

ATENȚIONARE!

Conținutul acestei platforme de instruire a fost elaborat în cadrul proiectului „Dezvoltarea resurselor umane în educație pentru administrarea rețelelor de calculatoare din școlile românești prin dezvoltarea și susținerea de programe care să sprijine noi profesii în educație, în contextul procesului de reconversie a profesorilor și atingerea masei critice de stabilizare a acestora în școli, precum și orientarea lor către domenii cerute pe piața muncii”. Conținutul platformei este destinat în exclusivitate pentru activități de instruire a membrilor grupului țintă eligibil în proiect.

Utilizarea conținutului în scopuri comerciale sau de către persoane neautorizate nu este permisă.

Copierea, totală sau parțială, a conținutului de instruire al acestei platforme de către utilizatori autorizați este permisă numai cu indicarea sursei de preluare (platforma de instruire eadmin.cpi.ro).

Pentru orice probleme, nelămuriri, sugestii, informații legate de aspectele de mai sus vă rugăm să utilizați adresa de email: proiect.eadmin@cpi.ro

Acest material a fost elaborat de Cristian Oftez, în cadrul S.C. Centrul de Pregătire în Informatică S.A., partener de implementare a proiectului POSDRU I/3/1.3/S/5.

Versiunea materialului de instruire: V2.0

9. Administrarea Serverelor

Drivere, module și dispozitive în Linux

Având în vedere faptul că numărul de producători de dispozitive hardware de toate felurile a crescut constant, există în prezent mii de astfel de dispozitive (și variază de la un simplu mouse până la plăci grafice de performanță care sunt foarte specializate), iar sistemele de operare existente au trebuit să țină pasul cu acestea, astfel încât să permită funcționarea corectă a acestora. Acest fapt este posibil datorită existenței driverilor, care oferă o interfață programabilă cu restul sistemului, astfel încât programatorii să poată folosi facilitățile oferite de diverse dispozitive atașate sistemului fără să fie nevoiți să cunoască în adâncime modul de funcționare al acestora. În esență, driverele pot fi privite ca un mecanism de simplificare și uniformizare a părții hardware, astfel încât aceasta să poată fi folosită mai ușor. Raționamentul din spatele acestei interfețe programabile foarte bine definite este simplu: acesta permite o dezvoltare modulară, în care drivere (fiind total independente de kernel) pot fi inițializate atunci când este nevoie de ele, simplificând astfel foarte mult procesul de creare al acestora. Aceste module sunt cunoscute în Linux ca **LKM**, sau **Loadable Kernel Module** și reprezintă extensii ale kernel-ului. O greșală comună în rândul utilizatorilor de Linux este concepția că aceste module fac parte din kernel în mod absolut, însă acesta poate fi împărțit în două părți importante: imaginea kernel-ului de „bază” (sau **base kernel**), adică imaginea care este încărcată în timpul procesului de **boot** și aceste extensii (**LKM**), care comunică în mod direct cu partea de bază.

Caracterul modular oferă sistemului Linux un avantaj foarte mare printr-un grad de flexibilitate crescut, modulele putând fi oprite când nu este nevoie de ele sau încărcate numai atunci când este nevoie de acestea. Există două alternative în implementarea LKM-urilor: încărcarea acestora ca extensii, sau includerea lor în kernel-ul de bază. De multe ori, prima abordare este și cea mai bună, din mai multe motive: în primul rând nu mai forțează utilizatorul să recompileze kernel-ul (și deci reduce semnificativ probabilitatea apariției erorilor sau greșelilor în kernel-ul de bază), nu necesită reîncărcarea sistemului și de foarte multe ori salvează memorie (**base kernel-ul** fiind încărcat mereu în memorie), încărcând modulele numai atunci când este nevoie de acestea.

Un alt avantaj major al LKM-urilor îl reprezintă ușurința cu care acestea pot fi depanate în cazul în care se întâmplă o eroare, de multe ori, erori în kernel-ul de bază ducând la imposibilitatea încărcării sistemului. Aceste module nu înregistrează o scădere în performanță comparativ cu cele incluse în kernel-ul de bază, însă există cazuri în care această abordare nu este posibilă și kernel-ul trebuie recompilat. Vom aborda pe larg această problemă în secțiunea dedicată recompilării unui kernel.

LKM-urile pot să îndeplinească o serie întreagă de funcții într-un sistem Linux, însă cele mai comune roluri sunt:

- **Drivere** – așa cum am menționa mai devreme, acestea reprezintă o interfață care permite kernel-ului să comunice cu dispozitivele pentru care au fost create, fără a fi nevoie de cunoștiințe detaliate despre modul în care acestea funcționează.
- **Funcții de sistem** – oferă programelor metode de a apela funcțiile puse la dispoziție de sistem cum ar fi crearea unui nou proces, citirea unui fișier sau chiar oprirea sistemului. Deși majoritatea apelurilor de sistem sunt deja incluse în kernel-ul de bază, însă se pot crea noi apeluri sau chiar înlocui unele deja existente prin intermediul LKM-urilor.
- **Interpretoare de executabile** – care fac posibilă încărcarea în memorie și rularea programelor executabile în diverse formate.
- **Drivere pentru sistemul de fișiere** – astfel încât sistemul de operare să poată interpreta conținutul fizic și implicit și cel logic (modul în care sunt organizate directoarele și fișierele pe un mediu de stocare sau de pe rețea)

Urmând această arhitectură, sistemul Linux pune la dispoziția utilizatorului o serie de utilitare prin intermediul cărora să poată controla modulele existente.

Pentru **vizualizarea modulelor încărcate**, utilizatorul poate să folosească comanda **lsmod** (probabil prescurtat de la „list modules”). Acest utilitar, după cum menționează și descrierea sa din manual (man), este unul oarecum trivial: unicul său rol este acela de **formatare a informațiilor** (aflate într-unul din fișierele speciale) într-o formă ce poate fi interpretată mai ușor, afișând pe ecran numele, dimensiunea și **dependințele** tuturor modulelor de kernel încărcate.

Module	Size	Used by
psmouse	32336	0
ppdev	6468	0
lp	8164	0
ipov6	235396	16
cpufreq_conservative	5968	0
cpufreq_stats	3776	0
cpufreq_userspace	3172	0
cpufreq_powersave	1856	0
cpufreq_ondemand	6476	0
freq_table	4224	2 cpufreq_stats,cpufreq_ondemand
vboxvfs	36704	0
nls_base	6820	1 vboxvfs
loop	12748	0
parport_pc	22500	0
parport	30988	3 ppdev,lp,parport_pc
serio_raw	4740	0
pcspkr	2432	0
battery	10180	0
ac	4196	0
button	6096	0
vboxadd	58848	4 vboxvfs
i2c_piix4	7216	0
i2c_core	19828	1 i2c_piix4
:_		

Să analizăm primul modul listat în poza de mai sus, cu numele de **psmouse**. Acesta reprezintă modulul responsabil de interfațarea cu dispozitive periferice de tip mouse, conectate prin PS2. Putem observa faptul că modulul respectiv are dimensiunea de **32336** bytes, iar câmpul următor, cu valoarea 0, reprezintă numărul de entități ce depind de acest modul.

Pentru **înlăturarea unui modul** se pot folosi atât comenzile **modprobe** dar și **rmmod**, însă doar prima este recomandată. Modul de utilizare al acesteia este destul de simplu, iar utilizatorul trebuie să specifice doar doi parametri: numele modului și parametrul care specifică faptul că acesta dorește înlăturarea modului (-r). Dacă ar fi să urmărim exemplul de mai sus și am dori să excludem modulul **psmouse**, ar trebui executată comanda **modprobe -r psmouse**. Efectul înlăturării unui modul este de regulă instant, având ca urmare (în acest caz) pierderea funcționalității mouse-ului.

În cazul în care utilizatorul dorește **punerea în funcțiune a unui modul**, din diverse motive, acesta poate utiliza tot comanda **modprobe [nume_modul]**. În exemplul de mai sus, putem folosi **modprobe psmouse** pentru a putea folosi mouse-ul din nou.

Încărcarea automată a modulelor

Sistemul poate fi configurat astfel încât să încarce un modul LKM în mod automat atunci când kernel-ul are nevoie de acesta prima oară. Această încărcare automată poate fi realizată în principiu prin două mecanisme: fie prin **kernel module loader** (care face parte din kernel-ul Linux), fie prin versiunea mai veche a acestuia: un daemon numit **kerneld**. Amândouă variantele folosesc utilitarul **modprobe** (deci implicit și **insmod**) pentru a-și duce sarcina la capăt.

Modul de funcționare al Kernel Module Loader

Așa cum am menționat mai devreme, KML (Kernel Module Loader) face parte directă din kernel și are scopul de a încărca module LKM în mod automat atunci când este nevoie de acestea. Când necesitatea unui modul devine aparență și acesta trebuie încărcat, KLM crează un nou proces (al cărui proprietar este contul root) care va încerca să încarce modulul dorit prin executarea utilitarului **modprobe** (acesta aflându-se la adresa **/sbin/modprobe** în mod implicit) cu o serie de parametri specifici: **-s** (forțează notarea oricăror erori în rapoartele de sistem), **-k** (modul „autoclean”) și desigur, numele modului ce trebuie încărcat.

De multe ori, ultimul parametru pasat utilitarului **modprobe** nu este neapărat numele unui fișier de tip modul, ci mai degrabă un nume simbolic, un alias pentru simplificarea numelor folosite. Lista acestora se găsește de regulă în fișierul **/etc/modprobe.d/aliases**.

Eliminarea automată a modulelor

În același fel în care modulele pot fi încărcate în mod automat, există și un mecanism de înlăturare a acestora, numit **autoclean**. Motivul principal pentru care acest mecanism a fost implementat, este înlăturarea modulelor care nu au mai fost folosite pentru perioade prelungite de timp (de regulă peste 1-2 minute) salvându-se altfel memorie și timp de procesor. Kernel Module Loader nu utilizează acest mecanism spre deosebire de **kernel**, însă poate fi forțat prin executarea comenzii **rmmod -all** la intervale diferite de timp. În cazul în care utilizatorul dorește să beneficieze de acest mecanism, dar nu dorește să folosească **kernel**, acesta poate programa execuția comenzii **rmmod -all** ca un **cron job**, astfel încât acesta să fie executat automat la intervale de timp fixe.

Interfațarea cu structurile de date ale kernel-ului

Kernel-ul Linux are două roluri majore în cadrul unui sistem: să intermedieze și să controleze accesul cu diferitele dispozitive conectate la sistem și să gestioneze interacțiunea proceselor (adică când și cum vor interacționa) cu aceste dispozitive.

Directorul **/proc** reprezintă o ierarhie de fișiere speciale, prin intermediul căreia se pot obține o multitudine de informații despre starea actuală a kernel-ului. Acest director constituie și principalul mecanism prin care aplicațiile și utilizatorii pot obține informații despre sistem.

Respectând structura unui sistem Linux (unde totul este un fișier) și analizând conținutul directorului **/proc** putem trage o concluzie interesantă: acesta este de fapt un pseudo sistem de fișiere dat fiind faptul că introduce un nou tip de fișier: fișierul virtual. Acesta are o calitate deosebită prin faptul că deși mare parte din acestea au dimensiunea de zero bytes, la afișare vor prezenta o cantitate mare de informații. În esență, aceste fișiere virtuale sunt mecanismul prin care kernel-ul Linux oferă informații despre sistem, folosind arhitectura cunoscută de „totul este un fișier”. Informația este structurată în directoare și subdirectoare în funcție de categoria din care face parte. Deși majoritatea fișierelor au dimensiunea zero, și sunt read-only, acestea sunt reactualizate de kernel la intervale regulate. Totuși, există câteva fișiere din directorul **/proc** care pot fi alterate (de regulă doar de contul root) modificând configurația kernel-ului. Pentru a afla de exemplu informații despre tipul de procesor instalat în sistem, utilizatorul trebuie doar să afișeze (în orice fel) fișierul **/proc/cpuinfo**, cu mențiunea că deși o parte din aceste informații sunt ușor de înțeles de către utilizator, există fișiere care pot fi interpretate de oameni cu greutate – aici devenind aparent rolul utilităților existente precum **top**, **free**, etc.

Pentru a putea înțelege mai bine cum funcționează acest pseudo sistem de fișiere și ce fel de informații ne oferă, vom analiza principalele zone ale acestuia. De regulă, fiind vorba de un alt sistem de fișiere decât cel de bază și în baza celor discutate în secțiunile anterioare, acesta trebuie „montat”

într-un director pentru a putea fi folosit, acesta fiind de regulă chiar **/proc** din directorul rădăcina. Principalele categorii din **/proc** sunt:

- **/proc/[identificador_numeric]/** - Directorul **/proc** va conține o serie de subdirectoare al căror nume este un identificador numeric. Acestea simbolizează de fapt procese, iar numele directorului este dat de PID-ului fiecărui proces ce rulează în momentul curent. Directorul va conține o serie fișiere și subdirectoare:
 - **/cmdline** conține linia de comandă cu care a fost lansat executabilul, incluzând și lista de parametri.
 - **/cwd** este un link simbolic către directorul activ al procesului în cauză.
 - **/environ** conține informații despre variabilele de stare și mediu cu care rulează procesul.
 - **/exe** reprezintă un link spre executabilul ce reprezintă procesul, având inclusă calea absolută.
 - **/fd/** este un subdirector ce conține link-uri către toate fișierele deschise de proces: 0 – standard input, 1 – standard output, 2 – standard error și așa mai departe. Fișierele sunt numite după descriptorul de fișiere folosit și pentru a vedea exact de ce fișiere sunt legate este suficientă comanda **ls -l**.
 - **/stat** conține informații despre procese, și este unul din fișierele pe care le consultă comanda **ps** pentru a obține informații despre procesele care rulează. Multe din informațiile conținute de acest fișier sunt de natură tehnică, și utile în special dezvoltatorilor de software. Pentru informații suplimentare, apălați comanda **man proc**.
- **/proc/apm** – prescurtarea de la advanced power management, ne oferă informații despre starea bateriei în cazul în care așa ceva există și kernel-ul a fost compilat corespunzător.
- **/proc/bus** – conține subdirectoare pentru toate magistralele instalate
 - **/pci** este un subdirector ce conține informații despre magistralele pci, dispozitive și drivere.
 - **/pci/devices** – informații despre dispozitivele conectate prin pci. Se recomandă folosirea comenzii **lspci** pentru vizualizarea conținutului acestui fișier.
- **/proc/cpuinfo** – reprezintă o sumarizare a arhitecturii sistemului.
- **/proc/diskstats** – este un fișier ce conține informații despre activitatea input/output a fiecărui disc fizic. Formatul acestui fișier este oarecum greu de citit, și se recomandă folosirea comenzii **df** pentru aflarea informațiilor legate de discuri.
- **/proc/filesystem** – este o listă cu toate tipurile de sisteme de fișiere cu care kernel-ul a fost compilat, fiind folosită de comandă **mount** pentru ciclare, atunci când nu se specifică în mod special un sistem de fișiere.
- **/proc/kmsg** - reprezintă o interfață prin care utilizatorul poate să citească mesaje generate de kernel. Pentru a putea citi acest fișier, este nevoie de

drepturi de root, și nu se recomandă accesarea simultană de către mai multe programe a acestui fișier. Pentru ușurință, conținutul acestui fișier este cel mai ușor de accesat prin intermediul comenzii **dmesg**.

