebbe essere una limitazione.

Una soluzione sviluppata commercialmente potrebbe essere altamente funzionale, ma allo stesso tempo potrebbe essere molto complessa per alcuni requisiti di una organizzazione. Detto questo, tale complessità potrebbe permettere di utilizzare sempre la stessa soluzione anche quando l'organizzazione va incontro ad alcuni cambiamenti.

Come potete notare, non vi è un metodo preciso per la scelta di un sistema di backup. L'unico avviso che possiamo offrirvi è quello di considerare i seguenti punti:

- La sostituzione del software di backup è molto complessa; una volta implementato, dovrete usare il suddetto software per molto tempo. Dopo tutto, avrete dei backup dell'archivio per un periodo molto lungo che dovrete essere in grado di leggere. Cambiare il software di backup significherebbe mantenere il software originale (per accedere ai backup dell'archivio), oppure convertire i backup del vostro archivio in modo da essere compatibile con il nuovo software.

A seconda del software di backup, l'impegno necessario per convertire i backup di un archivio potrebbe essere uguale (ma con una perdita di tempo considerevole) all'impegno necessario per eseguire i backup tramite un programma di conversione già esistente, oppure potrebbe richiedere l'esecuzione di una procedura di 'reverse-engineering' del formato di backup, creando un software personale per eseguire il compito.

- Il software deve essere affidabile al 100% — deve eseguire il backup di tutto ciò che è previsto, e quando gli è previsto.
- Quando bisogna eseguire il ripristino dei dati — sia di un singolo file che di un file intero — il software di backup deve essere affidabile al 100%.

8.2.3. Tipi di backup

Se domandate ad una persona non molto esperta sui backup, egli potrebbe considerare il backup come una copia identica di *tutti* i dati presenti sul computer. In altre parole, se è stato eseguito un backup martedì sera, e niente è stato modificato mercoledì, il backup eseguito mercoledì sera sarà uguale a quello eseguito martedì.

Anche se è possibile configurare il backup in questo modo, è molto probabile che voi non lo userete. Per poter ottenere più informazioni, bisogna prima capire i diversi tipi di backup che possono essere creati. Essi sono:

- Backup completi
- Backup progressivi
- Backup differenziali

8.2.3.1. Backup completi

Il tipo di backup affrontato all'inizio di questa sezione è conosciuto come *backup completo*. Un backup completo è quel backup dove ogni singolo file viene scritto sul media di backup. Come sopra riscontrato, se viene eseguito il backup di dati non modificati, ogni backup completo sarà uguale.

Questa somiglianza è dovuta al fatto che un backup completo non controlla se un file è stato modificato dall'ultimo backup; ma scrive ogni cosa sul media di backup senza tener in considerazione se il file è stato o meno modificato.

Ecco perchè il backup completo non viene sempre eseguito — ogni file viene scritto sui media di backup. Questo significa che vengono utilizzati molti media di backup anche quando niente è stato cambiato. Eseguire il backup di 100 gigabytes di dati ogni volta che in realtà sono stati modificati 10 megabytes non rappresenta un approccio corretto, ecco perchè i *backup progressivi* sono stati creati.

8.2.3.2. Backup progressivi

Diversamente dai backup completi, quelli progressivi controllano prima se l'orario di modifica del file sia più recente rispetto all'ultimo orario di backup. Se questo non è il caso, il file non è stato modificato dall'ultimo backup e quindi può essere saltato. Se la data della modifica è più recente rispetto alla data dell'ultimo backup, il file è stato modificato e quindi si può eseguire il suo backup.

I backup progressivi sono usati insieme con i backup completi (per esempio, un backup settimanale completo, insieme con backup progressivi giornalieri).

Il vantaggio primario riguardante l'uso dei backup progressivi è rappresentato dal fatto che questi ultimi sono più veloci dei backup completi. Mentre lo svantaggio primario è quello che il ripristino di ogni dato file potrebbe significare eseguire uno o più backup progressivi fino a che non viene trovato il file in questione. Quando si esegue il ripristino di un file system completo, è necessario ripristinare l'ultimo backup completo ed ogni successivo backup progressivo.

Per evitare il bisogno di adare attraverso ogni backup progressivo, è stato implementato un approccio leggermente diverso. Questo approccio è conosciuto come *backup differenziale*.

8.2.3.3. Backup differenziali

I backup differenziali sono simili a quelli progressivi e cioè essi modificano solo i file. Tuttavia, i backup differenziali sono *cumulativi* — in altre parole, con un backup differenziale, una volta che un file viene modificato esso continua ad essere incluso in tutti i backup differenziali successivi (fino ovviamente al successivo backup completo).

Ciò significa che ogni backup differenziale contiene tutti i file modificati fino all'ultimo backup completo, rendendo possibile l'esecuzione di un ripristino completo con solo l'ultimo backup completo e l'ultimo backup differenziale.

Come la strategia di backup utilizzata con i backup progressivi, normalmente i backup differenziali seguono lo stesso approccio: Un singolo backup periodico completo seguito da backup differenziali più frequenti.

L'effetto nell'uso in questo senso dei backup differenziali, è rappresentato dal fatto che i suddetti backup tendono ad aumentare attraverso il tempo (assumendo il fatto che diversi file possono essere stati modificati col tempo dall'ultimo backup). Questo posiziona i backup differenziali in una posizione compresa tra i backup progressivi e quelli completi, in termine di utilizzo dei media di backup e della relativa velocità, fornendo spesso file singoli più veloci ed un ripristino completo (dovuti ad un minor numero di backup da ricercare/ripristinare).

Date queste caratteristiche, vale la pena considerare i backup differenziali.

8.2.4. Media di backup

Durante le precedenti sezioni, abbiamo prestato molta attenzione nell'uso del termine "media di backup". Vi è una ragione per questo. Molti amministratori molto esperti generalmente considerano i backup, in termini di nastri di lettura e di scrittura, ma al giorno d'oggi vi sono altre opzioni.

I dispositivi a nastro rappresentavano i soli media disponibili in grado di essere rimossi, risultando quindi idonei all'uso per i processi di backup. Tuttavia, tutto questo è cambiato. Nelle sezioni seguenti affronteremo i media di backup più conosciuti, ed esamineremo i loro vantaggi insieme con i loro svantaggi.

8.2.4.1. Nastro

Il nastro è stato il primo mezzo, molto diffuso, di conservazione dati. Presenta il beneficio di un basso costo e di una buona capacità. Tuttavia, esso presenta qualche svantaggio — è soggetto a logoramento, ed il suo accesso dati è per natura sequenziale.

Questi fattori rendono necessario il mantenimento di informazioni sull'utilizzo del nastro (non utilizzando più i nastri che hanno raggiunto la fine del proprio ciclo), ed anche la ricerca di un file specifico può rappresentare un processo molto lungo.

D'altro canto, esso rappresenta il media di conservazione dati molto capiente, meno costoso ed estremamente affidabile attualmente disponibile. Ciò significa che la creazione di una libreria idonea non consuma una larga fetta del vostro budget, considerando anche la sua affidabilità per il futuro.

8.2.4.2. Disco

Negli anni passati le unità disco non venivano mai usate come mezzi di backup. Tuttavia, a causa della riduzione dei costi, in alcuni casi, utilizzare le unità disco come mezzo di backup potrebbe risultare il procedimento più idoneo.

Il motivo principale per l'utilizzo delle unità disco come mezzi di backup è la loro velocità. Non vi è altro mezzo di mass storage disponibile più veloce. La velocità potrebbe rappresentare un fattore

molto importante quando il vostro tempo di esecuzione di un backup è molto corto, e la quantità di dati da utilizzare per un backup è molto elevata.

Ma purtroppo tale mezzo di conservazione non rappresenta il mezzo ideale, questo è dovuto per diversi motivi:

- Generalmente le unità disco non possono essere rimosse. Un fattore importante per una strategia effettiva di backup, è quello di rimuovere i backup dal vostro centro dati e di inserirlo in un centro storage esterno. Un backup del vostro database di produzione che giace su di una unità disco vicina allo stesso database, non rappresenta un backup; è una copia. Generalmente le copie non sono molto utili se il centro dati ed il suo contenuto (incluso le vostre copie) vengono danneggiate o distrutte per qualche motivo.
- Le unità disco sono molto costose (se paragonate ad altri media di backup). Vi sono alcune circostanze dove l'aspetto economico non rappresenta un problema, ma in tutti gli altri casi, i costi associati con l'utilizzo delle unità disco per il backup, fanno sì che il numero di copie di backup siano il più basso possibile, in modo da mantenere al minimo i costi relativi. Un numero minore di copie di backup significa meno ridondanza se una copia dovesse risultare non leggibile.
- Le unità disco sono fragili. Anche se utilizzate più soldi per acquistare unità disco rimovibili, la loro fragilità potrebbe rappresentare un problema. Se fate cadere tali unità, significa praticamente aver perso il vostro backup. È possibile acquistare dei contenitori o astucci in grado di ridurre (ma non eliminare completamente) questo pericolo, aumentando però ulteriormente i costi.
- Le unità disco non rappresentano un media d'archiviazione. Anche se siete in grado di superare i problemi relativi ai backup sulle unità disco, è consigliato considerare quanto segue. Molte organizzazioni presentano diversi requisiti legali per mantenere le documentazioni per un tempo determinato. La possibilità di ottenere informazioni utili da un nastro vecchio di 20 anni, sono maggiori che ottenere informazioni da una unità disco vecchia di 20 anni



Nota Bene

Alcuni centri dati eseguono un backup utilizzando unità disco e poi, una volta terminato il backup, essi vengono salvati su nastro e quindi archiviati. Questo permette di avere il backup più veloce possibile durante il periodo di backup stesso. Il salvataggio dei backup su nastro, può essere eseguito durante il promemoria della giornata lavorativa; solo però se la "registrazione su nastro" termini prima del backup del giorno seguente.

Detto questo vi sono ancora alcuni esempi dove il backup su di una unità disco risulta essere l'alternativa migliore. Nella sezione successiva possiamo vedere come essi possono essere combinati con una rete creando una soluzione di backup fattibile (se costosa).

8.2.4.3. Rete

Da sola, una rete non è in grado di comportarsi come media di backup. Ma combinata con delle tecnologie di mass storage, potrebbe funzionare discretamente. Per esempio, combinando un link di una rete molto veloce, ad un centro dati remoto contenente grosse quantità di disk storage, gli svantaggi per un backup su dischi precedentemente indicati, non risulterebbero più essere dei veri e propri svantaggi.

Eseguendo un backup per mezzo della rete, le unità disco sono già esterne, in questo modo non vi è la necessità di trasportare le unità disco in un altro luogo. Con una larghezza di banda sufficiente, la velocità che potete ottenere dall'esecuzione del backup su di una unità disco viene conservata.

Tuttavia, questo approccio non è molto utile a risolvere il problema del mezzo di archiviazione (è possibile eseguire il "backup dei dati, passandoli successivamente su di un nastro", come indicato in precedenza). In aggiunta, i costi di un centro dati remoto con un link molto veloce per il centro dati

principale, renderebbe questa soluzione molto costosa. Ma per il tipo di organizzazione che necessita questo tipo di soluzione, i costi potrebbero rappresentare un fattore secondario.

8.2.5. Conservazione dei backup

Una volta completati i backup, cosa fare? La risposta più scontata è che i backup vengano conservati. Tuttavia cosa bisogna conservare — e dove.

Per rispondere a queste domande, bisogna prima considerare le circostanze con le quali è necessario usare i backup. Sono presenti tre diversi scenari:

1. Piccola, richieste di ripristino ad-hoc fatte dagli utenti
2. Ripristini enormi per riprendersi da un disaster
3. Storage d'archiviazione che non verrà più usato in futuro

Sfortunatamente vi sono delle differenze sostanziali tra numeri 1 e numeri 2. Quando un utente cancella accidentalmente un file, successivamente desidera ripristinarlo il più presto possibile. Ciò implica che il media di backup sia il più vicino possibile al sistema sul quale i dati devono essere conservati.

Nel caso di un disaster che necessita di un ripristino completo di uno o più computer nel vostro centro dati, se il suddetto è stato un disaster di natura fisica, esso potrebbe aver distrutto anche le impostazioni di backup vicino ai computer interessati. Questa situazione risulta essere tra le peggiori che si possano presentare.

Lo storage d'archiviazione è la parte che suscita meno controversie; in quanto le possibilità che possa venir utilizzato per altri compiti sono molto basse, se la posizione del media di backup risulta essere molto distante dal centro dati, non si dovrebbe riscontrare alcun problema.

Gli approcci seguiti per risolvere queste differenze variano a seconda delle necessità dell'organizzazione interessata. Un approccio possibile è quello di conservare internamente il lavoro di backup; questi backup vengono successivamente trasferiti in un luogo più sicuro dopo aver creato nuovi backup.

Un altro approccio potrebbe essere quello di mantenere due diversi gruppi di media:

- Un centro dati usato soltanto per richieste di ripristino ottimale
- Un gruppo esterno utilizzato per uno storage 'off-site' e per il disaster recovery

Ovviamente avere due gruppi implica la necessità di eseguire tutti i backup due volte, oppure eseguire una copia dei backup stessi. Tutto questo può essere fatto senza alcun problema, ma è da ricordare che un backup doppio richiede un tempo di esecuzione molto lungo, e l'esecuzione di una copia necessita unità di backup multiple per poter processare le copie stesse (e probabilmente un sistema apposito per eseguire la copia).

Il compito di un amministratore è quello di eseguire la scelta più idonea in grado di soddisfare le esigenze di tutti, assicurando in ogni momento la disponibilità dei backup in caso di necessità.

8.2.6. Problematiche riguardanti il processo di ripristino

Mentre i backup fanno parte di processi quotidiani, le fasi di ripristino invece risultano essere più rare. Tuttavia le suddette fasi sono necessarie, quindi è meglio farsi trovare sempre pronti.

La cosa più importante è quella di controllare tutti gli scenari possibili elencati in questa sezione, e determinare i modi per mettere alla prova la vostra abilità di risoluzione degli stessi. Ricordatevi che i più difficili sono sempre quelli che risultano essere i più critici.

8.2.6.1. Ripristino totale 'From Bare Metal'

Questo detto "ripristino totale o 'bare metal'" è il modo con il quale gli amministratori descrivono il processo di ripristino di un backup completo del sistema, su di un computer sprovvisto di dati di alcun genere — senza alcun sistema operativo, applicazioni ecc.

In generale vi sono due approcci per affrontare il suddetto tipo di ripristino:

Reinstallazione, seguita da un processo di ripristino

Con questo tipo di approccio il sistema operativo di base viene installato come se un nuovo computer fosse stato implementato. Una volta implementato il sistema operativo e configurato correttamente, le unità disco restanti possono essere partizionate e formattate, ripristinando tutti i backup tramite i media di backup.

Dischi di ripristino del sistema

Un disco di ripristino del sistema rappresenta un tipo di media avviabile (spesso un CD-ROM), che contiene un ambiente minimo del sistema, in grado di eseguire compiti di amministrazione di base del sistema stesso. L'ambiente per il processo di ripristino contiene le utility necessarie per partizionare e formattare le unità disco, le unità del dispositivo necessarie per accedere al dispositivo di backup, ed il software necessario per ripristinare i dati da un media di backup.



Nota Bene

Alcuni computer sono in grado di creare nastri di backup avviabili e di eseguire l'avvio degli stessi per iniziare il processo di ripristino. Tuttavia, questa capacità non è disponibile su tutti i computer. La maggior parte dei computer basati sull'architettura PC, non sono predisposti a tale approccio.

8.2.6.2. Test dei backup

Ogni tipo di backup deve essere testato periodicamente per assicurarsi che i dati in questione possano essere letti. In alcuni casi i backup, per qualche motivo, non risultano leggibili. La fase critica è rappresentata dal fatto che spesso questo processo viene ignorato fino a quando non si perdono effettivamente dei dati, avendo quindi il bisogno di ripristinarli tramite un processo di backup.

Tale errore si può verificare a causa di molteplici fattori, per esempio a causa di una modifica dell'allineamento della testina guida del nastro, oppure da una configurazione non corretta del software di backup, o da un errore dell'operatore. Non ha importanza quale sia la causa, senza una prova periodica, non sarete mai in grado di essere sicuri di generare dati di backup in grado di essere ripristinati in caso di necessità.

8.3. Disaster Recovery

La prossima volta che vi troverete nel vostro centro dati guardatevi intorno, e immaginatevi che niente sia a voi disponibile, non solo i computer, immaginate che l'intero edificio non sia più esistente. Successivamente immaginate che il vostro compito è quello di recuperare la maggior parte del lavoro eseguito precedentemente, il quale è andato perduto, il più presto possibile. Cosa fareste in questo caso?