- **/proc/meminfo** – conține informații despre memorie (atât liberă cât și utilizată, fizică și virtuală) și este interpretată cel mai ușor prin comanda **free**.
- **/proc/mounts** – oferă informații despre sistemele de fișiere montate în momentul afișării.
- **/proc/net** – oferă o serie de fișiere virtuale, prin care se pot obține informații despre starea curentă a rețelei. Deși mare parte din aceste informații sunt ușor de interpretat de către oameni, comanda **netstat** poate forma informațiile într-o măsură mai bună. Directorul **/proc/net** conține următoarele puncte de interes:
 - **/dev** – un fișier ce conține informații despre starea dispozitivelor de comunicare, precum plăci de rețea. Pe lângă acestea, fișierul mai conține și numărul de pachete trimise și primite pentru fiecare interfață, cât și numărul de erori și alte informații statistice. Acest fișier este citit de comandă **ifconfig** pentru afișarea și configurarea tuturor interfețelor de rețea.
 - **/tcp**, **/udp**, **/raw** conțin datele aferente socket-urilor de comunicație pentru fiecare protocol în parte. Aceste fișiere sunt de regulă folosite numai în scopul depanării sau monitorizării.
- **/proc/partitions** – afișează o tabelă ce evidențiază modul în care sunt partiționate hard-discurile.
- **/proc/sys** – acest director conține fișiere și subdirectoare ce alcătuiesc mediul și variabilele de kernel. Aceste variabile pot fi schimbate prin intermediul pseudo sistemului de fișiere **/proc** sau prin intermediul utilitarului **sysctl**.
- **/proc/uptime** – un fișier ce conține două valori numerice. Prima dintre acestea reprezintă timpul de funcționare al sistemului (măsurat în secunde zecimale) sau **uptime**, iar al doilea reprezintă timpul **idle** al sistemului (măsurat tot în secunde zecimale). Timpul de **idle** se referă la numărul cumulat de secunde în care procesul a fost inactiv (nu a executat nicio instrucțiune), și în cazul sistemele desktop, acesta poate să fie foarte mare, atingând chiar valori de peste 90% din uptime. În cazul unui server dedicat, această cifră este un indicator al eficienței de utilizare al acestuia, valori mai mici simbolizând un raport mai bun.
- **/proc/version** – informații despre versiunea de kernel ce rulează. Acest fișier concatenează informațiile din mai multe fișiere precum **/proc/sys/kernel/ostype** și **/proc/sys/kernel/osrelease**.
- **/proc/vmstat** – folosit pentru afișarea informațiilor legate de memoria virtuală.

Sistemul de gestiune a pachetelor

Unul dintre aspectele care sperie orice utilizator de calculatoare este gestiunea corectă a programelor: instalarea, dezinstalarea sau configurarea acestora, provoacă de multe ori probleme utilizatorilor, mai ales când aceștia nu au experiența cu fiecare program în parte.

În sistemele de tip UNIX, această problemă era chiar mai mare decât pe alte sisteme de operare, cel puțin în condițiile în care utilizatorul nu plănuia cu atenție ce și cum anume instalează, fiecare utilizator fiind forțat să compileze singur programele pe care dorea să le folosească. Acest detaliu însemna că în lipsa unui plan de desfășurare al software-ului instalat într-un calculator, utilizatorul putea instala programe de mai multe ori, cu versiuni diferite și în locuri diferite (sau cu suprascrieri parțiale), ceea ce făcea gestiunea corectă a acestora aproape imposibilă. Când un utilizator dorea să instaleze un program într-un sistem de tip UNIX, acesta de regulă achiziționa codul sursă al acestuia în format **tar.gz** (**tar** = **tape archive**, **gz** = **gzip**), adică o structură de fișiere și directoare, ce menținea informații despre sistemul de fișiere (permisiuni, date, și proprietari) care era arhivată pentru distribuirea ușoară. După acest pas, utilizatorul trebuia să extragă conținutul arhivei, și de regulă, să urmeze niște script-uri de compilare relativ automate, incluse în arhivă (prin intermediul comenzii **make**, care se ocupă atât de configurare cât și de compilarea programului). Dacă totul funcționa corect, și dependențele de alte programe și librării erau satisfăcute, programul era instalat - cu mențiunea că dacă utilizatorul ar fi vrut să facă o actualizare a acestui program la o versiune mai nouă, ar fi trebuit dezinstalat, și apoi reinstalat și reconfigurat.

Debian a fost gândit de la început să ofere o soluție pentru această problemă, și astfel a apărut un sistem complet de gestiune a programelor, unul care să ofere un sistem rapid și practic de instalare a pachetelor, abordând în mod automat problema dependențelor și a reconfigurărilor în cazul actualizărilor de versiuni. Astfel a apărut pachetul Debian (fișiere cu extensia **.deb**) și **apt** (Advanced Packaging Tool). Pe lângă toate aceste facilități, **apt** oferă și posibilitatea procurării pachetelor necesare din **repository**-uri (arhive de pachete), care pot fi surse locale (CD-ROM, hard-disc, etc) sau puse la dispoziție prin intermediul rețelei. Toate acestea fac din **apt** un excelent utilitar pentru gestiunea software-ului dintr-un sistem Linux.

Vom analiza în continuare compoziția pachetelor specifice distribuției Debian, fișierele **.deb**. Acestea conțin mai mult decât o serie de fișiere ce pot fi instalate pe sistemul nostru: pe lângă fișierele de bază acestea conțin și o informații ierarhizate despre dependențele pachetului (mai exact ce alte pachete trebuie instalate astfel încât cel în cauză să poată fi instalat). În baza acestor informații adiționale, sistemul responsabil de gestiunea pachetelor va ști ce pachete să descarce (din arhiva de pachete – repository) pentru ca instalarea să poată fi terminată cu succes. De altfel sistemul este capabil să decidă exact ordinea de instalare a dependențelor (cât și stabilirea acestora în funcție de pachetele deja instalate), cât și să țină evidența versiunilor. Vom continua prin a detalia conceptele de **dependințe** și **conflicte**.

Dependințele sunt în esență pachete necesare pentru instalarea unui alt pachet: dacă am lua trei pachete imaginare **a**, **b** și **c** ca exemplu, iar **c** ar avea ca dependențe pachetele **a** și **b**, atunci acesta nu ar putea fi instalat fără ca pachetele **a** și **b** să fie instalate în prealabil. Sistemul respectiv funcționează în ambele sensuri, atât pentru instalări dar mai ales pentru dezinstalări, când eliminarea unui pachet ar putea afecta multe altele.

Conflictele se referă la perechi de pachete, care instalate împreună ar putea cauza funcționarea anormală a unor programe. Din acest motiv, pachetele pot avea marcate alte pachete ca fiind **conflictuale**, și prin urmare acestea nu pot fi instalate în același timp.

Întreg sistemul este gestionat de programul **dpkg** (**Debian package management system**), existând mai multe interfețe pentru a facilita uzul acestei aplicații. Printre cele mai comune interfețe pentru **dpkg** putem aminti **apt-get** sau varianta mai nouă **aptitude** (tot o aplicație de consolă, însă care oferă o interfață mai ușor de folosit decât cea bazată pe linia de comandă).

Instalarea pachetelor într-un sistem Debian

Fișierele Debian pot fi instalate folosind atât utilitarul **dpkg**, cât și **apt-get**. În cazul în care utilizatorul dispune de un fișier Debian (.deb) în mod absolut (de exemplu **pachet.deb**) acesta ar putea dispune instalarea acestui pachet folosind comanda: **dpkg -i pachet.deb**

În cazul în care pachetul dorit nu există local pe sistemul în cauză, utilizatorul poate încerca să obțină acest pachet prin intermediul programului **apt-get** de la un server ftp, numit repository. Programul **apt-get** se va ocupa și de descărcarea automată a dependențelor, și tratarea oricăror conflicte ce pot să apară – ceea ce poate rezulta uneori în ștergerea unor pachete, cu o notificare foarte clară către utilizator. Fiind vorba de un utilitar în linie de comandă, sintaxa acestuia este una foarte simplă:

apt-get install nume_pachet

În urma executării comenzii anterioare, **apt** va determina o listă de pachete deja instalate pe sistem, pentru a putea stabili ce pachete lipsesc. Acest lucru este posibil prin interpretarea unui fișier esențial pentru sistemul Debian - **/var/lib/dpkg/status**. Compromiterea acestui fișier în orice fel, poate avea consecințe grave asupra unui sistem Debian, pornind de la incapacitatea de a folosi platforma de gestiune a pachetelor până la pierderea funcționalității sistemului de operare. Din acest motiv, sistemul ține o copie (ceva mai veche decât originalul, dar relativ similară) în **/var/lib/dpkg/status-old**.

În urma stabilirii dependențelor lipsă, **apt** va descărca atât acele pachete, cât și pachetul dorit și va continua instalarea acestora, în ordinea relevantă. Toate pachetele descărcate prin intermediul utilitarului **apt** vor fi stocate în directorul **/var/cache/apt/archives**. Sistemul de gestiune al pachetelor este gândit astfel încât după descărcarea și instalarea unui pachet, acesta să rămână stocat în sistem la adresa menționată anterior, pentru a nu forța descărcarea unui pachet în mod inutil în cazul instalărilor multiple. Ștergerea pachetelor are loc periodic, dar poate fi forțată de către utilizator fie prin

ștergerea manuală a conținutului directorului menționat anterior, fie prin rularea comenzii **apt-get clean**.

Arhive de pachete (repository)

Pentru a permite accesul facil la pachete comune, programul **apt-get** se folosește de arhive de pachete (numite repository). Acestea pot să varieze de la resurse locale (hard-disc, cdrom sau alte medii locale) la servere ftp de pe internet, acestea din urmă fiind și cea mai utilizată variantă, din moment ce de regulă conțin versiuni actualizate ale tuturor pachetelor dorite.

Lista de repositories poate fi modificată, putând fi adăugate sau șterse astfel de arhive, sistemul menținând un fișier text (**/etc/apt/sources.list**), cu un format bine stabilit în care aceste arhive de pachete sunt menționate. Este important de precizat faptul că acest fișier conține doar referințe către noi surse de software, și în cazul serverelor ftp, este necesară o conexiune fizică cu acestea.

```
user1@statie:~$ less /etc/apt/sources.list
#
# deb cdrom:[Debian GNU/Linux 5.0.3 _Lenny_ - Official i386 DVD Binary-1 20090908
5-08:48]/ lenny contrib main

deb cdrom:[Debian GNU/Linux 5.0.3 _Lenny_ - Official i386 DVD Binary-1 20090908-
08:48]/ lenny contrib main

deb http://ftp.ro.debian.org/debian/ lenny main
deb-src http://ftp.ro.debian.org/debian/ lenny main

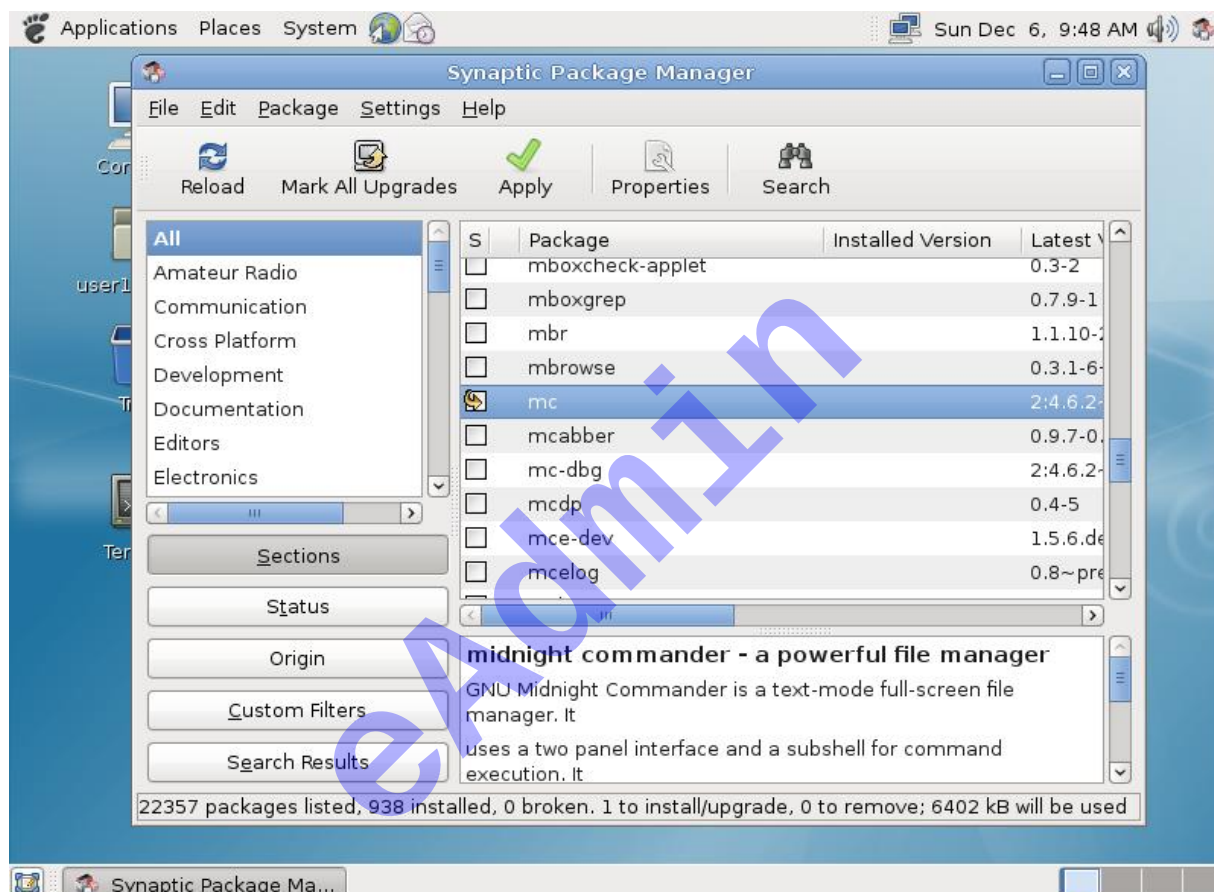
deb http://security.debian.org/ lenny/updates main contrib
deb-src http://security.debian.org/ lenny/updates main contrib

deb http://volatile.debian.org/debian-volatile lenny/volatile main contrib
deb-src http://volatile.debian.org/debian-volatile lenny/volatile main contrib
(END)
```

Analizând imaginea de mai sus putem observa un exemplu de cum ar putea să arate fișierul **sources.list**. Așa cum am menționat, acesta conține o listă cu referințe către toate arhivele de pachete de pe care am putea să instalăm ceva. Sistemul apt folosește o bază de date proprie pentru a păstra o listă cu toate pachetele disponibile pentru descărcare (inclusiv versiunile disponibile pentru acestea) și generează această bază de date în funcție de arhivele selectate prin intermediul fișierului **sources.list**. Deși acest proces este automat într-o oarecare măsură, se recomandă forțarea actualizării acestei baze de date ori de câte ori se adaugă un nou repository, sau când se încearcă instalarea unui nou program, astfel încât să ne asigurăm că pachetul dorit există în lista de pachete disponibile, actualizat la cea mai nouă versiune disponibilă. Pentru a forța sistemul de gestiune să reactualizeze această bază de date, trebuie folosită comanda: **apt-get update**

Pe lângă varianta instalării pachetelor folosind consola sistemului, utilizatorul poate opta pentru folosirea interfeței grafice, prin intermediul programului **Synaptic Package Manager**. Acesta nu este altceva decât o interfață pentru sistemul **apt**, și oferă în esență aceleași funcții: instalarea, ștergerea,

căutarea sau actualizarea pachetelor din sistem. Principala diferență dintre acesta și utilizarea în consolă este un sistem în care schimbările pot fi adăugate într-o coadă de comenzi, urmând să fie executate într-o ordine anume. Principalele avantaje ale acestui program, sunt ușurința cu care poate fi folosit de către oricine, mesaje clare (utilizatorul este atenționat foarte clar ce efect va avea orice schimbare), afișarea modificărilor altor pachete ca urmare a deciziilor luate și nu în ultimul rând facilitățile avansate de căutare a pachetelor. **Synaptic Package Manager** vine instalat în mod implicit cu Debian, și poate fi lansat foarte ușor prin accesarea meniului **System -> Administration -> Synaptic Package Manager**.