Il solo pensiero a questo problema rappresenta il primo passo nel processo di un disaster recovery. Tale processo rappresenta la capacità di ripristinare tutte le funzioni del centro dati della vostra orga-

nizzazione, nel modo più veloce e completo possibile. Il tipo di disaster varia, ma il risultato finale è sempre lo stesso.

Le fasi da intraprendere per un disaster recovery sono numerose. Ecco una panoramica riguardante tale processo, insieme con i passaggi più importanti da ricordare.

8.3.1. Creazione, prova ed implementazione di un piano per un disaster recovery

Un sito di backup risulta essere vitale, ma potrebbe essere inutile se non si ha un piano per un disaster recovery. Tale piano indica ogni aspetto del processo sopra indicato, incluso ma non limitato a:

- Eventi che denotano potenziali disastri
- Personale autorizzato a dichiarare la presenza di un disaster e quindi l'adozione del piano previsto
- Una volta dichiarata la presenza di un disaster, intraprendere la sequenza di eventi necessari per preparare il sito di backup
- I ruoli e le responsabilità del personale che ricopre incarichi importanti per l'esecuzione del piano previsto
- Un inventario hardware e software necessario per ripristinare il normale funzionamento
- Un programma che elenca il personale coinvolto nel sito di backup, incluso un turno di rotazione per il supporto delle operazioni in uscita, senza sovraccaricare così i membri del team.
- La successione degli eventi necessari per smistare le operazioni dal sito di backup al nuovo/ripristinato centro dati

I piani per il disaster recovery spesso possono essere molto dettagliati. Questo livello di dettaglio è molto importante nel caso in cui si verifichi una emergenza, infatti il suddetto piano potrebbe risultare l'unica cosa disponibile del vostro centro dati (oltre all'ultimo backup esterno), in grado di permettervi di ricreare e ripristinare quindi le normali funzioni.



Suggerimento

Mentre i piani riguardanti un disaster recovery devono essere leggibili da parte di tutti i membri del vostro ufficio, altre copie dovrebbero essere implementate esternamente. In questo modo, se si verifica un disastro all'interno del vostro ufficio, esso non sarà in grado d'intaccare le copie situate all'esterno. Il luogo più idoneo per conservare una copia potrebbe essere il vostro storage di backup esterno. Se non risulta andare contro la policy della vostra organizzazione, alcune copie possono anche essere conservate a casa di alcuni dei membri più importanti facenti parte del team di emergenza.

Un documento così importante necessita un'attenzione particolare (e possibilmente un'assistenza professionale per la creazione).

Una volta creato tale documento risulta essere determinante la sua periodica prova. La prova di un piano riguardante il disaster recovery implica l'esecuzione delle diverse fasi, e cioè: andare sul sito di backup e impostare il centro dati provvisorio, eseguire le applicazioni in modo remoto, e riprendere le normali operazioni dopo che il "disaster" sia terminato. Molte delle prove non implicano l'esecuzione completa delle fasi; al contrario vengono selezionati alcuni sistemi insieme alle applicazioni da riposizionare nel sito di backup, metterle quindi in produzione per un periodo di tempo limitato, per poi passare alla loro normale funzione dopo aver superato il test previsto.

**Nota Bene**

Possiamo dire che il piano per un disaster recovery deve essere un documento vivente, se il centro dati cambia, il piano deve essere aggiornato in modo da riflettere questi cambiamenti. In molti casi, un piano 'scaduto', potrebbe recare maggior danno rispetto ad una organizzazione che non ne implementa alcuno, per questo motivo è consigliato rivedere tale piano regolarmente (ogni tre mesi per esempio) e aggiornarlo a seconda dei cambiamenti.

8.3.2. Siti di backup: Cold, Warm, Hot

Un aspetto molto importante del disaster recovery è quello di avere un luogo dal quale è possibile eseguire il ripristino delle funzioni del vostro centro dati. Questo luogo è conosciuto come *sito di backup*. Nel caso in cui si verifichi un disastro, un sito di backup è il posto nel quale si può ricreare il vostro centro dati, e da dove sarete in grado di eseguire il vostro compito per la durata dello stesso disastro.

Sono presenti tre diversi siti di backup:

- Siti di backup di tipo 'Cold'
- Siti di backup di tipo 'Warm'
- Siti di backup di tipo 'Hot'

Ovviamente questi termini non si riferiscono alla temperatura del sito di backup. Invece, essi si riferiscono allo sforzo necessario per iniziare le operazioni all'interno del sito di backup nell'evento di un disastro.

Un sito di backup di tipo cold, potrebbe essere rappresentato da uno spazio di un edificio propriamente configurato. Ogni cosa necessaria per ripristinare il servizio ai vostri utenti, deve essere procurato e consegnato al suo interno, prima che tale processo possa iniziare. Come potete immaginare, il ritardo che va da un sito di backup di tipo cold, al funzionamento completo può essere rilevante.

I siti di backup di tipo cold sono i siti meno costosi.

Un sito di backup di tipo warm è un sito che già presenta un hardware simile a quello trovato nel vostro centro dati. Per il ripristino del servizio, è necessario fornire gli ultimi backup presenti nello storage esterno, insieme al completamento della procedura del ripristino completo 'bare metal', prima che il vero lavoro di recupero possa iniziare.

I siti di backup di tipo hot possiedono una riproduzione virtuale del vostro centro dati con tutti i sistemi configurati, ed in attesa solo degli ultimi backup di dati appartenenti ai vostri utenti e provenienti dal vostro storage esterno. Come potete immaginare il suddetto sito può essere portato in produzione in poche ore.

Un sito di backup di tipo hot rappresenta l'approccio più costoso ad un disaster recovery.

I siti di backup vengono originati da tre diverse fonti:

- Compagnie specializzate nella fornitura di servizi per il disaster recovery
- Altri luoghi posseduti e gestiti dalla vostra organizzazione
- Un comune accordo per condividere i servizi del centro con un'altra organizzazione, nel caso in cui si verificasse un disastro

Ogni singolo approccio presenta un lato negativo ed uno positivo. Per esempio, un accordo con un'azienda specializzata nei disaster recovery, vi dà la possibilità di accedere a servizi e personale specializzato, utili per la creazione, la prova e l'implementazione di un piano per il disaster recovery stesso. Come potete immaginare, questi servizi non sono gratuiti.

Utilizzare uno spazio diverso posseduto e gestito dalla vostra organizzazione, potrebbe rappresentare una soluzione che non implica alcun costo aggiuntivo, ma solo i costi riguardanti il sito di backup e quelli di gestione della sua capacità di essere letto.

Raggiungere un accordo per la condivisione di un centro dati con un'altra organizzazione, potrebbe essere poco costoso, ma vi ricordiamo che un accordo che implica un periodo molto lungo sotto tali condizioni, non sono generalmente possibili, in quanto il centro dati deve sempre mantenere la propria normale produzione.

Per concludere, la scelta di un sito di backup deve rappresentare un compromesso tra il costo e le esigenze della vostra organizzazione, in modo da poter continuare la produzione.

8.3.3. Disponibilità software e hardware

Il vostro piano di disaster recovery deve includere i diversi metodi su come reperire l'hardware ed il software necessari per le operazioni all'interno del sito di backup. Un sito di backup gestito in modo professionale, potrebbe avere tutto ciò che è a voi necessario (oppure potrebbe risultare essere necessario organizzare il reperimento e la consegna di materiale specializzato non disponibile all'interno del sito); mentre un sito di backup di tipo cold implica l'identificazione di una fonte affidabile per ogni singolo oggetto. Molto spesso le organizzazioni lavorano insieme con i fornitori per raggiungere un accordo sulla consegna di hardware e/o software nel caso si verifichi un disastro.

8.3.4. Disponibilità dei dati di backup

Quando si dichiara un disastro, è necessario notificarlo al vostro storage esterno per due motivi:

- In modo da avere gli ultimi dati all'interno del sito di backup
- Per organizzare la consegna o la raccolta continua di dati al sito stesso di backup (in supporto ai backup normali nel sito di backup)



Suggerimento

Nel caso in cui si verifichi un disastro, gli ultimi backup ottenuti dal vostro ultimo centro dati risultano essere molto importanti. Prima di ogni cosa eseguire delle copie, e inviate il prima possibile gli originali al sito esterno.