Ca și în cazul altor utilitare ce afectează sistemul la nivel global, pentru lansarea acestui program utilizatorul va fi rugat să furnizeze parola unui utilizator autorizat, în cazul de față fiind vorba de contul utilizatorului **root**. Utilizatorul poate să adauge comenzi de instalare, ștergere sau actualizare individual pentru pachete, însă schimbările nu vor fi procesate decât după ce acesta acționează butonul **Apply**. Ca un exercițiu, putem încerca să instalăm programul **mc (Midnight Commander)**, vizibil în poza de mai sus. Acesta este un program de gestiune a fișierelor în mod text, similar cu **Norton Commander** și oferă multe facilități pentru o gestiune mai ușoară a sistemului de fișiere. Pe lângă rolul său primar, programul mai conține și o serie de alte utilitare precum **mcedit**, un editor de text foarte util și foarte ușor de folosit.

Pentru a instala pachetul **mc**, utilizatorul trebuie să localizeze mai întâi pachetul respectiv, fie parcurgând lista de pachete (vizibilă în partea din dreapta a ecranului), fie folosind facilitățile de căutare de pachete. În cazul în care utilizatorul folosește funcția de căutare, acesta trebuie să specifice o serie de parametri: un șir de caractere ce reprezintă textul căutat, și câmpul în care programul să efectuează căutarea (opțiunile sunt destul de clare, însă dat fiind faptul că de regulă se caută nume de pachete, se recomandă căutarea după nume – și nu după descriere – pentru a reduce numărul de rezultate).

După ce utilizatorul a identificat cu succes pachetul dorit, acesta poate să acționeze butonul din stânga al mouse-ului peste pachetul respectiv și să selecteze una din opțiunile disponibile:

- **Mark for Installation** – pachetul va fi marcat pentru instalare în coadă de comenzi. În urma acționării butonului **Apply**, se va încerca instalarea acestuia și a tuturor dependențelor acestuia.
- **Mark for Reinstallation** – va marca pachetul pentru a fi deinstalat și instalat din nou.
- **Mark for Upgrade** – se va încerca actualizarea pachetului respectiv la o versiune mai nouă, în cazul în care acesta este disponibilă într-una din arhivele de pachete menționate.
- **Mark for Removal** – marchează pachetul pentru ștergerea acestuia din sistem.
- **Mark for Complete Removal** – spre deosebire de varianta precedentă, aceasta va elimina complet pachetul cât și eventuale fișiere de configurație.

Utilizatorul poate înșirui mai multe comenzi, iar **Synaptic** se va asigura că acestea vor fi executate în ordinea corespunzătoare și va avea grijă să instaleze și să afișeze numele tuturor dependențelor ce vor fi instalate. Este important de reamintit faptul că nici o schimbare nu va fi dusă la capăt până când utilizatorul nu acționează butonul **Apply**.

Ștergerea unui pachet

Pe lângă instalarea pachetelor, administratorul unui sistem Debian trebuie să poată și să le elimine din sistem. Pentru ștergerea pachetelor se poate folosi tot interfața **apt-get** cu o serie de parametri: **apt-get remove nume_pachet**

În urma executării comenzii de mai sus, sistemul **apt** va afișa o serie de informații esențiale precum o atenționare, lista tuturor pachetelor ce vor fi șterse (**apt** se va ocupa automat și de înlăturarea tuturor pachetelor ce depind de pachetul pe care încercăm să-l ștergem, utilizatorul neavând altă opțiune), cât și informații generale precum numărul de pachete ce vor fi instalate, șterse sau actualizate și spațiul eliberat în urma acțiunii. Ultimul rând este de regulă o cerere de confirmare din partea utilizatorului – de regulă în formatul **Do you want to continue [Y/n]** -, specificând dacă acesta dorește cu adevărat executarea comenzii.

OBSERVAȚIE: Urmând un fel de tradiție în Linux, marea parte a confirmărilor de tipul de mai sus, pot avea una, două sau mai multe opțiuni. Printre opțiunile prezente în exemplul de mai sus, putem observa faptul că litera **Y** (simbolizând yes) este scrisă cu majuscule, pentru a semnală opțiunea implicită; dacă utilizatorul nu specifică altă opțiune, cea implicită va fi selectată automat la apăsarea tastei ENTER.

Exemplul de mai sus poate fi tradus foarte simplu în programul **Synaptic Package Manager** prin selectarea opțiunii **Mark for Removal**. Pentru specificarea ștergerii complete a pachetului (adică inclusiv a fișierelor de configurație, în cazul în care acestea există), trebuie adăugat parametrul **--purge**: **apt-get --purge remove nume_pachet**

Căutarea și afișarea informațiilor referitoare la pachete

Urmărind arhitectura prezentată până acum, sistemul de gestiune al pachetelor funcționează prin folosirea unei baze de date proprii cu informații despre pachetele disponibile, aceasta putând fi reactualizată prin comanda **apt-get update**. Programul folosit pentru menținerea bazei de date interne actuală este **apt-cache**, acesta permițând utilizatorilor (prin intermediul consolei) să afișeze și să caute în lista de pachete disponibile. De regulă, principalul motiv pentru care se dorește căutarea unui pachet este atunci când utilizatorul nu știe exact numele pachetului, sau pur și simplu dorește afișarea tuturor pachetelor ce conțin unul sau mai mulți termeni în descriere.

Căutarea se poate face foarte simplu folosind programul **apt-cache** cu parametrul **search**: **apt-cache search sir_de_căutare** unde șirul de căutare reprezintă o expresie regulată (REGEX). În urma executării acestei comenzi, utilitarul va afișa toate pachetele ce conțin expresia căutată fie în numele pachetului, fie în descrierea acestuia. Informațiile afișate pot fi de multe ori derutante, doar prin prisma numărului de rezultate găsite, însă există metode pentru reducerea numărului de rezultate.

O altă metodă de căutare poate implica căutarea după mai multe expresii; abordarea poate să difere, fie prin folosirea expresiilor regulate (REGEX – accesați comanda **man regex** pentru mai multe informații), fie prin filtrarea rezultatului intermediar cu ajutorul comenzii **grep**: **apt-cache search mc | grep midnight**

În exemplul de mai sus putem observa prima parte a comenzii, care va efectua o căutare după șirul de caractere „mc”. Chiar dacă în cazul de față, dorim să aflăm informații doar despre aplicația Midnight Commander (prezentată mai sus), **apt-cache** va afișa numele și descrierile tuturor pachetelor ce conțin șirul „mc” ducând astfel la un număr foarte mare de rezultate. Partea a doua va filtra rezultatul intermediar (cu ajutorul utilitarului **grep**); astfel se vor afișa doar pachetele care conțin atât textul „mc” cât și textul „midnight”. Această modalitate de căutare poate fi extinsă în mai multe feluri, de exemplu, pentru a căuta o listă cu pachetele editoarelor de text disponibile: **apt-cache search editor | grep text**

Odată știut numele pachetului dorit, utilizatorul poate afla în continuare mai multe informații despre acestea, tot prin intermediul programului **apt-cache**. Acesta poate de exemplu să solicite informații detaliate despre un pachet, precum versiunea, dimensiunea, arhitectură, conflictele sau chiar dependențele acestuia: **apt-cache show mc**

Comanda de mai sus va afișa toate informațiile disponibile despre pachetul „mc”. A se observa faptul că în loc de directiva **search**, s-a folosit directiva **show** pentru a specifica faptul că utilizatorul dorește informații, și nu o căutare. Comanda poate fi folosită alternativ cu utilitarul **less** pentru ca rezultatul să fie mai ușor de citit.

Alternativ, utilizatorul poate să afișeze dependențele unui pachet (acestea fiind menționate și în secțiunea de informații adiționale a pachetului) într-o formă mai ușor de citit folosind directiva **depends** în loc de **show**.

Căutarea pachetelor în linie de comandă poate fi evitată în cazul în care mediul de lucru beneficiază de o interfață grafică. Programul **Synaptic Package Manager** prezentat anterior, poate fi folosit cu ușurință de către aproape oricine, fiind foarte ușor să se afle informații adiționale precum dependențe sau conflicte.

Actualizarea surselor de pachete (Repositories)

Sistemul de gestiune al pachetelor **apt** folosește surse de pachete pentru descărcarea acestora, pentru a putea fi instalate, indiferent de mediul pe care care se află acestea (extern sau intern). Principalul avantaj al acestei infrastructuri este că atât utilizatorii cât și dezvoltatorii de software pot apela la o sursă comună de programe. Această abordare nu este utilă numai pentru descărcarea inițială a unui sau mai multor programe, ci și pentru păstrarea acestora la versiuni actualizate; un dezvoltator se va asigura de regulă că pachetele sale sunt prezente fie pe servere ftp proprii, fie pe arhive comune, facilitând astfel accesul către acestea.

Arhitectura folosită este foarte flexibilă, permițând astfel utilizatorului să adauge și să elimine astfel de surse de pachete după nevoi. Lista de arhive poate fi alterată de utilizator pentru a adăuga sau a scoate anumite surse în două feluri: fie alterând fișierul ce conține lista de arhive (**/etc/apt/sources.list**) manual cu un editor de text, fie folosind un utilitar – care în esență face același lucru, însă oferă o interfață mai ușor de folosit.

Formatul fișierului **/etc/apt/sources.list** este prestabilit, fiecare sursă de pachete fiind menționată pe o linie, și este gândit să suporte o gamă foarte largă de medii și surse. Fiecare linie descrie așadar o sursă de pachete și respectă următoarea sintaxă: **deb URI distribuție [componenta 1] [componenta 2] ...**

Primul cuvânt poate fi **deb** sau **deb-src**, și descrie fie o arhivă cu două nivele tipică distribuției Debian, fie o arhivă ce conține codul sursă în aceeași formă ca tipul **deb**.

URI specifică baza distribuției, sau mai exact locația de unde apt va putea descărca pachetele necesare. Aceasta poate fi și sub forma unei adrese (de internet de exemplu), caz în care componentele menționate trebuie omise și distribuția trebuie terminată cu ajutorul caracterului „/”. Principalele tipuri de surse recunoscute sunt următoarele:

- file – permite folosirea unui fișier oarecare din sistem drept o arhivă.
- cdrom – permite folosirea suportului CDROM pentru surse.
- http – folosește un server HTTP drept arhivă
- ftp – specifică un server FTP drept arhivă

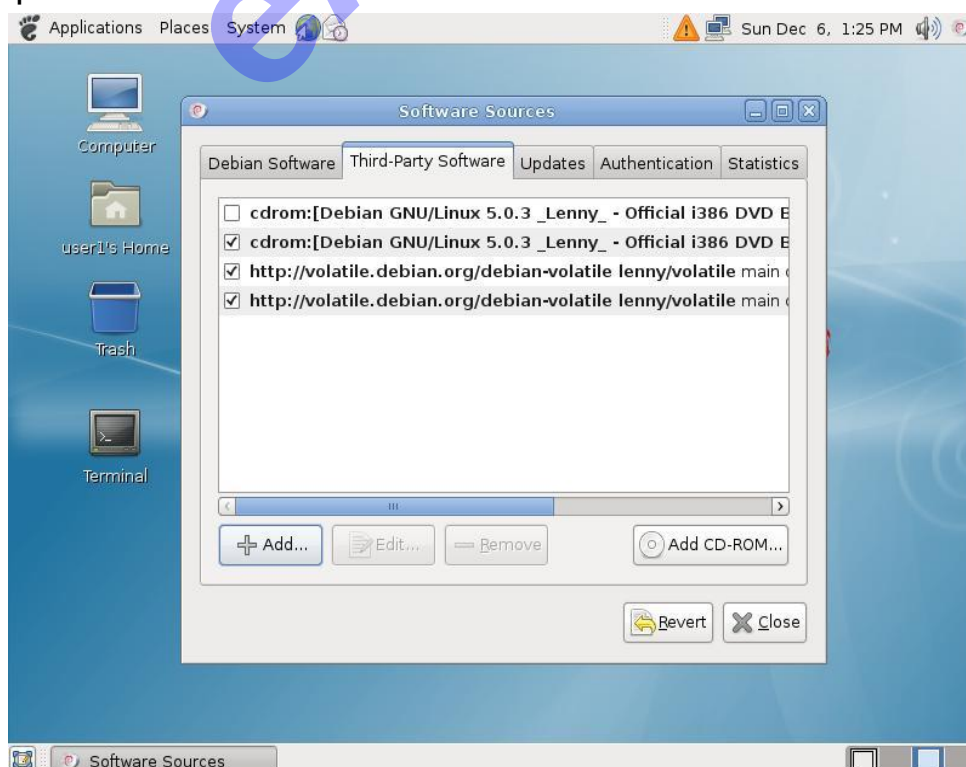
Pentru a clarifica tipurile menționate anterior, putem da câteva exemple:

Linia **deb http://archive.debian.org/debian-archive main** va menționa că sursă de pachete serverul HTTP de la adresa menționată, însă folosind doar zona **main**.

Linia **deb ftp://ftp.debian.org/debian stable main** va folosi ca sursă serverul ftp de la adresa menționată, sub directorul **/debian**, accesând doar zonele **stable** și **main**.

Pentru mai multe informații, puteți accesa pagina de manual corespunzătoare comenzii apt prin intermediul comenzii **man sources.list**. Ca observație, utilizatorul trebuie să aibă grijă mereu de unde descărca pachete, fiindcă folosind un repository necunoscut acesta poate permite instalarea unor pachete compromise pe sistem. De asemenea, orice modificare a fișierului **sources.list** trebuie să fie urmată de comandă **apt-get update** astfel încât baza de date folosită de apt să poată fi actualizată.

Alternativ, pe lângă editarea manuală a fișierului **sources.list** putem folosi programul **Software Sources** (accesibil prin meniul **System -> Administration -> Software Sources**) care oferă o interfață mult mai ușor de folosit pentru modificarea listei de arhive.

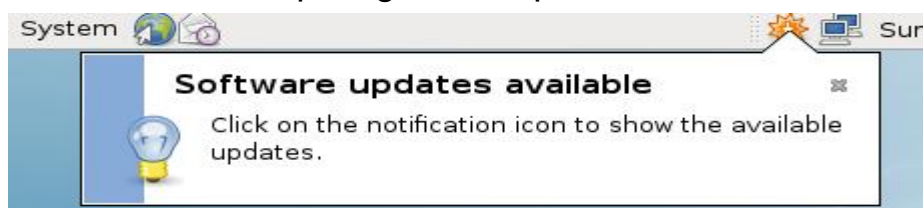


Observăm în lista din poză de mai sus, o înșiruire a tuturor surselor care sunt incluse de către apt, mai exact, vedem o variantă puțin interpretată a fișierului **sources.list**. Fiecare linie poate fi activată sau dezactivată prin bifarea, respectiv debifarea acesteia și se traduce prin adăugarea caracterului „#” la începutul liniei respective în cadrul fișierului **sources.list**. Acest caracter simbolizează existența unui comentariu în cadrul fișierului, ducând la ignorarea liniei respective. Mecanismul respectiv este foarte util pentru a nu forța utilizatorul să piardă astfel de surse, oferindu-i în același timp posibilitatea excluderii acestora. Pe lângă dezactivarea unui repository, mai există și opțiunea de adăugare de noi arhive prin acționarea butonului **Add**, fiind necesară doar linia apt, ce descrie sursa (formatul acesteia este cel prezentat mai devreme, în cadrul fișierului **sources.list**).

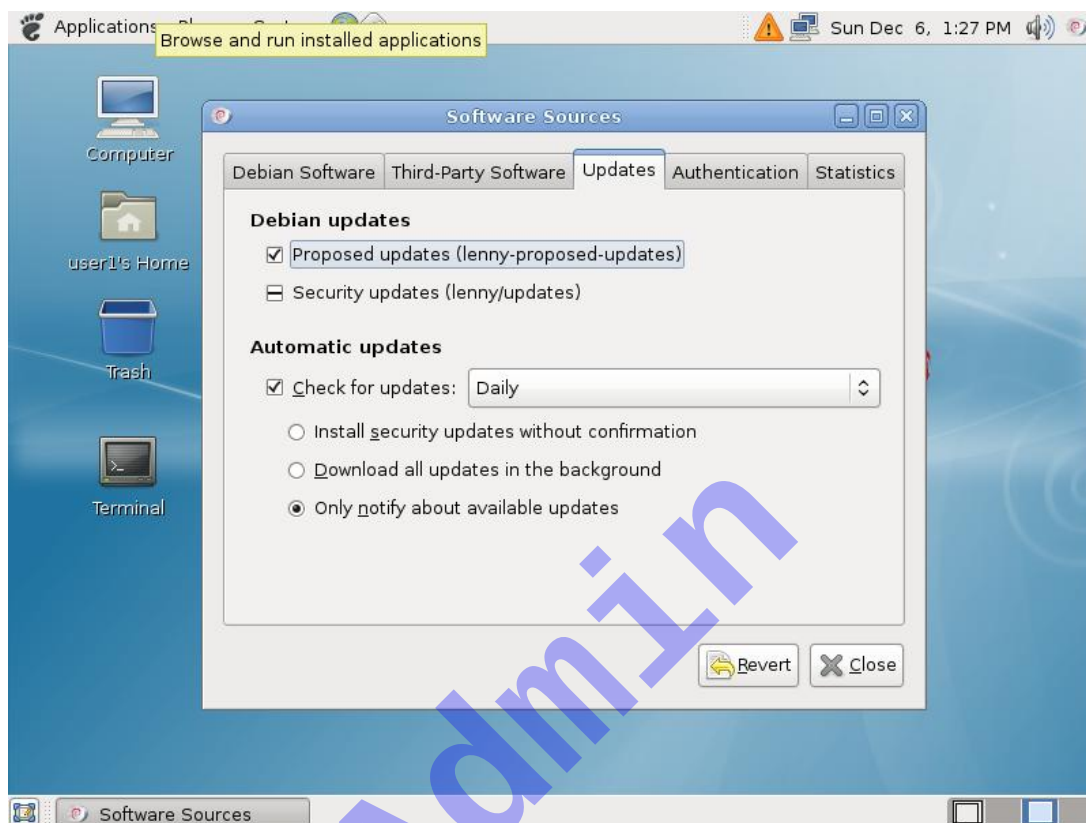
Actualizarea sistemului (Update)

Actualizarea versiunilor programelor ce rulează pe un sistem la un moment dat poate fi o sarcină dificilă chiar și pentru utilizatorii experimentați: aceștia trebuie să monitorizeze permanent o serie de producători pentru a afla când vor lansa o versiune mai nouă (și în general îmbunătățită) a aplicațiilor pe care le folosesc. Monitorizarea poate fi dificilă, dat fiind faptul că pot interacționa multe entități cu politici diferite: unii dezvoltatori de aplicații pot lansa versiuni noi zilnic, pe când alții o pot face și la intervale de ani de zile. Pe lângă această problemă, administratorii mai au și sarcina de a menține sistemul în stare funcțională cât mai mult timp posibil, iar actualizarea versiunii unui serviciu – precum cel de http – presupune mai întâi oprirea serviciului respectiv, și apoi actualizarea sa. Combinată cu alte sarcini pe care administratorii le pot avea în cadrul unui sistem Linux, pot transforma această sarcină într-una dificilă, dar esențială – mare parte din îmbunătățirile versiunilor mai noi se referă la neajunsuri anterioare sau vulnerabilități de securitate care nu pot fi ignorate.

Continuând această idee, s-a dezvoltat – inițial pentru distribuția Ubuntu, apoi inclusă și în Debian – aplicația **Update Manager**, care are rolul de a automatiza procesul de actualizare a programelor instalate într-un sistem Linux. Aceasta își îndeplinește rolul printr-un daemon de fundal, și poate atenționa utilizatorul atunci când apar noi versiuni ale pachetelor existente, și opțional le poate descărca și instala. **Update Manager** este ca și celelalte aplicații de gestiune a pachetelor, o interfață pentru sistemul APT, însă nu poate fi folosit pentru adăugarea sau ștergerea pachetelor. Programul rulează automat și discret, în distribuțiile mai recente Debian, necesitând un minim de interacțiune cu utilizatorul, însă poate fi lansat accesând meniul **System -> Administration -> Update Manager**. Utilizatorul este atenționat mereu prin intermediul unei pictograme în panoul de sistem:



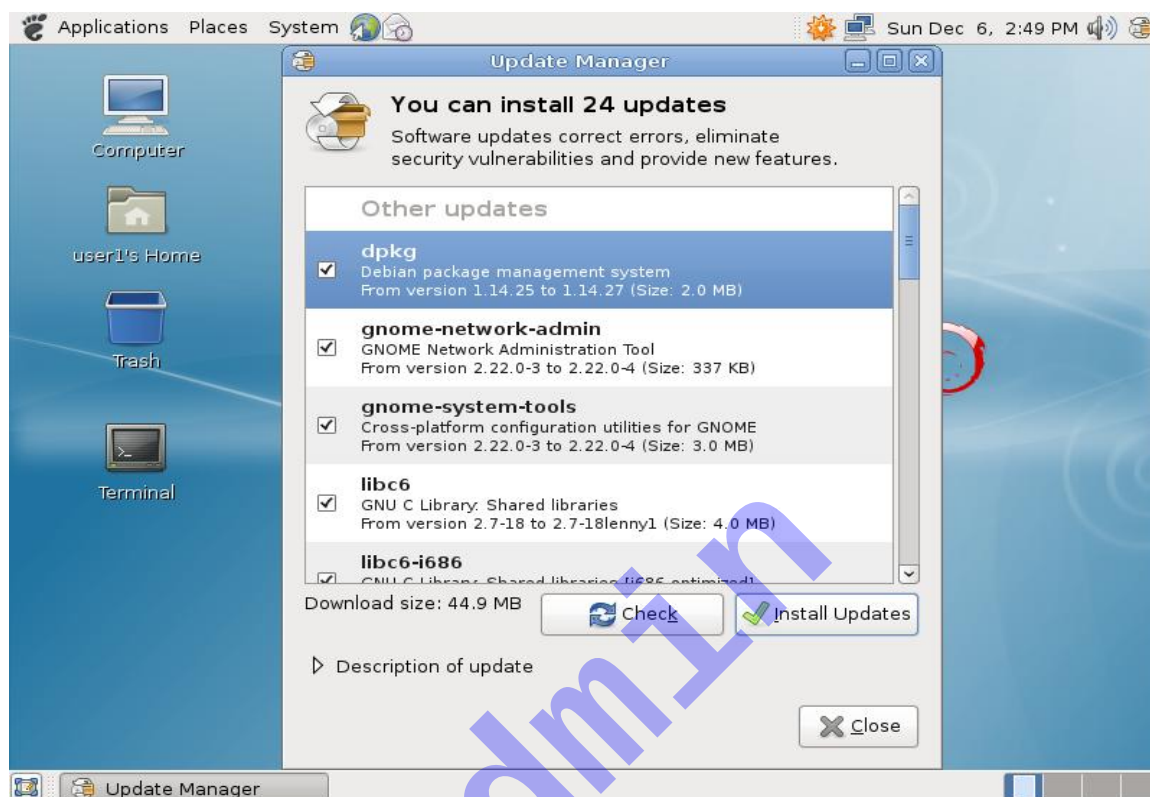
Este important de știut faptul că, precum și celelalte programe bazate pe sistemul APT, **Update Manager** detectează existența unor noi versiuni și descarcă, precum și instalează noile pachete folosind APT. Mai mult, pentru a căuta noi versiuni, acesta trebuie să aibă acces la arhivele de pachete menționate în fișierul **sources.list**, sau prin cadrul programului **Software Sources**.



În imaginea precedentă, putem observa câteva dintre opțiunile pe care **Software Sources** ni le oferă în relația cu sistemul automat de actualizare:

- **Proposed updates** - în esență, pachetele noi trebuie să treacă printr-o perioadă de testare ceva mai lungă până să fie incluse în arhivele majore. Pachetele marcate ca fiind „proposed” sunt de regulă ceva mai noi, și nu au beneficiat de întregul mecanism de testare. Ca atare, o primă consecință a bifării acestei opțiuni va fi aceea că sistemul va beneficia de mai multe actualizări, mai repede. Deși pentru un sistem normal de tip „desktop” această opțiune poate fi folosită cu încredere, nu se recomandă bifarea în cazul unui server aflat în producție.
- **Security updates** – marchează pentru verificare, descărcare și instalare, toate actualizările referitoare la securitate. Această opțiune este una foarte importantă, deoarece se va ocupa de repararea oricăror probleme de securitate și nu poate fi dezactivată prin mijloace convenționale.
- **Check for updates** – setează intervalul la care daemon-ul responsabil cu actualizarea să caute pachete noi.
- **Install security updates without confirmation** – orice versiune nouă de pachet va fi instalat în mod automat, fără nici un fel de confirmare din partea utilizatorului.

- **Download all updates în the background** – va specifica faptul că utilizatorul dorește să descarce noile pachete în fundal
- **Only notify about available updates** – este și modul implicit de funcționare, în care utilizatorul va fi atenționat de existența unor noi pachete, acesta putând să aleagă ulterior dacă dorește să le instaleze sau nu.



Utilizatorului îi este prezentată o fereastră în care poate obține informații detaliate despre actualizarea fiecărui pachet, prin selectarea pachetului dorit, și vizualizarea secțiunii „**Description of update**”, localizată în partea de jos a ferestrei. Pe lângă aceste informații generale, mai se precizează și dimensiunea totală ce va fi descărcată. Pachetele pot fi actualizate preferențial, prin bifarea sau debifarea fiecărei intrări din lista vizibilă și nu se va efectua nici o schimbare asupra sistemului până nu se acționează butonul „**Install Updates**”.

OBSERVAȚIE: Actualizarea oricărui program presupune întreruperea acestuia în cazul în care rulează. Ca urmare imediată, se recomandă salvarea oricăror documente, și oprirea tuturor programelor pe toată durata actualizării.

Actualizarea pachetelor se poate face și fără existența unei interfețe grafice, prin intermediul utilitarului **apt-get**, folosind directiva **upgrade**. Se recomandă folosirea parametrului **-u** pentru afișarea informațiilor despre pachetele ce vor fi actualizate – în caz contrar, aceasta făcându-se în fundal:

apt-get -u upgrade

O altă caracteristică foarte importantă a sistemului de gestiune APT este aceea că permite actualizarea întregii distribuții, prin executarea comenzii:

apt-get dist-upgrade

Procese Linux

Așa cum am menționat anterior, sistemul Linux este un sistem de operare multi-user cu capacități multi-tasking. Dar ce înseamnă de fapt acești doi termeni? Pentru a răspunde la această întrebare, aceștia trebuie analizați.

Termenul multi-user se referă în mod concret la capacitatea accesării unui sistem informatic sau al unui program de către mai mulți utilizatori concurenți. Putem trage concluzia că Linux este un sistem de operare multi-user din moment ce permite utilizarea resurselor disponibile de către mai mulți utilizatori simultani, însă astfel se ivește o nouă problemă: nevoia de multi-tasking.

Spre deosebire de alte sisteme de operare tip DOS (Disk Operating System), precum MS-DOS, care sunt single-user, single-tasking, mai precis nu permit accesul mai multor utilizatori simultani, și nici execuția în paralel a mai multor programe, o platformă multi-user ar avea nevoie de capacități multi-tasking pentru a opera eficient, altfel utilizatorii conectați la sistem ar trebui să își desfășoare activitatea pe rând, unul după altul. Deși DOS este în esență un sistem single-tasking (adică nu permite execuția mai multor programe simultan) și numai un singur program poate fi încărcat în memorie la un moment dat, există mecanisme de evitare a acestei probleme: TSR sau Terminate and Stay Resident.

Acest mecanism forța sistemul de operare să păstreze un program în memorie și să modifice vectorul de întreruperi al sistemului astfel încât acesta să poată fi apelat chiar și după terminarea sa, și a fost considerat o formă primitivă de multi-tasking, reprezentând o platformă pentru construirea unor driveri ce nu erau suportați nativ de sistem, asigurând totuși funcționarea sistemului. Dezavantajul acestui sistem era faptul că programele TSR trebuiau păstrate într-o zonă de memorie foarte limitată – primii 640KB de memorie convențională, rezultatul fiind reprezentat atât de numeroase conflicte cu alte programe de tip TSR cât și lupta continuă de a micșora dimensiunea acestora astfel încât să păstreze cât mai multă memorie liberă.

Deși o soluție simplă la nivelul sistemelor de atunci, TSR nu este nici pe departe un sistem adevărat de gestionare a proceselor, cum este cel oferit de kernel-ul Linux.

În realitate, conceptul de multi-tasking este de fapt mult mai complicat și poate fi abordat în multe feluri, fiind de fapt execuția mai multor procese individuale folosind o resursă comună – un procesor. În cazul unui calculator cu un singur procesor, este imposibil ca acesta să execute două instrucțiuni diferite în exact același moment în timp, însă multe sisteme de operare moderne permit execuția simultană a mai multor aplicații - precum vizualizarea paginilor de interne în timp ce utilizatorul asculta muzică – folosind un singur procesor. Mult simplificat, mecanismul poate fi explicat prin programarea în timp a permisiunii de execuție a fiecărui proces în parte – fiecare fiind executat pe rând de către procesor pentru un timp stabilit. Prin

alocarea timpului de procesor între procese diferite la intervale foarte mici, se crează iluzia de execuție în paralel – însă, realitatea este că părți din fiecare program sunt executate secvențial.

Există mai multe implementări ale acestui sistem de partajare a timpului de execuție, însă un sistem de gestiune al întregului proces este esențial. Primele tentative de execuție în paralel se numeau „multiprogramming”, și au apărut ca urmare a nepotrivirii dintre viteza mare de execuție a unui procesor comparativ cu echipamentele periferice. Dezavantajul major al acestui sistem era însă faptul că nu se putea garanta în niciun fel execuția în timp util al unei aplicații, maximizând însă gradul de utilizare al procesorului.