8.3.5. Connettività della rete al sito di backup

Un centro dati potrebbe risultare inutile se lo stesso non è collegato con il resto dell'organizzazione. A seconda del piano implementato per il disaster recovery e della natura di quest'ultimo, la vostra community potrebbe essere molto lontana dal sito di backup, una buona connettività risulta essere quindi molto importante per il ripristino delle funzioni normali.

Un altro tipo di collegamento da ricordare è quello telefonico. Assicuratevi di avere un numero sufficiente di linee telefoniche in grado di poter gestire le comunicazioni con i vostri utenti; per questo motivo pianificate sempre un numero leggermente maggiore di collegamenti rispetto al numero previsto.

8.3.6. Assunzione di personale per il sito di backup

Il problema riguardante la designazione del personale di un sito di backup risulta avere diverse sfaccettature. Uno degli aspetti potrebbe essere la determinazione del personale necessario per poter gestire il centro di backup dei dati. Mentre il nucleo principale del team designato per la gestione del suddetto centro, potrebbe risultare idoneo alla sua gestione per un breve periodo di tempo, se un disastro si protrae per un periodo più lungo, è necessario designare un numero maggiore di personale, in modo da poter far fronte a tale problema.

Questo include la considerazione di un tempo sufficiente per il personale a smontare dal proprio turno e tornare nelle proprie abitazioni. Se si è verificato un disastro tale da influenzare anche la vita privata degli impiegati, allora è necessario considerare un periodo di riposo aggiuntivo. Da tenere in valida considerazione l'affitto di camere/alloggi situati il più vicino possibile al sito di backup, insieme all'affitto di mezzi di trasporto per facilitare lo spostamento del personale.

Spesso un piano per il disaster recovery include una rappresentativa sul posto, di tutti i membri facenti parte delle diverse community dell'organizzazione. Questo dipende anche dall'abilità della vostra organizzazione, di operare con un centro dati remoto. Se le diverse rappresentative hanno la necessità di lavorare nel sito di backup, allora sono necessari arrangiamenti simili a quelli precedentemente descritti.

8.3.7. Ritorno alla normalità

Tutti i disastri hanno una fine. Il piano per un disaster recovery deve affrontare anche quest'ultima fase. Il nuovo centro dati deve essere completo di tutto l'hardware ed il software necessario; anche se questa fase non rappresenta un momento critico durante la preparazione per affrontare un disastro, i siti di backup hanno un loro costo, per questo motivo è necessario smistarsi al centro dati appena possibile.

Gli ultimi dati del sito di backup, devono essere disponibili e consegnati al nuovo centro dati. Dopo averli ripristinati nel nuovo hardware, la produzione può essere smistata nel nuovo centro dati.

A questo punto il centro di backup dei dati può essere smantellato, insieme con tutto l'hardware provvisorio necessario per la parte finale del piano. Infine, è necessario eseguire una revisione sull'efficacia del piano stesso, con una eventuale implementazione di alcune modifiche da parte del comitato esaminatore, con la creazione di una versione aggiornata.

8.4. Informazioni specifiche di Red Hat Enterprise Linux

Non vi sono molte informazioni generali riguardanti i diversi disastri e disaster recovery, che possiedono un certo legame con un sistema operativo specifico. Dopo tutto, se si verifica un allagamento di un centro dati, i computer non potranno essere operativi sia che essi eseguano Red Hat Enterprise Linux, o un diverso sistema operativo. Tuttavia, vi sono alcune parti di Red Hat Enterprise Linux che hanno una certa relazione con alcuni aspetti specifici del disaster recovery; esse verranno affrontate in questa sezione.

8.4.1. Supporto software

In quanto rivenditore software, Red Hat offre una gamma di supporto per i suoi prodotti, incluso Red Hat Enterprise Linux. In questo momento possiamo dire che il lettore, leggendo questo manuale, sta utilizzando il toll di supporto di base. La documentazione per Red Hat Enterprise Linux è disponibile su CD di Documentazione di Red Hat Enterprise Linux (il quale può essere installato sul vostro sistema per un accesso rapido), può essere stampata, o presente sul sito web di Red Hat su <http://www.redhat.com/docs/>.

Le opzioni per il Self Support, sono disponibili tramite le diverse mailing list disponibili con Red Hat (presenti su <https://www.redhat.com/mailman/listinfo>). Queste mailing list offrono il vantaggio dalla conoscenza della community di Red Hat, molte delle suddette mailing list, vengono controllate dal personale di Red Hat, il quale può contribuire al loro sviluppo. Altre risorse sono disponibili tramite la pagina di supporto principale di Red Hat su <http://www.redhat.com/apps/support/>.

Esistono altresì altre opzioni di supporto più complete; informazioni in merito sono disponibili sul sito web di Red Hat.

8.4.2. Tecnologie di backup

Red Hat Enterprise Linux contiene diversi programmi per il backup ed il ripristino dei dati. Da soli questi programmi utility non costituiscono una soluzione di backup. Tuttavia, tale soluzione può essere utilizzata come parte principale delle soluzioni stesse.



Nota Bene

Come indicato nella Sezione 8.2.6.1, la maggior parte dei computer basati su di una architettura PC standard, non possiedono le funzionalità necessarie per eseguire l'avvio direttamente da un nastro di backup. Di conseguenza, Red Hat Enterprise Linux non risulta essere in grado di eseguire un avvio tramite il nastro quando esegue un hardware di questo tipo.

Tuttavia, è possibile utilizzare il CD-ROM di Red Hat Enterprise Linux come ambiente per il system recovery; per maggiori informazioni consultate il capitolo sul system recovery di base nella *Red Hat Enterprise Linux System Administration Guide*.

8.4.2.1. tar

La utility `tar` è molto conosciuta tra gli amministratori di sistema UNIX. Rappresenta la scelta del metodo di archiviazione per la condivisione ottimale del codice sorgente e dei file tra i sistemi. L'implementazione di `tar` incluso con Red Hat Enterprise Linux è GNU `tar`, una delle implementazioni `tar` più ricche.

Utilizzando `tar`, il backup dei contenuti di una directory può essere semplice quanto all'emissione di un comando simile al seguente:

```
tar cf /mnt/backup/home-backup.tar /home/
```

Questo comando crea un file d'archivio chiamato `home-backup.tar` in `/mnt/backup/`. L'archivio presenta i contenuti della directory `/home/`.

Il file d'archivio risultante avrà una grandezza simile ai dati verso i quali è stato eseguito il processo di backup. A seconda del tipo di dati, la compressione del file d'archivio potrebbe risultare in riduzioni significative della misura. Il file d'archivio può essere compresso aggiungendo una singola opzione al comando precedente:

```
tar czf /mnt/backup/home-backup.tar.gz /home/
```

Il file d'archivio `home-backup.tar.gz` risultante, è ora compresso `gzip`⁵.

Sono disponibili diverse altre opzioni su `tar`; per saperne di più, consultate la pagina `man di tar(1)`.

5. L'estensione `.gz` viene utilizzata tradizionalmente per indicare che il file è stato compresso con `gzip`. Talvolta `.tar.gz` viene abbreviato in `.tgz` per mantenere i nomi dei file con una misura accettabile.

8.4.2.2. cpio

La utility `cpio` è un altro programma tradizionale UNIX. È un programma generale molto efficiente, idoneo per lo spostamento dei dati da un luogo ad un altro, per questo motivo, potrebbe essere utile come programma di backup.

Il comportamento di `cpio` risulta essere leggermente diverso da `tar`. Diversamente da `tar`, `cpio` è in grado di leggere i nomi dei file da processare tramite un input standard. Un metodo comune per la creazione di un elenco di file per `cpio`, è quello di utilizzare i programmi come ad esempio `find`, che grazie al suo output viene collegato a `cpio`:

```
find /home/ | cpio -o > /mnt/backup/home-backup.cpio
```

Questo comando crea un file d'archivio `cpio` (in grado di contenere tutto nella `/home/`), chiamato `home-backup.cpio`, il quale risiede nella directory `/mnt/backup/`.



Suggerimento

Poiché `find` possiede un ricco set di test per la selezione dei file, è possibile creare facilmente backup molto complessi. Per esempio, il seguente comando esegue un backup di quei file che non sono stati visitati nell'arco di tempo di un anno:

```
find /home/ -atime +365 | cpio -o > /mnt/backup/home-backup.cpio
```

Sono disponibili molte altre opzioni su `cpio` (e `find`); per maggiori informazioni consultate le pagine man di `cpio(1)` e `find(1)`.

8.4.2.3. dump/restore: Non consigliato per i file systems montati!

I programmi `dump` e `restore` sono gli equivalenti ai programmi UNIX di Linux. Per questo motivo, molti amministratori di sistema con una certa esperienza in UNIX, possono considerare `dump` e `restore` come possibili candidati per un buon programma di backup sotto Red Hat Enterprise Linux. Tuttavia, un metodo di utilizzo di `dump`, potrebbe causare alcuni problemi. Ecco riportato il commento da parte di Linus Torvald su tale problema:

```
From: Linus Torvalds
To: Neil Conway
Subject: Re: [PATCH] SMP race in ext2 - metadata corruption.
Date: Fri, 27 Apr 2001 09:59:46 -0700 (PDT)
Cc: Kernel Mailing List <linux-kernel At vger Dot kernel Dot org>
```

```
[ linux-kernel added back as a cc ]
```

On Fri, 27 Apr 2001, Neil Conway wrote:

```
> > I'm surprised that dump is deprecated (by you at least ;-)). What to
> use instead for backups on machines that can't umount disks regularly?
```

Note that `dump` simply won't work reliably at all even in 2.4.x: the buffer cache and the page cache (where all the actual data is) are not coherent. This is only going to get even worse in 2.5.x, when the directories are moved into the page cache as well.

So anybody who depends on "dump" getting backups right is already playing Russian roulette with their backups. It's not at all guaranteed to get the right results - you may end up having stale data in the buffer cache that

ends up being "backed up".

Dump was a stupid program in the first place. Leave it behind.

```
> I've always thought "tar" was a bit undesirable (updates atimes or
> ctimes for example).
```

Right now, the cpio/tar/xxx solutions are definitely the best ones, and will work on multiple filesystems (another limitation of "dump"). Whatever problems they have, they are still better than the `_guaranteed_*` data corruptions of "dump".

However, it may be that in the long run it would be advantageous to have a "filesystem maintenance interface" for doing things like backups and defragmentation..

Linus

(*) Dump may work fine for you a thousand times. But it `_will_` fail under the right circumstances. And there is nothing you can do about it.

A causa di tale problema, l'utilizzo di `dump/restore` su file system montati è fortemente sconsigliato. Tuttavia, `dump` è stato creato originariamente per il backup di file system di tipo unmounted, e quindi in situazioni dove è possibile sospendere un file system con `umount`, `dump` resta una tecnologia di backup fattibile.

8.4.2.4. L'Advanced Maryland Automatic Network Disk Archiver (AMANDA)

AMANDA è una applicazione di backup basata su di un client/server creata nell'Università del Maryland. Avendo un'architettura client/server, un server di backup singolo (normalmente un sistema potente che presenta molto spazio sui dischi più veloci e configurato con il dispositivo di backup desiderato), è in grado di eseguire il backup di molti sistemi client, i quali non necessitano altro che del solo software client AMANDA.

Il suddetto approccio rappresenta il metodo migliore, in quanto concentra le risorse necessarie per i backup in un sistema, invece di richiedere hardware aggiuntivo per ogni sistema che necessita servizi di backup. Il design di AMANDA serve anche a centralizzare la gestione dei backup, facilitando i compiti dell'amministratore di sistema.

Il server di AMANDA gestisce un gruppo di media di backup e alterna l'utilizzo all'interno dello stesso, in modo da assicurare che ogni backup venga conservato per il periodo di memorizzazione dedicato all'amministratore. Tutti i media sono preformattati con dati che permettono ad AMANDA di rilevare se il media corretto è disponibile o meno. In aggiunta, AMANDA può essere interfacciata con un media robotizzato, cambiando le unità, automatizzando così completamente i processi di backup.

AMANDA è in grado di utilizzare sia `tar` che `dump` per eseguire i backup (anche se è preferibile con Red Hat Enterprise Linux utilizzare `tar`, a causa delle problematiche dovute all'uso di `dump`, accennate nella Sezione 8.4.2.3). Per questo motivo, i backup di AMANDA non richiedono AMANDA per ripristinare i file —

Normalmente AMANDA viene programmato per essere eseguito solo una volta al giorno durante la finestra di backup del centro dati. Il server di AMANDA si collega ai sistemi client indicando di effettuare delle previsioni dei backup da eseguire. Una volta ottenute tali previsioni, il server crea una tabella, determinando automaticamente l'ordine di backup dei sistemi stessi.

Una volta avviata la procedura di backup, i dati vengono inviati attraverso la rete dal client al server, dove vengono conservati in un disco. Una volta completato il backup, il server inizia la sua scrittura esternamente al suddetto disco, sul media di backup. Allo stesso tempo altri client sono in procinto di inviare i propri backup al server per la conservazione all'interno del disco. Ciò ne determina un flusso

continuo di dati disponibili alla scrittura sul media di backup. Man mano che i dati vengono scritti sul media di backup, essi vengono cancellati dal disco del server.

Una volta completati tutti i backup, viene inviato all'amministratore di sistema sotto forma di una email, un rapporto contenente lo stato dei backup, rendendo il controllo facile e veloce.

Se risulta necessario ripristinare i dati, AMANDA contiene un programma che permette all'operatore di identificare il file system, la data, ed i file name. Una volta fatto questo, AMANDA identifica il media di backup corretto, trovando e ripristinando i dati desiderati. Come precedentemente accennato, il design di AMANDA rende possibile il ripristino dei dati anche senza l'aiuto di AMANDA, anche se l'identificazione del media corretto sarà un processo manuale e molto lento.

Questa sezione ha affrontato solo i concetti di base di AMANDA. Per saperne di più, consultate le pagine `man di amanda(8)`.

8.5. Risorse aggiuntive

Questa sezione contiene le diverse risorse che possono essere usate per ottenere maggiori informazioni su di un disaster recovery, e sugli argomenti specifici di Red Hat Enterprise Linux affrontati in questo capitolo.

8.5.1. Documentazione

Le seguenti risorse sono state installate nel corso di una installazione tipica di Red Hat Enterprise Linux, e possono essere utilizzate per ottenere maggiori informazioni sugli argomenti affrontati in questo capitolo.

- Pagina `man di tar(1)` — Imparate ad archiviare i dati
- Pagina `man di dump(8)` — Imparate a scaricare i contenuti dei file system.
- Pagina `man di restore(8)` — Imparate a riprendere i contenuti dei file system salvati da `dump`.
- Pagina `man di cpio(1)` — Imparate a copiare i file da e sugli archivi.
- Pagina `man di find(1)` — Imparate a cercare i file.
- Pagina `man di amanda(8)` — Per saperne di più sui comandi che fanno parte del sistema di backup di AMANDA.
- File in `/usr/share/doc/amanda-server-<version>/` — Per saperne di più su AMANDA, controllate nuovamente questi documenti insieme con gli esempi di file.

8.5.2. Siti web utili

- <http://www.redhat.com/apps/support/> — La pagina di supporto di Red Hat fornisce un accesso facile alle diverse risorse relative al supporto di Red Hat Enterprise Linux.
- <http://www.disasterplan.com/> — Una pagina molto interessante con collegamenti a diversi siti relativi ad un disaster recovery. Include un piano molto semplice per un disaster recovery.
- <http://web.mit.edu/security/www/isorecov.htm> — La homepage del Massachusetts Institute of Technology Information Systems Business Continuity Planning, contenente diversi link informativi.
- <http://www.linux-backup.net/> — Una panoramica molto interessante riguardante diverse problematiche di backup.

- http://www.linux-mag.com/1999-07/guru_01.html — Un articolo molto interessante tratto dal Linux Magazine sugli aspetti tecnici della produzione di backup sotto Linux.
- <http://www.amanda.org/> — La homepage dell'Advanced Maryland Automatic Network Disk Archiver (AMANDA). Contiene alcuni pointer per le mailing list relative ad AMANDA e altre risorse online.

8.5.3. Libri correlati

I seguenti libri affrontano le diverse problematiche relative ad un disaster recovery, e rappresentano delle risorse molto utili per gli amministratori di sistema Red Hat Enterprise Linux:

- *Red Hat Enterprise Linux System Administration Guide*; Red Hat, Inc. — Include un capitolo sul system recovery, il quale potrebbe essere molto utile nell'esecuzione dei ripristini completi 'bare metal'.
- *Unix Backup & Recovery* di W. Curtis Preston; O'Reilly & Associates — Anche se non è stato scritto in modo specifico per i sistemi Linux, esso affronta in modo completo le diverse problematiche riguardanti il backup, includendo anche un capitolo sul disaster recovery.