O altă variantă de multi-tasking, ceva mai evoluată decât precedentă, este modelul time-sharing (sau multi-tasking cooperativ), unde utilizatorul putea executa programe diferite simultan, ca și când acestea erau singurele care rulau de pe sistem. Printre sistemele de operare care folosesc acest mecanism se numără primele variante de Windows și Mac OS. Prin acest model, fiecare program în execuție pe sistem se angajează să acorde timp de proces și celorlalte procese, dezavantajul major fiind acela că în cazul execuției unei aplicații prost gândite, întregul sistem se poate bloca.

Pentru a sumariza cele spuse, putem concluziona că un proces reprezintă un program în execuție și are atașate o serie de informații specifice precum instrucțiunile programului, resurse folosite (precum fișiere deschise), unul sau mai multe fire de execuție și alte informații necesare procesului de execuție în paralel.

Caracteristicile proceselor

Orice proces Linux va avea un set de caracteristici comune, ce oferă informații despre acesta:

- **PID** – sau Process ID, este un identificator de proces sub forma unui număr întreg unic.
- **PPID** – similar cu PID, cu excepția că reprezintă identificatorul procesului care a dat naștere procesului curent (cunoscut și ca proces părinte)
- **Terminalul atașat** – prescurtat și TTY, reprezintă terminalul la care procesul curent este legat.
- **RUID** – Real User ID, reprezintă identificatorul utilizatorului care a lansat aplicația. Similar există și EUID (sau Effective User ID) pentru identificarea drepturilor reale la resursele sistemului.
- **RGID și EGID** – similar cu RUID și EUID, doar că se referă la identificatorul grupului de utilizatori.
- **factorul nice** – folosit pentru a determina, așa cum sugerează și numele, „factorul de prietenie” al procesului cu scopul stabilirii priorității de execuție (bazat pe factorul nice și istoricul de utilizare al procesorului)

Procesele pot rula în mod normal în două moduri: **foreground** respectiv **background**. Dintre procesele interactive, marea majoritate funcționează în modul foreground. Diferența dintre cele două moduri se rezumă strict la interacțiunea cu interfața cu utilizatorul, mai exact la ocuparea consolei din

care respectivul proces a fost lansat. În acest mod, consola nu va fi eliberată până la terminarea programului care o ocupă. În cazul comenzilor prezentate până acum, vorbim tot despre procese ce rulează în foreground – de exemplu, și comanda **ls** face parte din această categorie, însă datorită timpului foarte scurt de rulare acest fapt este doar o iluzie. Un contra exemplu ar fi execuția comenzii **less** care nu va termina execuția odată ce a terminat afișarea pe ecran, ci va aștepta în permanență comenzi de la utilizator.

Ca alternativă, programele pot rula și în modul **background** pentru a permite terminalului de pe care rulează să primească și alte comenzi în același timp, execuția având loc în fundal și fiind guvernată de aceleași reguli ca orice alt proces. Acestea sunt procesele care de regulă au un timp de execuție mare sau nu necesită interacțiunea cu utilizatorul. Serviciile ce rulează pe un calculator fac și ele parte din această categorie, însă există câteva diferențe.

Dacă interacțiunea cu programe care rulează în foreground este evidentă (acestea de multe ori așteptând gesturile utilizatorului), interacțiunea cu programele ce rulează pe fundal (background) poate fi ceva mai dificilă. Pentru a putea gestiona aceste procese, shell-ul ne pune la dispoziție o serie de utilitare:

- **comanda jobs** va afișa toate comenzile ce rulează în fundal. Fiecare proces ce rulează în fundal va avea asociat un număr întreg reprezentând identificatorul și poate fi folosit pentru referențierea procesului respectiv. Acest identificator este diferit de câmpul PID (sau process identifier) și se referă doar la consola curentă.
- comanda **kill** este folosită pentru terminarea oricărui proces, specificat prin menționarea PID-ului procesului țintă ca parametru.
- comanda **fg %n** va aduce procesul cu identificatorul **n** înapoi pe fundal.
- similar, comanda **bg %n** va reactiva un proces suspendat în fundal.

Procesele pot fi afectate și prin combinații de taste precum **CTRL+Z** care va suspenda (dar nu termină) procesul ce rulează în foreground. Similar, combinația **CTRL+C** este foarte cunoscută pentru funcția sa de terminare a procesului curent.

Orice comandă executată în consolă va rula implicit în modul foreground, însă pentru a forța o comandă să ruleze în fundal trebuie adăugat caracterul **&** după comandă:

```
user1@computer1:~$ top &  
[1] 3180  
user1@computer1:~$ jobs  
[1]+  Stopped                  top  
user1@computer1:~$ _
```

În exemplul de mai sus observăm cum prin apelarea unui program ce nu își termină execuția decât la comanda utilizatorului (**top**) împreună cu **&** acesta este trecut pe fundal. Prin apelarea comenzii **jobs** se vor afișa o serie de informații despre procesele din fundal.

Servicii

O categorie foarte importantă în orice sistem de operare modern, serviciile oferă diverse funcționalități fie local, fie unei rețele de calculatoare. Acestea sunt ca orice alt program într-un sistem Linux, cu mențiunea că sunt denumiți **Daemons** și că rulează în continuu, de regulă pornind odată cu sistemul. Programele de acest tip rulează de regulă în fundal și așteaptă cereri către serviciile pe care le oferă, însă despre servicii vom vorbi pe larg în capitolele următoare.

Ciclul de viață al proceselor

Ca orice entitate dinamică, procesele au un ciclu de viață foarte bine definit, acestea nascându-se, evoluând și terminându-se după reguli bine puse la punct. Nevoia creării unui proces nou este evidentă: acest lucru se întâmplă ori de cât ori este nevoie să fie instantiat un nou program.

Procesele în Linux, urmează o structură de organizare foarte bine definită: aceea de arbore, cu procese ce reprezintă rădăcina arborelui, nodurile și frunzele acestuia. Procesele în cadrul sistemului Linux pot fi create numai din alte procese, acestea devenind **proces fii** (sau **children**). Urmând structura de arbore, fiecare proces poate fi reprezentat printr-un nod, rădăcina fiind părintele tuturor proceselor create. Nodurile din arborele de procese, pot avea la rândul lor alte procese de tip fiu, creându-se astfel o structură ierarhică foarte bine definită.

Crearea unui proces implică apelarea mecanismului numit **fork**, prin care acesta crează o copie identică a sa, însă cu un **PID** diferit. După terminarea mecanismului de forking, spațiul de adresare al procesului este înlocuit cu cel reprezentativ pentru noul proces (procedeu numit **exec**) și astfel, noul proces este gata de execuție. În ceea ce privește execuția propriu-zisă, o practică comună este aceea de transformare în **daemon** a proceselor fiu ale unui program, pentru ca acestea să poată rula chiar și după terminarea procesului părinte. Cazuri excepționale pot avea loc atunci când un proces fiu își termină execuția, în timp ce procesul părinte nu așteaptă acest lucru; acest tip de procese se numesc procese **zombie**.

Terminarea execuției unui proces în mod normal (dacă nu este terminat forțat) are loc printr-un mecanism ce returnează un rezultat procesului părinte (numit **exit status**). Acest rezultat reprezintă un număr, ce poate oferi indicii despre modul în care procesul și-a derulat activitatea, mecanismul fiind împrumutat din limbajul **C**, unde orice funcție poate returna un rezultat. Codurile de ieșire pot fi interpretate de către procesul tată și este specific fiecărui program. În momentul în care execuția unui proces tată este terminată, procesele sale fii pot fi terminate sau nu, în funcție de natura acestora, însă este bine de știut că nu există un o regulă fixă.

Afișarea structurii ierarhice în care sunt organizate procesele poate fi făcută în mai multe feluri, probabil cea mai comună fiind apelarea comenzii **ps tree**:

```

user1@computer1:~$ pstree
init--NetworkManager--{NetworkManager}
      --NetworkManagerD
      --acpid
      --atd
      --avahi-daemon--avahi-daemon
      --bonobo-activati--{bonobo-activati}
      --cron
      --cupsd
      --2*[dbus-daemon]
      --dbus-launch
      --dhcdd--dhclient
      --exim4
      --gconfd-2
      --gdm--gdm--Xorg
                        --x-session-manag--bluetooth-apple
                                                --gnome-panel
                                                --gnome-settings--{gnome-settings-}
                                                --kerneloops-appl
                                                --metacity
                                                --nautilus
                                                --nm-applet
                                                --seahorse-agent
                                                --system-config-p
                                                --update-notifier

```

Durata de funcționare a sistemului; uptime și downtime

Unul din avantajele majore ale sistemului de operare Linux este stabilitatea sa, motiv pentru care este și atât de răspândit în aplicații unde este nevoie de siguranță și durată mare de funcționare. O astfel de aplicație ar putea fi un calculator care se ocupă de servicii bancare și pe care banca respectivă nu-și poate permite să-l oprească deloc, sau unde siguranța datelor procesate este foarte importantă. În terminologia folosită pentru servere, apar astfel doi termeni legați de disponibilitatea acestora: **uptime** și **downtime**. Primul se referă la durata de funcționare a sistemului de când a fost pornit, pe când cel de-al doilea vine în completare și măsoară timpul în care serverul respectiv nu a fost operațional. Deși nu există o modalitate corectă de a măsura downtime-ul local (acesta fiind măsurat extern de către alte sisteme), se poate determina uptime-ul unui sistem, știindu-se exact momentul în care sistemul a pornit cât și ora curentă. Pentru afișarea uptime-ului se folosește comanda **uptime**:

```

computer1:/etc/rc3.d# uptime
16:57:28 up 2 days, 18:38, 3 users, load average: 0.00, 0.00, 0.00
computer1:/etc/rc3.d# _

```

Pe lângă comanda **uptime**, se mai poate folosi și comanda **top**, utilă nu numai pentru determinarea timpului de funcționare al calculatorului, cât și pentru afișarea proceselor care rulează în acel moment, totul în timp real. În plus, pe lângă aceste informații, comanda **top** va mai afișa și alți parametri importanți ai sistemului, precum numărul de procese și informații despre memorie.


```
top - 16:58:06 up 2 days, 18:39, 3 users, load average: 0.00, 0.00, 0.00
Tasks: 86 total, 1 running, 85 sleeping, 0 stopped, 0 zombie
Cpu(s): 0.0%us, 0.0%sy, 0.0%ni, 99.9%id, 0.0%wa, 0.0%hi, 0.0%si, 0.0%st
Mem: 1036092k total, 247112k used, 788980k free, 32004k buffers
Swap: 409616k total, 0k used, 409616k free, 110760k cached
```

PID	USER	PR	NI	UIRT	RES	SHR	S	%CPU	%MEM	TIME+	COMMAND
1	root	20	0	2100	688	588	S	0.0	0.1	0:02.32	init
2	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	kthreadd
3	root	RT	-5	0	0	0	S	0.0	0.0	0:00.00	migration/0
4	root	15	-5	0	0	0	S	0.0	0.0	0:00.74	ksoftirqd/0
5	root	RT	-5	0	0	0	S	0.0	0.0	0:00.00	watchdog/0
6	root	15	-5	0	0	0	S	0.0	0.0	0:04.98	events/0
7	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	khelper
39	root	15	-5	0	0	0	S	0.0	0.0	0:00.06	kblockd/0
41	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	kacpid
42	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	kacpi_notify
106	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	kseriod
143	root	20	0	0	0	0	S	0.0	0.0	0:00.00	pdflush
144	root	20	0	0	0	0	S	0.0	0.0	0:00.46	pdflush
145	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	kswapd0
146	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	aio/0
585	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	ata/0
586	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	ata_aux
594	root	15	-5	0	0	0	S	0.0	0.0	0:00.00	ksuspend_usbd

Terminarea forțată a proceselor

În ciuda stabilității foarte mari a sistemelor Linux, uneori se întâmplă că procesele ce rulează să se blocheze sau să funcționeze anormal, făcând astfel oprirea lor prin mijloacele convenționale imposibilă. Linux oferă o stabilitate foarte mare a sistemului, acesta putând să continue operarea normală chiar și în cazul în care unul din procesele în curs de derulare se blochează. În acest caz, pentru a aduce sistemul la o stare de normalitate, se folosesc mecanisme de terminare forțată a proceselor.

Unic pentru fiecare proces ce rulează la un moment dat, este atributul PID (sau identificatorul numeric unic) al procesului, care poate fi folosit pentru a comanda terminarea sa. Pentru aflarea PID-ului unui proces, se pot folosi mai multe tehnici care implică folosirea următoarelor comenzi:

- comanda **ps** afișează o listă a tuturor proceselor ce rulează în mod curent. Dat fiind faptul că foarte multe procese pot rula la un moment dat (având ca rezultat o listă de procese ce se întinde pe mai multe ecrane), se recomandă folosirea împreună cu comanda **less**. Parametrii cei mai utili sunt **-ax** (unde parametrul **a** forțează afișarea tuturor proceselor iar **x** elimină restricția de afișare numai a proceselor ce rulează dependent de o consolă)

EXEMPLU: **ps ax | less** va lista toate procesele, câte unul pe linie, afișând totodată și PID-ul fiecăruia.

- comanda **top** prezentată anterior poate fi și ea folosită pentru afișarea proceselor ce rulează în momentul curent, însă este mai puțin folositoare datorită cantității limitate de informații pe care le poate afișa la un moment dat.
- o altă metodă de folosire a comenzii **ps** este în concordanță cu comanda **grep** pentru filtrarea rezultatelor la un număr mai mic: **ps ax | grep „nume_proces”**. Acest mod de apelare va filtra rezultatele, afișând informații numai despre procesele al căror nume conține șirul de caractere menționat între ghilimele.

```
computer1:/etc/rc3.d# ps ax i grep "bash"
2707 pts/0      Ss+  0:00 bash
3376 tty2        S    0:01 -bash
3422 tty2        R    0:04 bash
3545 tty3        S    0:00 -bash
3549 tty3        S+   0:00 bash
3631 tty2        S+   0:00 grep bash
computer1:/etc/rc3.d# _
```

Această comandă poate fi extrem de utilă dat fiind faptul că mai multe procese cu nume identice pot rula la un moment dat.

- comenzile **pidof** și **pgrep** pot fi folosite pentru afișarea tuturor identificărilor PID ce rulează sub numele menționat în parametru. Urmărind exemplul anterior, comanda **pgrep bash** sau **pidof bash** ar fi afișat PID-urile vizibile în poză. De regulă, comanda **pgrep** este mai utilă, fiindcă permite o libertate mai mare în specificarea șirului de caractere ce reprezintă numele procesului căutat.

OBSERVAȚIE: folosirea comenzilor împreună cu caracterul „|” simbolizează folosirea nu a uneia, ci a două comenzi. Pe scurt, acest caracter specifică faptul că datele primei comenzi afișate pe ecran(output-ul), vor reprezenta datele de intrare (sau input) pentru comanda **grep**, care are rolul de a filtra text. Ca urmare, deși comanda **ps ax** ar returna informații despre toate procesele existente, afișarea propriu-zisă pe ecran se rezumă la doar câteva, datorită comenzii **grep**.

Odată aflat PID-ul unui proces, se poate trece la terminarea acestuia prin intermediul comenzii **kill**. Aceasta poate primi o serie de parametri, însă de regulă este folosit doar cu unul singur: PID-ul procesului țintă (dat fiind exemplul anterior, comanda **kill 3376** ar comanda terminarea procesului cu PID-ul respectiv).

Modul în care **kill** funcționează este prin folosirea sistemului de semafoare aferent sistemului de control al proceselor. Comandă va trimite o cere de terminare normală procesului țintă, care va returna (ca orice alt proces) un rezultat. Există cazuri în care procese blocate nu vor răspunde unor astfel de comenzi, când este necesară terminarea forțată prin includerea parametrului **-9** (**kill -9 3376**). Terminarea proceselor în acest mod nu este recomandată, dat fiind faptul că poate întrerupe procese importante în mod necontrolat.