Indice

Simboli

/etc/cups/, 148
/etc/passwd file
 gruppo, ruolo in , 134
 user account, ruolo in , 134
/etc/printcap, 148
/etc/shadow file
 gruppo, ruolo in , 135
 user account, ruolo in , 135

A

abuso delle risorse, 131
abuso, risorse, 131
account
 (Vd. user account)
amministratore di sistema
 filosofia di , 1
 automatizzare, 1
 compagnia, 6
 comunicazione, 3
 documentazione, 2
 eventi imprevisti, 8
 ingegneria sociale, i rischi di, 7
 pianificare, 8
 risorse, 5
 sicurezza, 7
 utenti, 6
attivazione della vostra sottoscrizione, iv
automatizzare, 9
 panoramica di , 1

B

backup
 acquistare un software, 172
 circostanze relative alle problematiche dei dati, 171
 conservazione di, 176
 creare un software, 172
 introduzione a , 170
 problematiche sul ripristino, 176
 ripristino totale 'bare metal', 177
 test del processo di ripristino, 177
 programma, modifica, 94
 software di backup AMANDA, 184
 tecnologie utilizzate, 182
 cpio, 183
 dump, 183
 tar, 182
 tipi di , 173
 backup completi, 173

 backup differenziali, 174
 backup progressivi, 173
 tipi di media, 174
 disco, 174
 nastro, 174
 rete, 175
bash shell, automazione e, 9

C

CD-ROM
 file system
 (Vd. file system ISO 9660)
comando vmstat, 56
comando chage, 137
comando chfn, 137
comando chgrp, 138
comando chmod, 138
comando chown, 138
comando chpasswd, 137
comando df, 103
comando free, 18, 56
comando gnome-system-monitor , 20
comando gpasswd, 137
comando groupadd, 137
comando groupdel, 137
comando groupmod, 137
comando grpck, 137
comando iostat, 22, 38
comando mpstat, 22
comando passwd, 137
comando raidhotadd, utilizzo di, 115
comando sa1, 22
comando sa2, 22
comando sadc, 22
comando sar, 22, 24, 39, 42, 57
 report, lettura, 24
comando top, 18, 19, 40
comando useradd, 137
comando userdel, 137
comando usermod, 137
comando vmstat, 18, 21, 38, 41
comando watch, 18
compagnia, conoscenza di, 6
comunicazione
 Informazioni specifiche su Red Hat Enterprise
 Linux, 10
 necessità di , 3
configurazione della stampante, 147
 applicazione di testo, 148
 CUPS, 147
controllo
 prestazione del sistema, 14
 risorse, 13
controllo della prestazione del sistema, 14

- controllo delle risorse , 13
 - capacità del sistema, 14
 - concetti, 13
 - cosa controllare, 14
 - larghezza della banda, 16
 - memoria, 17
 - pianificazione delle capacità, 14
 - Potenza della CPU, 15
 - prestazione del sistema, 14
 - storage, 17
- tool
 - free, 18
 - GNOME System Monitor, 20
 - iostat, 22
 - mpstat, 22
 - OPProfile, 25
 - sa1, 22
 - sa2, 22
 - sadc, 22
 - sar, 22, 24
 - Sysstat, 22
 - top, 19
 - vmstat, 21
- tool utilizzati, 18
- convenzioni
 - documento, ii
- CUPS, 147

D

- dati
 - accesso condiviso a, 129, 130
 - problematiche inerenti la ownership globale, 131
- devlabel, 98
- disaster recovery
 - backup, disponibilità di , 180
 - disponibilità hardware, 180
 - disponibilità software, 180
- fine di, 181
- introduzione a , 177
- pianificare, creare, provare, implementare, 178
- sito di backup, 179
 - assunzione di personale per, 181
 - connettività della rete al, 180
- dischi fissi, 50
- disk quota
 - abilitare, 111
 - gestione di, 112
 - introduzione a , 109
 - panoramica di , 109
 - controllo sull'utilizzo del blocco, 110
 - controllo sull'utilizzo dell'inode, 110
 - grace period, 111
 - limiti hard, 110

- limiti soft, 110
 - specifico al file-system, 109
 - specifico al gruppo, 110
 - specifico all'utente, 110
- dispositivo
 - accesso all'intero dispositivo, 97
 - alternative ai nomi del dispositivo, 97
 - convenzione del naming, 96
 - devlabel, naming con, 98
 - file name, 96
 - file system label, 98
 - label, file system, 98
 - nomi del dispositivo, alternative a, 97
 - partizione, 97
 - processo di naming con devlabel, 98
 - tipo, 96
 - unità, 96
- documentazione
 - Informazioni specifiche su Red Hat Enterprise Linux, 10
- documentazione, necessità di , 2

F

- file /etc/fstab
 - aggiornamento, 107
 - montaggio dei file system con, 104
- file /etc/group
 - gruppo, ruolo in , 136
 - user account, ruolo in , 136
- file /etc/gshadow
 - gruppo, ruolo in , 136
 - user account, ruolo in , 136
- file /etc/mtab, 101
- file /proc/mdstat, 114
- file /proc/mounts, 102
- file name
 - dispositivo, 96
- file system
 - label, 98
- file system EXT2 , 99
- file system EXT3 , 99
- file system ISO 9660, 99
- file system MSDOS, 100
- file system VFAT, 100

G

gestione
 stampanti, 141
 GID, 133
 group ID
 (Vd. GID)
 gruppo
 accesso di dati condivisi usando, 130
 file che controllano, 134
 /etc/group, 136
 /etc/gshadow, 136
 /etc/passwd, 134
 /etc/shadow, 135
 gestione di, 119
 GID, 133
 permessi relativi a, 132
 esecuzione, 132
 lettura, 132
 scrittura, 132
 setgid, 132
 setuid, 132
 sticky bit, 132
 struttura, determinazione, 130
 system GID, 133
 system UID, 133
 tool per la gestione, 137
 comando gpasswd, 137
 comando groupadd, 137
 comando groupdel, 137
 comando groupmod, 137
 comando grpck, 137
 UID, 133

H

hardware
 competenze necessarie per eseguire la riparazione, 151
 contratti sul servizio, 153
 budget per, 156
 disponibilità delle parti di ricambio, 156
 drop-off service, 154
 hardware ricoperto dal contratto, 156
 orario di copertura, 154
 servizio depot, 154
 Servizio di accesso 'walk-in service', 154
 tecnico interno, 155
 tempo di risposta, 155
 problemi di, 151
 unità di riserva
 mantenere, 151
 scambio di hardware, 153
 stock, quantità, 152
 stock, selezione di, 152
 home directory

centralizzata, 131
 home directory centralizzata, 131

I

imprevisti, preparazione per, 8
 informazioni specifiche di Red Hat Enterprise Linux
 controllo delle risorse
 larghezza di banda, 38
 potenza della CPU, 38
 disaster recovery, 181
 supporto software, 181
 supporto, software, 181
 tecnologie di backup
 AMANDA, 184
 cpio, 183
 dump, 183
 panoramica di, 182
 tar, 182
 tool di controllo delle risorse
 iostat, 38
 sar, 39, 42
 top, 40
 vmstat, 38, 41
 Informazioni specifiche su Red Hat Enterprise Linux
 automatizzare, 9
 bash shell, 9
 comunicazione, 10
 documentazione, 10
 intrusion detection systems, 10
 PAM, 10
 perl, 9
 RPM, 10
 Script della shell, 9
 sicurezza, 10
 tool per il controllo della risorsa, 18
 free, 18
 OProfile, 18
 Sysstat, 18
 top, 18
 vmstat, 18
 ingegneria sociale, i rischi di, 7
 ingegneria, sociale, 7
 interfaccia IDE
 panoramica di, 66
 interfaccia SCSI
 panoramica di, 67
 intrusion detection systems, 10

L

la filosofia di un amministratore di sistema, 1
 logical volume management
 (Vd. LVM)

lpd, 149

LVM

- contrastato con RAID, 84
- grouping dello storage, 83
- logical volume resizing, 84
- migrazione dati, 84
- migrazione, dati, 84
- panoramica di , 83
- resizing, logical volume, 84

M

memoria

- controllo di , 17
- memoria virtuale, 51
 - backing store, 52
 - page faults, 53
 - panoramica di, 52
 - prestazione, migliore delle ipotesi, 55
 - prestazione, peggiore delle ipotesi, 55
 - prestazioni della, 55
 - spazio dell'indirizzo virtuale, 53
 - swapping, 54
 - working set, 54
- utilizzo delle risorse della, 47

memoria cache, 48

memoria fisica

(Vd. memoria)

memoria virtuale

(Vd. memoria)

modifica del controllo, 168

montaggio di file system

(Vd. storage, file system, montaggio)

mount point

(Vd. storage, file system, mount point)

multiprocessazione simmetrica, 37

N

NFS, 103

nome utente, 119

cambiamenti, 121

collision tra , 120

naming convention, 119

O

OProfile, 18, 25

P

page description languages (PDL), 146

Interpress, 146

PCL, 146

PostScript, 146

page faults, 53

PAM, 10

partizione, 97

attributi di , 72

campo del tipo , 73

geometria, 73

tipo, 73

creazione di, 93, 105

estese, 73

logiche, 73

panoramica di , 72

primarie, 73

password, 122

corta, 123

da ricordare, 126

deboli, 123

forte, 125

informazioni personali usate in, 124

invecchiamento, 126

parole usate in , 124

più lunghe, 125

scritta, 125

set di caratteri ampio usato in , 125

set di caratteri minuscoli usato in , 123

trucchi semplici usati in , 124

usata ripetutamente, 125

perl, automazione e , 9

permessi, 132

tool per la gestione

comando chgrp, 138

comando chmod, 138

comando chown, 138

permesso di esecuzione, 132

permesso di lettura, 132

permesso di scrittura, 132

permesso setgid, 10, 132

permesso setuid, 10, 132

permesso sticky bit, 132

pianificare, importanza di , 8

pianificazione delle capacità, 14

pianificazione di un disaster, 151

alimentazione, backup, 162

blackout, esteso, 165

generatore, 164

insieme motore-generatore, 163

UPS, 163

tipi di disaster, 151

applicazioni usate in modo improprio, 166

aria condizionata, 165

crash del sistema operativo, 157

- elettriche, 160
- elettricità, qualità di, 161
- elettricità, sicurezza di, 160
- errori causati dal tecnico, 169
- errori dell'amministratore del sistema, 167
- errori di un utente, 166
- errori dovuti all'operatore, 167
- errori durante l'esecuzione delle procedure, 167
- errori nella configurazione, 168
- errori procedurali, 167
- errori relativi alla manutenzione, 169
- errori umani, 166
- HVAC (Heating, Ventilation, Air Conditioning), 165
- integrità dell'edificio, 160
- problematiche relative al sistema operativo, 157
- problematiche riguardanti l'ambiente, 160
- problemi con le applicazioni, 158
- problemi hardware, 151
- problemi software, 157
- relativo alle condizioni atmosferiche, 165
- riparazioni incorrette, 170
- riscaldamento, 165
- sospensioni del sistema operativo, 158
- ventilazione, 165

Pluggable Authentication Modules

(Vd. PAM)

potenza della CPU

(Vd. risorse, sistema, potenza di processazione)

potenza di processazione, risorse relative a

(Vd. risorse, sistema, potenza di processazione)

printconf

(Vd. configurazione della stampante)

printtool

(Vd. configurazione della stampante)

Q

quota, disk

(Vd. disk quota)

R

RAID

array

comando raidhotadd, utilizzo di, 115

creazione, 115

gestione di, 114

stato, controllo, 114

array, creazione, 113

al momento dell'installazione, 113

dopo il periodo d'installazione, 113

contrastato con LVM, 84

creazione degli array

(Vd. RAID, array, creazione)

implementazioni di, 82

hardware RAID, 82

software RAID, 83

introduzione a , 78

livelli di , 79

RAID 0, 79

RAID 0, svantaggi, 79

RAID 0, vantaggi, 79

RAID 1, 80

RAID 1, svantaggi, 80

RAID 1, vantaggi, 80

RAID 5, 80

RAID 5, svantaggi, 81

RAID 5, vantaggi, 81

RAID nidificati, 81

panoramica di , 78

RAID nidificati, 81

RAM, 49

Red Hat Enterprise Linux-informazioni specifiche

controllo delle risorse

memoria, 56

tool di controllo delle risorse

free, 56

sar, 57

vmstat, 56

registrazione della sottoscrizione, iv

registrazione della vostra sottoscrizione, iv

ricursion

(Vd. recursion)

risorse del sistema

(Vd. risorse, sistema)

risorse relative alla larghezza di banda

(Vd. risorse, sistema, larghezza di banda)

risorse, importanza di , 5

risorse, sistema

larghezza della banda

controllo di , 16

larghezza di banda, 31

bus, esempi di , 32

capacità, aumento, 33

carico, distribuzione, 33

carico, riduzione, 33

datapath, esempi di , 32

datapath, ruolo in, 32

panoramica di , 31

problemi relativi a , 32

regola per i bus in , 31

soluzioni ai problemi con , 33

memoria

(Vd. memoria)

potenza di processazione, 31

applicazioni, eliminazione, 36

capacità, aumento, 37

carenza di, migliorare, 35

carico, riduzione, 36

controllo di , 15

- CPU, migliorare, 37
- dati relativi a , 34
- migliorare, 37
- multiprocessazione simmetrica, 37
- overhead dell'applicazione, riduzione, 36
- overhead O/S, riduzione, 36
- panoramica di , 34
- SMP, 37
- utenze di , 34
- utilizzo da parte del sistema operativo, 35
- utilizzo da parte dell'applicazione di , 35

storage

(Vd. storage)

RPM, 10

RPM Package Manager

(Vd. RPM)

S

Script della shell, 9

sicurezza

importanza di, 7

Informazioni specifiche su Red Hat Enterprise Linux, 10

SMB, 104

SMP, 37

software

supporto per

documentation, 158

panoramica, 158

supporto autonomo, 158

supporto email, 159

supporto on-site, 159

supporto telefonico, 159

Supporto web, 159

spazio del disco

(Vd. storage)

spazio dell'indirizzo virtuale, 53

stampanti

colore , 144

CMYK, 144

inkjet, 144

laser, 145

considerazioni, 141

duplex, 141

gestione, 141

linguaggi

(Vd. page description languages (PDL))

locale, 147

networked, 147

risorse aggiuntive, 149

tipi, 141

cera termica 'thermal wax', 146

daisy-wheel, 142

dot-matrix, 142

impatto, 142

inchiostro solido 'solid ink', 146

inkjet, 144

laser, 145

laser a colori, 145

linea, 142

sublimazione 'dye-sublimation', 146

stampanti ad impatto, 142

daisy-wheel, 142

dot-matrix, 142

linea, 142

materiali di consumo, 144

stampanti daisy-wheel

(Vd. stampanti ad impatto)

stampanti di linea

(Vd. stampanti ad impatto)

stampanti dot-matrix

(Vd. stampanti ad impatto)

stampanti inkjet, 144

materiali di consumo, 144

stampanti laser, 145

colore , 145

materiali di consumo, 145

stampanti laser a colori, 145

statistiche di controllo

relativo alla CPU, 15

relativo alla larghezza della banda, 16

relativo alla memoria, 17

relativo allo storage, 17

selezione di , 14

storage

accessibile attraverso la rete, 77, 103

NFS, 103

SMB, 104

aggiunta, 90, 105

/etc/fstab, aggiornamento, 107

configurazione, aggiornamento, 94

formattare, 93, 106

hardware, installazione, 90

partizionamento, 93, 105

programma di backup, modifica, 94

unità disco ATA, 90

unità disco SCSI, 91

basato su RAID

(Vd. RAID)

controllo di , 17

disk quota, 87

(Vd. disk quota)