Pe lângă comanda **kill** (care necesită neapărat PID-ul procesului țintă), se mai poate folosi și comanda **kill** care va lua ca parametru direct numele procesului. Spre deosebire de comandă **pgrep**, comanda **kill** nu va afișa PID-urile proceselor, ci le va comanda terminarea.

Procesul INIT

După cum se vede în structura de procese prezentată în poza anterioară, procesul rădăcină, care este părintele tuturor proceselor ce rulează la un moment dat, se numește **init**. Procesul **init** are rolul de a inițializa sistemul într-un mod prestabilit, după ce kernel-ul sistemului a fost încărcat în

totalitate. Mai exact, init va efectua sarcini precum montarea sistemului de fișiere, pornirea serviciilor aferente sistemului, sau chiar încărcarea interfeței grafice.

În Linux, init este executat cu un parametru, cunoscut și ca **run-level**, prin care se specifică modul în care se dorește inițializarea sistemului. Valoarea acestui parametru este un număr întreg între 1 și 6, și semnalează ce componente ale sistemului vor fi inițializate, fiecare nivel având o serie de script-uri ce vor fi executate. Aceste scripturi sunt de obicei stocate în directoare dedicate fiecărui nivel (de regulă în **/etc/rc[n].d** unde **[n]** reprezintă run level), iar fișierul principal de configurare în **/etc/inittab**.

De menționat este faptul că script-urile conținute în directoarele ce reprezintă fiecare run-level nu sunt de fapt fișiere propriu-zise, ci legături către alte scripturi (care se află de regulă în **/etc/init.d**). Motivul acestei implementări este că se poate modifica comportamentul fiecărui run-level independent de celelalte, fără să existe 6 copii diferite ale scripturilor folosite.

```
computer1:/etc/rc3.d# ls /etc/rc3.d/
README                S20exim4                S30gdm
S05loadcpufreq        S20kerneloops          S30system-tools-backends
S10rsyslog             S20nfs-common           S89anacron
S12acpid               S20openbsd-inetd       S89atd
S12dbus                S24dhcdd                S89cron
S14avahi-daemon        S24hal                  S99rc.local
S19cpufrequtils        S26network-manager     S99rmnologin
S20cups                S26network-manager-dispatcher S99stop-bootlogd
computer1:/etc/rc3.d# _
```

În poza de mai sus, se poate observa lista tuturor scripturilor atribuite nivelului 3. Se poate observa că acestea urmează un algoritm de denumire caracteristic, pentru a putea determina ordinea în care aceste sunt executate (alfabetic). În cazul în care se dorește executarea unui script pentru pornirea unui serviciu înaintea altuia, numele primului trebuie să fie înaintea celui de-al doilea din punct de vedere alfabetic.

După terminarea inițializării procesului init, acesta este clonat prin procedeul **fork**, devenind un nou proces, numit **getty**. Acesta are rolul de a porni consolele virtuale aferente sistemului, de regulă în număr de 6 (acest număr fiind configurabil în **/etc/inittab**).

Fișierul de configurare **inittab** poate de asemenea să specifice modul în care sistemul va reacționa la combinația de taste CTRL+ALT+DEL (reset în mod implicit) – se dorește închiderea corectă a sistemului spre deosebire de închiderea imediată.

Nivele run-level pot să însemne lucruri diferite în funcție de distribuția folosită, însă în Debian, există următoarele nivele:

- **nivelul 0** – halt (shutdown)
- **nivelul 1** – single user
- **nivelele 2-5** – multi-user cu console și mod grafic (distribuția Debian nu face nici o diferență între modurile 2,3,4 și 5)
- **nivelul 6** – reboot

```
# /etc/inittab: init(8) configuration.
# $Id: inittab,v 1.91 2002/01/25 13:35:21 miquels Exp $

# The default runlevel.
id:2:initdefault:

# Boot-time system configuration/initialization script.
# This is run first except when booting in emergency (-b)
mode.
si::sysinit:/etc/init.d/rcS

# What to do in single-user mode.
~~:S:wait:/sbin/sulogin

# /etc/init.d executes the S and K scripts upon change
# of runlevel.
..
```

Fire de execuție (THREADS)

Firele de execuție, cunoscute și sub numele de threads (din limba engleză), ca și procesele, reprezintă un mecanism de execuție în paralel a instrucțiunilor, facilitate esențială pentru orice sistem de operare modern, cu capabilități multi-tasking.

Deși termenul este folosit foarte des drept sinonim pentru proces, există diferențe foarte mari între acestea. Procesele în Linux se referă la instanțe ale unor programe în execuție, adică inclusiv variabile, informații de stare și în general întregul spațiu de memorie al respectivului program, pe când firele de execuție se referă la execuția în cadrul unui proces. Deși firele de execuție și procesele funcționează în mod diferit în funcție de sistemele de operare pe care rulează, există totuși un punct comun: firele de execuție pot să împartă accesul la informațiile de stare caracteristice unui proces fără să fie nevoie să apeleze la mecanisme de comunicare complicate precum comunicarea prin fișiere, pipes sau sockets. Firele de execuție pot efectua schimbul de informații în mod direct prin intermediul variabilelor sau al altor structuri de memorie. Consecința directă a acestui fapt este că schimbarea contextului de execuție se poate face mai rapid pentru fire de execuție decât pentru procese.

Principalul avantaj al firelor de execuție este acela al performanței în cazul aplicațiilor; o bază de date de exemplu, care este caracterizată de un volum foarte mare de operații de intrare/ieșire poate obține o îmbunătățire foarte mare a performanței dacă folosește o abordare orientată pe fire de execuție. Un alt exemplu de utilizare al firelor de execuție ar putea fi un program care necesită interacțiunea permanentă cu utilizatorul în timp ce se execută operații pe fundal; abordarea normală în acest caz este implementarea proceselor costisitoare (din punct de vedere al timpului de execuție) folosind conceptul de „worker thread” (sau fir de execuție „lucrător”), care va rula în mod paralel cu firul de execuție principal al programului, păstrând însă responsivitatea acestuia.

Principalul dezavantaj al firelor de execuție este dificultatea de a le implementa în mod corect, fapt datorat mecanismului de funcționare mai greu de înțeles, cât și diverselor conflicte ce pot să apară datorită accesului mai multor entități la resurse comune.

Concluzionând, procesele pot fi reprezentate drept un fir de execuție alături de spațiul sau de adresare și alte informații, și pot fi compuse din mai multe fire de execuție.

Conceptul de pipe

În cadrul sistemelor de operare de tip UNIX, conceptele de **pipe** și **pipeline** se referă la mecanisme de redirectare ale fluxurilor de date între procese, având ca principală întrebuintare comunicarea între procese diferite, sau implementarea unor filtre. Un exemplu bun de folosire al conceptului poate fi regăsit în cadrul unor comenzi prezentate anterior: **ps ax | less**. În acest caz, avem de-a face cu două comenzi separate, **ps** (cu parametrii menționați, listează procesele existente în momentul execuției) și **less** (care permite parcurgerea și navigarea în cadrul textelor lungi). Caracterul „|” poate fi observat între cele două comenzi, și reprezintă mecanismul obișnuit de creare al unui pipe. Comanda anterioară ar avea ca efect, afișarea numelor tuturor proceselor ce rulează la momentul execuției în sistem și ar depăși cu ușurință cantitatea maximă de informații ce poate fi afișată într-un singur ecran.

Comanda anterioară poate fi rezumată astfel: redirecționează fluxul standard output (ceea ce se afișează pe ecran) rezultat din execuția comenzii **ps** direct către aplicația **less**, care îl va putea prezenta utilizatorului într-o formă ușor de citit. Un alt exemplu de folosire al pipe-urilor este filtrarea datelor, această abordare fiind uneori mai eficientă decât căutarea manuală. Filtrarea se poate face redirectând fluxul standard output spre un program care în loc să ofere facilități de navigare al conținutului primit, va afișa numai informațiile dorite, pe baza unor parametrii. Un astfel de program este **grep**, care folosit simplu cu un singur parametru va afișa doar liniile ce conțin textul specificat: comanda **ps ax | grep bash** va redirecționa rezultatul (lista de procese) către comandă **grep**, care în baza parametrului specificat, va afișa doar acele linii care conțin șirul de caractere „bash”.

Redirecționarea fluxurilor de intrare/ieșire

Urmând filozofia Linux conform căreia totul este un fișier, chiar și tastatura, mouse-ul sau alte dispozitive sunt reprezentate prin fișiere. Principalul avantaj al acestei abordări este ușurința cu care se pot interfața dispozitivele de către programatori, cât și posibilitatea de a trimite, a primi sau a manipula datele primite de la dispozitive.

De exemplu, fișierul **/dev/input/mice** este folosit pentru a primi informații de la mouse, și folosind comanda **cat /dev/input/mice** (urmat de acționarea mouse-ului) într-un terminal vom putea observa schimbările. Comanda **cat** are rolul de a concatena conținutul unui fișier pe standard output (adică pe ecran).

Linux folosește 3 fluxuri importante de intrare/ieșire, codificate numeric: 0 – standard input (tastatură), 1 – standard output (ecran), 2 – standard error.

Shell-ul ne permite să redirectionăm aceste fluxuri pentru o flexibilitate foarte mare, folosind operatorii „>” (redirectarea fluxului de ieșire) și „<” (redirectarea fluxului de intrare). Comanda **ps ax > fișier.txt** ar avea ca rezultat, un efect similar cu cel observat la pipe-uri: **ps ax** ar afișa toate procesele ce rulează, însă rezultatul în loc să fie afișat pe standard output, va fi redirectat către fișierul **fișier.txt** pentru examinare ulterioară. În cazul în care se apelează o comandă care poate genera o eroare (de exemplu, comanda **ls** pe un fișier inexistent), ar trebui redirectat fluxul de standard error: **ls fișier_inexistent.txt 2> fișier.txt**.

În cadrul sistemului Linux există o serie de fișiere speciale, ce pot fi folosite în diverse moduri. Unul dintre aceste fișiere este **/dev/null**, un fișier cu dimensiunea 0, în care orice modificare este imediat ștearsă. O utilizare a acestui fișier ar putea fi ascunderea mesajelor de eroare: **ls fișier_inexistent 2> /dev/null** ar trimite mesajul de eroare către **/dev/null**, fiind total ignorat.

O altă abordare a redirectării poate fi folosită pentru tipărirea unui fișier prin intermediul imprimantei: comanda **cat fișier.txt > /dev/lp0** ar trimite conținutul fișierului **fișier.txt** către fișierul **lp0** care reprezintă imprimanta sistemului.

În mod invers, fluxul de intrare poate fi redirectat, astfel încât să reprezinte un fișier. Operatorul de redirectare al fluxului de ieșire către un fișier este în mod absolut distructiv, în sensul că va rescrie tot fișierul țintă cu noul conținut. Pentru a adăuga date la un fișier deja existent (modul append) trebuie folosit operatorul „>>”.

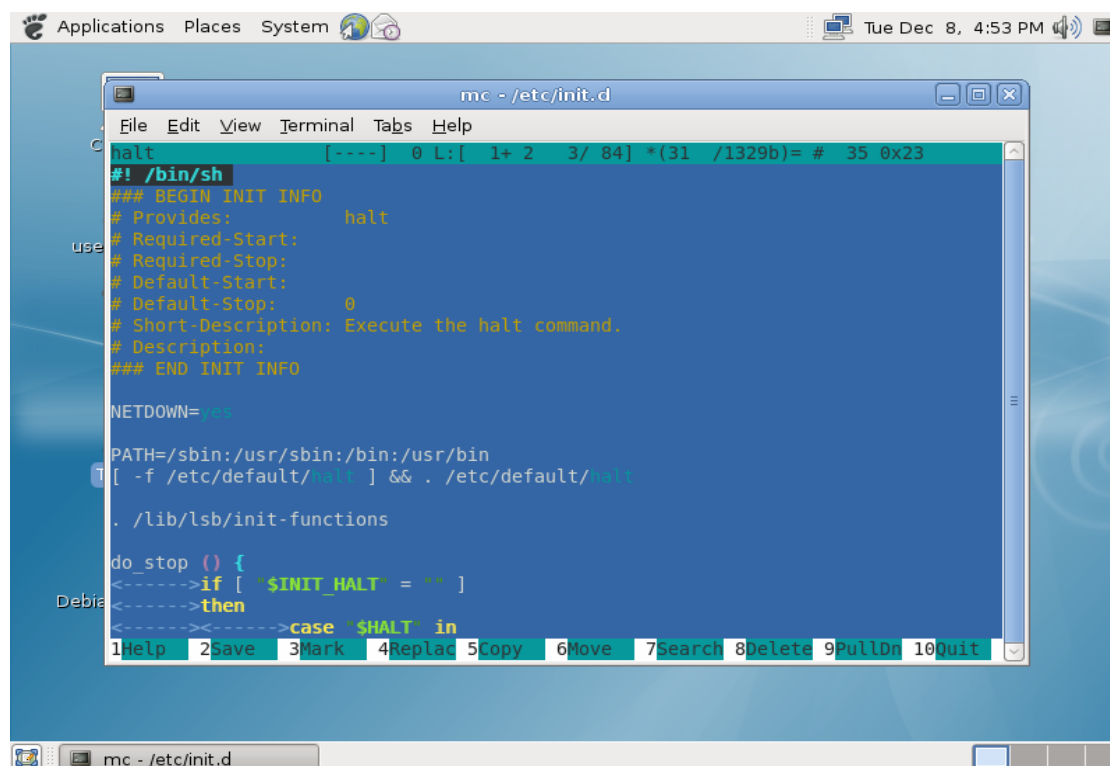
Script-uri bash

Bash, sau Bourne-Again Shell, este un interpretor de comenzi pentru sistemul de operare Linux. Acesta este shell-ul implicit pentru o serie foarte mare de distribuții Linux, printre care și Debian, și pe lângă facilitățile de interpretare a comenzilor lansate prin intermediul consolei, oferă și posibilitatea scrierii unor script-uri pentru diverse sarcini, în general în scopul automatizării unor procese, sau pentru îmbunătățirea facilităților existente ale unui program. Un script este un fișier text, ce poate fi executat, al cărui conținut respectă un limbaj de scripting (în cazul de față **bash**). Fiind vorba de fișiere text, și nu binare (compilate în cod mașină), scripturile nu pot fi executate propriu-zis; acestea sunt de fapt interpretate de un program specializat, de cele mai multe ori fiind vorba de interpretorul în a cărui limbaj a fost conceput respectiv script. În continuare vom studia procesul de creare al script-urilor **bash**, cât și o parte din facilitățile pe care acest interpretor le oferă.

Conceptul din spatele scripturilor este simplu: comenzile din script vor fi interpretate linie cu linie, secvențial (sau cu salturi în cazul script-urilor mai complicate). Orice script bash, este de fapt un fișier, cu drepturi de execuție, și pot fi lansate în mai multe feluri. Să presupunem în acest scop, că avem un script numit **s1** în directorul curent, pe care îl putem rula folosind unul din mijloacele de mai jos:

- Apelând una din comenzile **bash s1** sau **sh s1**. Această modalitate dezvăluie faptul că se folosește un interpretor pentru execuție. În acest caz, fișierul nu trebuie să aibă neapărat drepturi de execuție.
- Lansând direct scriptul cu ajutorul comenzii **./s1**. În acest caz, shell-ul își va da seama că este vorba de un script bash, și va instantia interpretorul pentru acest fișier. Pentru a seta drepturi de execuție asupra unui script pentru orice utilizator este necesară rularea comenzii **chmod a+x nume_script**.