dispositivi mass-storage

braccia d'accesso, 62

carichi I/O, letture, 71

carichi I/O, prestazione, 71

carichi I/O, scritture, 71

cilindro, 63

concetti relativi, 63

disk platter, 61

- geometria, problemi con, 64
- interfacce per , 65
- interfacce standard del settore, 66
- interfacce, standard del settore, 66
- interfacce, storico, 65
- interfaccia IDE, 66
- interfaccia SCSI, 67
- lettori contro scrittori, 71
- limitazioni elettriche di, 69
- limitazioni meccaniche di, 69
- luogo I/O, 72
- metodo basato sul blocco, 65
- metodo basato sulla geometria, 63
- metodo, basato sul blocco, 65
- metodo, basato sulla geometria, 63
- movimento del braccio d'accesso, 70
- movimento, braccio d'accesso, 70
- panoramica di , 61
- piatti, disco, 61
- prestazione di, 69
- processazione del comando, 70
- processazione, comando, 70
- settore, 64
- tempo di rotazione, 70
- tempo, rotazione, 70
- testina, 64
- testine, 61
- testine di lettura, 70
- testine di scrittura, 70
- file system, 74, 98
 - abilitare l'accesso allo , 77
 - basato sul file, 74
 - controllo dell'accesso, 75
 - controllo dello spazio, 75
 - controllo, spazio, 75
 - directory, 74
 - directory gerarchica, 74
 - EXT2, 99
 - EXT3, 99
 - file /etc/mntab, 101
 - file /proc/mounts, 102
 - ISO 9660, 99
 - montaggio, 100
 - montaggio con il file /etc/fstab, 104
 - mount point, 100
 - MSDOS, 100
 - struttura, directory, 76
 - tempi di accesso, 75
 - tempi di creazione, 75
 - tempi di modifica, 75
 - utilizzo, comando df, 103
 - VFAT, 100
 - visualizzare, 101
- gestione di, 61, 84
 - aumento, normale, 87
 - controllo dello spazio disponibile, 85
 - problematiche riguardanti un utente, 85
 - uso eccessivo di, 85
 - utilizzo di una applicazione, 87
- impiego, 72
- partizione
 - attributi di , 72
 - campo del tipo , 73
 - estese, 73
 - geometria di, 73
 - logiche, 73
 - panoramica di , 72
 - primarie, 73
 - tipo di , 73
- percorsi d'accesso, 47
- problematiche relative al file, 88
 - accesso del file, 88
 - file sharing, 89
- rimozione, 94, 107
 - /etc/fstab, rimuovere da, 107
 - cancellare i contenuti, 95, 108
 - comando umount, utilizzo di , 108
 - dati, rimozione, 95
- tecnologie, 47
 - disco fisso, 50
 - L1 cache, 49
 - L2 cache, 49
 - memoria cache, 48
 - memoria principale, 49
 - RAM, 49
 - Registri CPU, 48
 - storage di backup, 51
 - storage off-line, 51
 - unità disco, 50
- tecnologie, avanzate, 77
- Strumento di configurazione della stampante
 - (Vd. configurazione della stampante)
- swapping, 54
- Sysstat, 18, 22
- system-config-printer
 - (Vd. configurazione della stampante)

T

tool

- controllo delle risorse , 18
 - free, 18
 - GNOME System Monitor, 20
 - iostat, 22
 - mpstat, 22
 - OProfile, 25
 - sa1, 22
 - sa2, 22
 - sadc, 22
 - sar, 22, 24
 - Sysstat, 22
 - top, 19
 - vmstat, 21
- gruppi, gestione
 - (Vd. gruppo, tool per la gestione)
- user account, gestione
 - (Vd. user account, tool per la gestione)

U

- UID, 133
- unità disco, 50
- unità disco ATA
 - aggiunta, 90
- unità disco SCSI
 - aggiunta, 91
- user account
 - accesso ai dati condivisi, 129
 - controllo dell'accesso, 127
 - file che controllano, 134
 - /etc/group, 136
 - /etc/gshadow, 136
 - /etc/passwd, 134
 - /etc/shadow, 135
 - gestione di , 119, 119, 127
 - modifiche degli incarichi , 129
 - nuovi dipendenti, 127
 - Termine di un rapporto lavorativo, 128
- GID, 133
- home directory
 - centralizzata, 131
- nome utente, 119
 - cambiamenti a, 121
 - le collision in naming, 120
 - naming convention, 119
- password, 122
 - corta, 123
 - da ricordare, 126
 - deboli, 123
 - forte, 125
 - informazioni personali usate in, 124
 - invecchiamento, 126
 - parole usate in , 124

- più lunghe, 125
- scritta, 125
- set di caratteri ampio usato in , 125
- set di caratteri minuscoli usato in , 123
- trucchi semplici usati in , 124
- usata ripetutamente, 125
- permessi relativi a , 132
 - esecuzione, 132
 - lettura, 132
 - scrittura, 132
 - setgid, 132
 - setuid, 132
 - sticky bit, 132
- risorse, gestione di, 129
- system GID, 133
- system UID, 133
- tool per la gestione, 137
 - comando chage, 137
 - comando chfn, 137
 - comando chpasswd, 137
 - comando passwd, 137
 - comando useradd, 137
 - comando userdel, 137
 - comando usermod, 137

UID, 133

user ID

(Vd. UID)

utenti

importanza di, 6

W

working set, 54

Colophon

I manuali sono scritti in formato DocBook SGML v4.1. I formati HTML e PDF vengono prodotti usando i fogli stile DSSSL personali e script wrapper jade personali. I file SGML DocBook sono scritti in **Emacs** con l'aiuto della modalità PSGML.

Garrett LeSage ha creato le grafiche di ammonizione (nota, suggerimento, importante attenzione e avviso). Essi possono essere ridistribuiti liberamente con la documentazione Red Hat.

Il team di documentazione del prodotto di Red Hat é composto dalle seguenti persone:

Sandra A. Moore — Scrittore principale/Maintainer della *Red Hat Enterprise Linux Installation Guide per x86, Itanium™, AMD64, e Intel® Extended Memory 64 Technology (Intel® EM64T)*; Scrittore principale/Maintainer della *Red Hat Enterprise Linux Installation Guide per l'Architettura IBM® POWER*; Scrittore principale/Maintainer della *Red Hat Enterprise Linux Installation Guide per le architetture IBM® S/390® e IBM® eServer™ zSeries®*

John Ha — Scrittore Principale/Maintainer della *Red Hat Cluster Suite Configurazione e gestione di un Cluster*; Scrittore/Maintainer della *Red Hat Enterprise Linux Security Guide*; Maintainer delle stylesheet DocBook personali e degli script

Edward C. Bailey — Scrittore primario/Controllore della *Red Hat Enterprise Linux Introduzione al System Administration*; Scrittore primario/Controllore delle *Release Note*; Contributing Writer alla *Red Hat Enterprise Linux Installation Guide per x86, Itanium™, AMD64, e Intel® Extended Memory 64 Technology (Intel® EM64T)*

Karsten Wade — Scrittore primario/Maintainer della *Red Hat SELinux Application Development Guide*; Scrittore primario/Maintainer della *Red Hat SELinux Policy Writing Guide*

Andrius Benokraitis — Scrittore principale/Maintainer della *Red Hat Enterprise Linux Reference Guide*; Scrittore/Maintainer della *Red Hat Enterprise Linux Security Guide*; Assistente per la *Red Hat Enterprise Linux System Administration Guide*

Paul Kennedy — Scrittore principale/Maintainer della *Red Hat GFS Administrator's Guide*; Assistente alla *Red Hat Cluster Suite Configurazione e gestione di un Cluster*

Mark Johnson — Scrittore principale/Maintainer della *Red Hat Enterprise Linux Configurazione Desktop e Administration Guide*

Melissa Goldin — Scrittore principale/Maintainer della *Red Hat Enterprise Linux Guida passo dopo passo*

Il team di localizzazione del prodotto di Red Hat é composto dalle seguenti persone:

Amanpreet Singh Alam — Traduttore della versione Punjabi

Jean-Paul Aubry — Traduttore della versione francese

David Barzilay — Traduttore della versione brasiliana/portoghese

Runa Bhattacharjee — Traduttrice della versione Bengali

Chester Cheng — Traduttore della versione cinese tradizionale

Verena Fuehrer — Traduttrice della versione tedesca

Kiyoto Hashida — Traduttore della versione giapponese

N. Jayaradha — Traduttrice della versione Tamil

Michelle Jiyeen Kim — Traduttrice della versione coreana

Yelitza Louze — Traduttrice della versione spagnola

Noriko Mizumoto — Traduttrice della versione giapponese

Ankitkumar Rameshchandra Patel — Traduttore della versione Gujarati

Rajesh Ranjan — Traduttore della versione Hindi

Nadine Richter — Traduttrice della versione tedesca

Audrey Simons — Traduttrice della versione francese

Francesco Valente — Traduttore della versione italiana

Sarah Wang — Traduttrice della versione cinese semplificato

Ben Hung-Pin Wu — Traduttore della versione cinese tradizionale