În continuare vom analiza compoziția unui fișier de tip script bash.



Se recomandă folosirea editorului **mcedit** pentru vizualizarea, modificarea sau crearea scripturilor bash (sau a multor altor tipuri de script-uri) datorită facilităților de subliniere colorată a sintaxei (syntax highlighting) oferite de acesta.

Orice script bash, va avea pe prima linie secvența de caractere „#!” urmată de locația interpretorului pentru care a fost scris. În scopuri didactice, vom crea în continuare un script simplu pornit de la reguli menționate anterior:

Pasul 1. Crearea fișierului se poate face cu ajutorul comenzii **touch**:

touch exemplu

Pasul 2. Setarea permisiunilor. Reamintim pe această cale, că se recomandă ca orice script să aibă drepturi de execuție, cel puțin pentru utilizatorii care ar trebui să-l poată rula.

chmod a+x exemplu

Pasul 3. Editarea fișierului în vederea introducerii comenzilor ce vor alcătui script-ul. Se recomandă folosirea editorului **mcedit** pentru acest pas. Mai jos, se poate observa un script foarte simplu prin care se șterge conținutul directorului **tmp** (presupunând că acesta există) din cadrul directorului home al utilizatorului curent, urmat de un mesaj de atenționare:

#!/bin/bash

rm ~/tmp/* 2>/dev/null

echo „Conținutul directorului temporar a fost golit.”

Se poate observa că întregul script este compus din comenzi obișnuite. Prima linie a script-ului semnalează faptul că avem de-a face cu un script și include și calea către interpretorul bash. A doua linie este ceva mai greu de înțeles la început, dar poate fi interpretată în felul următor: **rm** este comanda folosită pentru ștergerea fișierelor, caracterul **~** reprezintă simbolul pentru directorul home al ușerului curent, iar calea **~/tmp/*** se referă la toate

fișierele din cadrul directorului **tmp**, conținut în directorul home al utilizatorului curent. Ultima parte a liniei este o redirectionare; în secțiunile anterioare am discutat despre redirectionarea fluxurilor principale: standard out (1), standard error (2), respectiv standard în (3). Utilitarele pot trimite un mesaj de eroare către fluxul standard pentru a fi afișate pe ecran. În cazul de față, eroarea cea mai probabilă ar fi fost lipsa oricărui fișier din directorul **tmp**, utilizatorul fiind avertizat de acest lucru. Redirectionarea **2>/dev/null** nu face altceva decât să redirectioneze toate mesajele de eroare generate de către **rm** către dispozitivul special **/dev/null** (care are rolul de a ignora orice fel de date scrise în el) astfel încât acestea să nu mai fie afișate pe ecran. Ultima linie face uz de comandă **echo** folosită pentru afișarea unui text pe ecran, în cazul de față fiind vorba despre o atenționare simplă.

În cazul în care exemplul de mai sus este greu de înțeles la început (mai ales pentru începători) o soluție bună este simplificarea sa:

```
#!/bin/bash
rm ~/tmp/*
```

În varianta simplificată, am scăpat de orice fel de tratare a mesajelor de eroare, cât și de atenționarea asupra terminării operației de golire a directorului **tmp**. Așa cum am menționat mai devreme, comenzile vor fi executate secvențial, iar **prompt**-ul (controlul tastaturii) va fi redat consolei după terminarea completă a execuției.

În continuare vom trata câteva din facilitățile mai avansate oferite de interpretorul bash.

Variabilele sunt nume simbolice al căror rol este să păstreze o valoare ce poate fi schimbată. Exemplul următor arată cum se pot folosi variabilele pentru stocarea și afișarea unei valori:

```
#!/bin/bash
var="un text"
echo $var
```

Ca urmare a execuției script-ului anterior, se va afișa pe ecran prin comanda **echo** conținutul variabilei **var**. Ca observație, semnul \$ trebuie folosit în afișarea acesteia (acesta este modul prin care interpretorul bash își dă seama că are de-a face cu o variabilă și nu un text oarecare).

Parametrii în linie de comandă reprezintă o facilitate importantă oferită de interpretor. Ca și multe alte executabile ce se folosesc de parametri primiți prin cadrul liniei de comandă, la fel și scripturile pot beneficia de parametri pasați în linia de comandă. Deși linia de comandă este una singură, parametrii folosiți sunt de regulă separați prin spații. Accesul la acești parametri din cadrul unui script bash se face prin intermediul unor variabile speciale \$1, \$2, până la \$9. \$0 va reprezenta numele script-ului mereu.

```
#!/bin/bash
echo „Parametrul 1 este $1”
echo „Parametrul 2 este $2”
```

Lansând scriptul prin **./nume_script exemplu script**, interpretorul va afișa pe ecran:

Parametrul 1 este exemplu

Parametrul 2 este script

Citirea datelor de la tastatură se poate face în momentul execuției unui script folosind comanda **read**.

```
#!/bin/bash
echo „Introduceți numele dvs:”
read nume
echo „Bine ai venit $nume!”
```

Verificarea condițiilor prin cuvintele cheie if / else / fi oferă utilizatorilor posibilitatea de a verifica anumite condiții în cadrul script-urilor și de a lua decizii particularizate pentru fiecare caz în parte. În exemplul următor vom putea examina verificarea de condiții multiple, cât și verificarea existenței unui fișier:

```
#!/bin/bash
fișier="exemplu"
echo „Doriți afișarea rezultatului pe ecran? DA/NU”
read ecran
if [ -f $fișier ]; then
if [ $ecran == „DA” ]; then
    echo „Fișierul $fișier exista.”
fi
else
if [ $ecran == „DA” ]; then
    echo „Fișierul $fișier nu exista”
fi
fi
```

Exemplul anterior demonstrează cum se pot pune condiții multiple și cum se poate verifica existența unui fișier (a se observa spațiile dintre parantezele pătrate atât după cât și înaintea terminării condiției) - **fără aceste spații script-ul nu va funcționa**.

Verificarea se face cu ajutorul cuvântului cheie **if** care poate fi folosit cu o serie de parametri. Șirul **-f \$fișier** este folosit pentru verificarea existenței fișierului al cărui nume este conținut de variabila \$fișier. În cazul în care se dorește verificarea existenței unui director, trebuie folosit parametrul **-d**.

O structură **if** va urma un set de reguli: începe cu **if**, continua cu instrucțiunile care se vor executa dacă condiția este împlinită, opțional poate continua cu **else** (care va conține o secțiune de instrucțiuni executate în cazul în care condiția nu este împlinită) și obligatoriu se va termina în **fi** care semnalizează sfârșitul structurii.

După cum se poate observa în exemplul anterior, în cadrul structurilor **if** pot exista alte structuri **if**, numite „nested”.

Crearea copiilor de rezervă. Backup.

Copiile de rezervă sunt folosite în general pentru salvarea datelor importante, pentru a putea fi recuperate în urma unor evenimente neașteptate – precum defectarea unui disc, erori software sau chiar furt.

În general se recomandă includerea tuturor datelor în copiile de rezervă, în cazul în care spațiul dedicat acestei soluții permite acest lucru, și în cazul în care timpul folosit este disponibil. Copiile întregi reprezintă cea mai sigură soluție pentru siguranța datelor în eventualitatea unor dezastre, însă acestea pot dura foarte mult, în funcție de dimensiunea lor.

Pentru planificarea unui backup parțial, utilizatorul trebuie să aleagă în primul rând ce fișiere dorește să salveze. Printre fișierele importante (și posibil de neînlocuit) într-un sistem Linux, putem aminti:

- Baze de date – în cazul serverelor MYSQL, acestea pot fi găsite în **/var/lib/mysql**
- Fișiere personale și documente – localizate de obicei în **/home**
- Fișiere de configurare – directorul **/etc**

Copiile de siguranță pot fi realizate în mai multe feluri, fie salvând manual toate fișierele dorite, fie prin intermediul unor utilitare specializate. Indiferent de metoda de backup folosită, utilizatorul va trebui să planifice în prealabil ce resurse dorește să salveze.

Salvarea manuală a datelor poate fi realizată în mai multe feluri și variază de la copierea manuală a fișierelor, până la folosirea unor utilitare precum **dd** sau **tar**. În cazul copiilor de siguranță realizate manual, se recomandă însă folosirea unor utilitare precum cele menționate mai devreme în defavoarea copierii manuale, motivul principal fiind alterarea informațiilor legate de sistemul de fișiere – permisiuni, proprietar, etc.

Utilitarul tar, prezentat anterior poate fi folosit (împreună cu facilitatea de compresie **gzip**) pentru realizarea unei copii fidele a unei structuri de fișiere și directoare. Un exemplu bun ar fi: **tar -czf nume_copie.tar.gz /etc /home**

Comanda anterioară va crea o arhivă comprimată, reprezentând copia de siguranță a directoarelor **/etc /home**. Recuperarea datelor în eventualitatea pierderii lor, se poate realiza prin comanda: **tar -zxfp nume_copie.tar.gz**

Diferența față de dezarhivarea și decomprimarea normală a unei arhive este parametrul **-p**, prin care utilitarul va fi forțat să dezarhiveze structura de fișiere menținând permisiunile originale.

Utilitarul dd este unul dintre programele originale UNIX folosite pentru scrierea datelor în mod „raw”. Acesta este și în ziua de azi un program foarte util, datorită facilităților pe care le oferă: poate fi folosit pentru scrierea preferențială pe discuri, pentru extragerea parțială a unor secțiuni din fișiere binare și chiar pentru copierea la nivel de dispozitiv.

Modul de folosire al utilitarului **dd** este unul ceva mai complicat, prin faptul că se bazează pe o serie de parametri specificați în formatul **parametru=valoare**. Restrâns la cazul de față, **dd** poate fi folosit pentru crearea copiilor de siguranță a unor partiții întregi (sau secțiuni din acestea).

Copia de siguranță în acest caz, va fi un fișier imagine, care reprezintă o copie fidelă a partiției după care a fost făcută. Acest fișier imagine poate fi comprimat, urmând să fie păstrat, sau poate fi scris pe o partiție. Această metodă de backup este ceva mai cuprinzătoare, putând fi folosită pentru reîncărcarea întregului sistem, însă, datorită faptului că se folosește pentru crearea unei copii fidele a unei partiții, procesul poate dura considerabil mai mult.

Principiul de funcționare al utilitarului este simplu: acesta se bazează pe o sursă de intrare a datelor, folosită pentru citirea datelor ce vor fi copiate (fiind folosit foarte des pentru partiții întregi, sursa de intrare va fi de multe ori un fișier de tip dispozitiv) și o sursă pentru datele de ieșire, sau mai exact destinația unde se vor scrie datele copiate. Sursa de ieșire poate fi un fișier, în cazul în care se dorește crearea unei imagini, sau poate fi fișier de tip dispozitiv, în cazul în care se dorește scrierea directă pe o partiție.

Parametrii cei mai folosiți de către **dd** sunt **if (input file)** și **of (output file)**. Un exemplu bun de creare a unei imagini fidele a partiției **/dev/sda1** ar fi:
dd if=/dev/sda1 of=/(nume_imagine) bs=4096 conv=notrunc,noerror

Aceasta poate fi tradusă în felul următor: crează o imagine cu numele specificat, având ca sursă partiția **/dev/sda1**, cu dimensiunea blocului de date de 4096 bytes, și cu parametrii **notrunc** (adică fișierul destinație nu va fi trunct) și **noerror** (adică procesul va continua și în cazul în care apare o eroare). Pentru a facilita stocarea imaginii astfel create, aceasta poate fi comprimată cu ajutorul utilitarului **gzip**. Comanda anterioară poate fi adaptată astfel încât fișierul rezultat va fi direct o arhivă: **dd if=/dev/sda1 bs=4096 conv=notrunc,noerror | gzip -c > /(nume_imagine).gz**

Comanda poate fi interpretată în felul următor: redirecționează datele de ieșire ale utilitarului **dd** spre utilitarul **gzip** (care va redirecționa la rândul său fișierul comprimat către standard output prin parametrul **-c**) care va redirecționa din nou ieșirea către fișierul final, reprezentând imaginea comprimată a partiției.

Restaurarea unei partiții având fișierul imagine inițial (dacă acesta nu este comprimat) se poate face cu ajutorul comenzii: **dd if=/(nume_imagine) of=/dev/sda1 bs=4096 conv=notrunc,noerror**

În cazul în care fișierul imagine a fost comprimat după creare, comanda devine: **gzip -dc /(nume_imagine).gz | dd of=/dev/sda1**

OBSERVAȚIE: Principalul dezavantaj al utilitarului **dd** față de utilitare special gândite pentru back-up este acela că nu cunoaște detalii despre sistemul de fișiere folosit, și ca urmare, va face o copie fidelă a întregii partiții, inclusiv a blocurilor de date nefolosite. De asemenea, indiferent de utilitarul folosit pentru crearea copiei de siguranță, se recomandă alegerea unor nume sugestive pentru acestea, astfel încât administratorul să-și poată da seama cu ușurință care este conținutul unei copii și cât de actuală este aceasta.

Cron. Automatizarea procesului de creare a copiilor de siguranță.

Procesul de crearea a copiilor de siguranță poate fi automatizat într-o foarte mare măsură, prin programarea evenimentelor folosind daemon-ul **cron**. Rolul acestuia este de a programa și executa comenzi la intervale și date prestabilite.

Cron funcționează prin interpretarea fișierului **/etc/crontab**, un fișier text ce descrie când și ce comenzi sunt programate pentru executare automată. Fiecare utilizator își poate defini singur sarcinile automatizate, folosind utilitarul **crontab**, dat fiind faptul că aceștia nu au acces direct la fișierul **/etc/crontab**. Deși în mod implicit toți utilizatorii își pot defini propriul set de comenzi programate, administratorul poate să restricționeze accesul la acest serviciu prin intermediul fișierelor **cron.allow** și **cron.deny** localizate în directorul **/etc**. Utilizatorii listați în primul fișier vor avea dreptul de a programa sarcini automate, pe când cei listați în cel de-al doilea fișier nu vor avea acest drept.

Programarea sarcinilor se face de regulă prin intermediul utilitarului **crontab**. Pentru afișarea listei de comenzi planificate trebuie executată comanda **crontab -l**. Formatul tipic pentru fișierele **crontab** este simplu: fiecare linie reprezintă programarea pentru o comandă, și are 6 câmpuri separate prin spații. Acestea sunt în ordine:

- Numărul de minute după oră [0;59]
- Ora [0;23]
- Ziua lunii [1;31]
- Luna [1;12]
- Ziua săptămânii [0 și 7 pentru Duminică, progresiv pentru restul]
- Comandă ce va fi executată

Fiecare coloană poate conține mai multe valori separate prin virgulă sau o linie orizontală și se acceptă caracterul „*” cu semnificația **oricare**. Fiecare utilizator își poate edita propriul fișier **crontab**, prin executarea comenzii **crontab -e**. Utilizatorul root este singurul care are dreptul să editeze fișierele **crontab** ale altor utilizatori prin adăugarea parametrului **-u nume_utilizator**, de exemplu **crontab -u user1 -e**. Fiecare utilizator își poate programa astfel execuția unor comenzi la intervale prestabilite. Urmează câteva exemple de linii **crontab** explicate:

- **40 * * * * echo „salut” | wall** - va emite un mesaj de atenționare către toți utilizatorii la fiecare 40 de minute (0:40, 1:40, 2:40, etc...)
- **1-59 * * * * echo „salut” | wall** - va emite un mesaj de atenționare către toți utilizatorii o dată la fiecare minut.
- **0 12 * * 5 /bin/salut** - va executa scriptul / programul **/bin/salut** în fiecare zi de vineri (5 = vineri) la ora 12:00 AM.

Administratorul poate programa și activați la nivel de sistem, prin editarea fișierului **/etc/crontab**, păstrând sintaxa anterioară, cu mențiunea că acesta are o coloană în plus inserată între coloanele „ziua săptămânii” și „comanda”. Acest câmp adițional reprezintă utilizatorul pentru care se va executa comanda programată.

În această idee, procesul de creare a copiilor de siguranță poate fi automatizat, la ore prestabilite. În cazul în care administratorul dorește executarea unei copii de siguranță (de orice tip) folosind oricare dintre utilitarele prezentate mai sus, tot ce are de făcut este de adăugat linia corespunzătoare în fișierul **/etc/crontab**.

Gestionarea datelor

Datele de pe un sistem Linux se supun la aceleași reguli și riscuri ca orice date: proasta gestionare, pierdere sau furt. În continuare vom studia câteva tehnici diferite pentru o mai bună gestiune a datelor pe un sistem Debian.

Primul pas în vederea asigurării confidențialității datelor, pentru a preveni accesul neautorizat sau furtul este **criptarea**. Procesul implică configurarea unei partiții astfel încât conținutul acesteia să fie criptat în mod transparent și în timp real. Configurarea unei partiții pentru criptare presupune pierderea datelor existente de pe acea partiție, astfel încât se recomandă un back-up înainte de acest proces.

În continuare, vom prezenta pașii necesari pentru obținerea unei partiții criptate:

Pasul 1. Instalarea pachetului cryptsetup

Pachetul cryptsetup conține utilitarul necesar pentru realizarea partiției criptate. Instalarea se face apelând comanda **apt-get install cryptsetup**.

Pasul 2. Pregătirea partiției pentru criptare

Acest pas presupune scrierea unor seturi de date pe partiția ce urmează să fie creată, pentru a ajuta procesul de criptare. Urmând acest pas înaintea criptării efective, partiția va fi protejată împotriva atacurilor. Comanda recomandată pentru acest pas este una care îndeplinește raportul optim securitate / timp generare: **badblocks -vfw (nume_partiție) 4096**

Aplicația **badblocks** este folosită pentru examinarea discurilor fizice în scopul detectării blocurilor defecte. În acest scop, acest program va genera niște șabloane pe discul ce urmează să fie testat. Parametrii **-vfw** forțează scrierea acestor șabloane pe disc, generând în esență un punct de plecare pentru facilitarea procesului de criptare.

Pasul 3. Criptarea propriu-zisă

În acest moment, partiția prelucrată este gata pentru a fi criptată. Pentru a continua, vom stabili câțiva parametri necesari pentru criptare: tipul de cifru va fi „**aes**” și dimensiunea cheii în biți (în cazul de față vom folosi o cheie pe 128 de biți). Pentru folosirea acestor parametri, și pentru realizarea criptării este necesară executarea comenzii (atenție, partiția în cauză NU trebuie să fie montată): **cryptsetup luksFormat (nume_partiție) -c aes -s 128**

În cazul în care utilizatorul primește un mesaj de eroare în genul „**Failed to setup dm-crypt key mapping**”, înseamnă că modulul de kernel corespunzător nu este încărcat. Pentru remedierea acestei probleme, trebuie încărcat modulul **dm-crypt** prin executarea comenzii: **modprobe dm-crypt**.

Dacă totul a decurs normal, utilizatorul va fi rugat să introducă o frază de

autentificare ce va fi folosită în procesul de criptare. Se recomandă alegerea unei fraze cât mai lungi pentru garantarea unui nivel minim de securitate.

Pasul 4. Deschiderea partiției

În acest moment, avem pregătită o partiție criptată, însă aceasta nu este încă gata pentru folosire. Pentru deschiderea partiției, putem folosi tot utilitarul **cryptsetup** cu parametrii corespunzători:

cryptsetup luksOpen (nume_partiție) (nume_criptat)

În urma execuției comenzii anterioare, utilizatorul va fi rugat să introducă frază de autentificare stabilită la pasul anterior. Dacă fraza este corectă, atunci partiția va fi prelucrată de către sistemul **device mapper**. Acesta este un cadru prin care se poate face o asociere între două dispozitive – în esență crează un fel de legătură către un fișier de tip dispozitiv ce va fi folosit pentru montarea partiției criptate. Numele fișierului creat este cel specificat ca parametrul în cadrul comenzii anterioare, și va putea fi localizat la adresa: **/dev/mapper/(nume_criptat)**.

Pasul 5. Crearea sistemului de fișiere pe partiția criptată

În acest moment partiția criptată este gata, însă nu poate fi folosită până nu va fi formatată. Pentru a crea sistemul de fișiere standard Linux (ext3) pe noua partiție, trebuie folosită comanda:

mke2fs -j /dev/mapper/(nume_criptat)

Pasul 6. Montarea partiției

Odată generată, partiția poate fi montată ca orice altă partiție, însă cu avantajul că toate datele vor fi criptate într-un mod transparent pentru utilizator. Montarea se face în mod normal, cu mențiunea că nu se va folosi numele real al partiției pentru montare, ci legătura pusă la dispoziție prin **device mapper**.

Pentru automatizarea procesului, astfel încât partiția criptată să fie montată în mod automat la fiecare pornire a sistemului trebuie modificate două fișiere. Primul este **/etc/crypttab**, folosit pentru generarea device map-ului, prin adăugarea liniei următoare la sfârșitul fișierului.

(nume_criptat) (nume_partiție) none luks

Al doilea fișier ce trebuie modificat este **/etc/fstab** pentru ca procesul de montare automată să poată fi finalizat. Montarea fără fișierul generat de device mapper nu este posibilă, dar dincolo de acest punct, procesul de montare este unul obișnuit prin includerea liniei următoare în **/etc/fstab**

/dev/mapper/(nume_criptat) (punct_de_montare) ext3 defaults 0 0

În cazul în care utilizatorul optează pentru montarea automată a partiției în timpul procesului de boot, acesta va fi rugat să introducă frază de autentificare. În cazul în care totul a decurs corect, sistemul beneficiază în acest moment de o partiție criptată conform parametrilor aleși.

OBSERVAȚII: Toate comenzile lansate mai devreme necesită drepturi de root. Deși criptarea partiției principale (directorul root) este posibilă, acest procedeu este mai complicat, o alternativă bună la această problemă fiind stocarea datelor importante pe partiția criptată. În cazul examinării hard disc-

ului cu partiția criptată cu utilitare externe, aceasta va figura ca fiind ne-formatata.

Compresia și arhivarea datelor

Conceptul de arhivare și compresare a datelor nu este unul nou. Datele comprimate și/sau arhivate sunt mult mai ușor de gestionat și de transmis pe diverse medii de comunicare, distribuțiile Linux folosindu-se în mod curent de ambele procedee, datorită multiplelor avantaje oferite.

Putem începe prin a sublinia diferențele dintre cele două procese: cel de arhivare și cel de compresie a datelor. Arhivarea datelor înseamnă unirea mai multor fișiere și directoare în cadrul unui singur fișier, pentru o gestiune mai ușoară, păstrând totuși o serie de informații despre structura de fișiere anterioară, astfel încât aceasta să poată fi dez-arhivată. Probabil cel mai cunoscut format de arhive în lumea Linux este formatul **tar** (tape archive), creat la începuturile sistemelor UNIX, și standardizat de abia în 1988. Acesta a fost gândit inițial pentru scrierea secvențială pe dispozitive de stocare cu bandă magnetică, în scopul creării copiilor de rezervă, dar este folosit în prezent pentru gruparea unei structuri de fișiere și directoare într-un singur fișier principal, păstrând informații despre sistemul de fișiere inițial (informații precum permisiuni, proprietari, grupuri, date, etc).

Compresia în schimb, reprezintă un proces de codificare reversibilă a fișierelor, astfel încât acestea să ocupe mai puțin spațiu fizic decât ar ocupa în mod normal. Deși există multe tipuri de algoritmi de comprimare a datelor, unele specializate pe anumite tipuri de date (fișiere video, foto, sunet, etc), toate au în comun același principiu de funcționare: codificarea datelor existente într-o manieră mai eficientă din punct de vedere al spațiului folosit pentru stocare. În general utilizatorii vor observa că programele de compresie vor funcționa cel mai bine când sunt aplicate pe surse de date al căror conținut nu este foarte aleatoriu (text, imagini, etc). În cadrul sistemelor Linux, cele mai răspândite formate compresate sunt **zip**, **bz2** și **gzip**.

O tactică comună este combinarea celor două procedee pentru a beneficia de facilitățile oferite de acestea; de aici vine și cunoscutul format **tar.gz**, care nu este altceva decât o arhivă comprimată.

Crearea arhivelor se poate realiza cu ajutorul programului **tar**, inclus în toate distribuțiile Linux. Sintaxa este una destul de simplă: **tar [parametri] [fișiere]**. Vom examina în continuare câteva exemple în care se folosește comanda **tar**:

Pentru arhivarea unor fișiere: **tar -cf arhiva.tar fișier_1 fișier_2 ...**

În exemplul anterior, programul va arhiva toate fișierele menționate într-o singură arhivă numită **arhiva.tar**. Parametrul **-c** semnalează faptul că dorim să creăm o nouă arhivă, iar **-f** înseamnă că se va lucra cu fișiere.

Pentru arhivarea unui director: **tar -cf arhiva.tar cale_director**

Comprimarea fișierelor se realizează de obicei cu ajutorul utilitarului **gzip**. Sintaxa este foarte simplă: **gzip (nume_fișier)**. Este important de menționat faptul că programul **gzip** va înlocui fișierul original cu varianta sa

comprimată. Dacă numele fișierului inițial ar fi fost (de exemplu) **test.txt** în urma acționării comenzii **gzip test.txt**, fișierul test.txt ar fi fost înlocuit cu fișierul **test.txt.gz**.

Decompresia datelor se face prin includerea parametrului **-d** în linia de comandă: **gzip -d nume_fișier.gz**

Comprimarea arhivelor se poate face printr-o combinație a utilităților prezentate mai devreme. Utilitarul **tar** știe să arhiveze fișierele arhivă cu ajutorul utilitarului **gzip**: **tar -czf nume_arhiva.tar.gz director**

Parametrii **-c** și **-f** au aceeași semnificație ca și în exemplele anterioare, iar parametrul **-z** specifică faptul că se dorește comprimarea arhivei folosind **gzip**. Despachetarea unui fișier de tip **tar.gz** (acest lucru implică mai întâi decompresia și apoi dezarhivarea sa) se poate face prin comanda:

tar -xzf nume_arhiva.tar.gz

Parametrul **-x** specifică faptul că se dorește dezarhivarea fișierului tar, rezultat în urma decomprimării.

Sistemul Debian nu oferă în mod nativ compresie la nivel de partiție, însă dacă luăm în considerare progresele tehnologice din ultimii ani, care au redus costurile dispozitivelor de stocare semnificativ, în timp ce capacitățile de stocare au crescut foarte mult, atunci putem concluziona că pierderea de performanță asociată cu compresia datelor în timp real, raportată la nivelul de spațiu recuperat nu este justificabilă.

Implementarea cotelor de disc

Implementarea cotelor de disc este o necesitate în orice sistem multi-user. Riscurile asociate cu supraîncărcarea discurilor sistemului nu pot fi ignorate, putând să genereze o serie întreagă de erori, mergând de la imposibilitatea despachetării fișierelor temporare (necesare pentru foarte multe activități) și ajung până la cazuri de blocare a sistemului. Din acest motiv, activitatea utilizatorilor poate fi restricționată într-un sistem Linux, astfel încât aceștia – conform politicii de sistem curente – să nu poată depăși un anumit nivel de spațiu ocupat.

Sistemul de cote permite configurarea cotelor de disc pentru două categorii: stabilirea cotelor individuale pentru fiecare utilizator cât și stabilirea cotelor de grup – ambele vor trebui luate în considerare în momentul stabilirii cotelor, pentru o gestiune corectă a sistemului de fișiere.

În cazul cotelor individuale, administratorul va stabili o limită superioară de spațiu de stocare alocat pentru fiecare utilizator, pe când în cazul cotelor de grup, se va stabili limita superioară de spațiu alocat întregului grup.

Pentru configurarea sistemului cu suport pentru stabilirea cotelor de disc, trebuie urmați următorii pași:

Pasul 1. Montarea partițiilor în mod corespunzător

Pentru ca sistemul de cote să funcționeze corespunzător, partițiile pe care se dorește stabilirea cotelor de disc trebuie montate cu o serie de parametri speciali. Aceștia se vor seta în cadrul fișierului **/etc/fstab**, în cea de-a cincea

coloană „<options>”, prin includerea cuvintelor **usrquota**, respectiv **grpquota**. În exemplul de mai jos, vom monta partiția principală a sistemului cu suport pentru cote prin modificarea liniei corespunzătoare în fișierul **/etc/fstab**: `/dev/sda1/ ext3 errors=remount-ro,usrquota,grpquota 0 1`
Semnificația liniei de mai sus a fost explicată în cadrul procesului de montare automată, singura modificare a procesului de montare pentru partiția **/dev/sda** fiind includerea parametrilor **usrquota** și **grpquota**.

Pasul 2. Încărcarea modului kernel corespunzător

Pentru ca întregul sistem să funcționeze, modulul de kernel **quota_v2** trebuie încărcat. Acest proces se poate realiza temporar, pentru instanța curentă a sistemului prin intermediul comenzii **modprobe quota_v2**, iar în cazul în care se dorește încărcarea modului în mod automat la fiecare reinițializare a sistemului, linia „**quota_v2**” trebuie adăugată la sfârșitul fișierului **/etc/modules** prin intermediul comenzii `echo `quota_v2` >> /etc/modules`.

Pasul 3. Crearea fișierelor de configurare

Pentru ca întreg procesul de instalare să decurgă în mod normal, administratorul trebuie să creeze două fișiere ce vor fi folosite pentru descrierea cotelor. În exemplul curent, dat fiindcă vom implementa cotele pe partiția principală montată în rădăcină, fișierele trebuie create în același loc. În acest scop, trebuie executate următoarele comenzi:

```
touch /quota.user  
touch /quota.group  
chmod 600 /quota.*
```

Aceste fișiere trebuie create ca root, și vor fi folosite pentru determinarea cotelor individuale și de grup. Ultima comandă are rolul de a bloca orice fel de acces la fișierele respective, prin setarea permisiunilor 600 (adică read-write doar pentru owner – root).

Pasul 4. Instalarea pachetelor necesare

Pachetele necesare pentru implementarea sistemului de cote sunt **quota**, respectiv **quotatool**. Acestea pot fi instalate cu ajutorul comenzii: **apt-get install quota quotatool**.

Pasul 5. Pregătirea sistemului

Dacă dorim ca sistemul de cote să fie implementat fără o repornire a sistemului, partiția în cauză va trebui să fie remontată prin comanda **mount –o remount /**. Pentru a verifica dacă aceasta a fost remontată cu parametrii necesari (menționați la punctul 1), este suficientă executarea comenzii **mount**, care va afișa toate partițiile montate, cât și opțiunile individuale pentru fiecare din acestea.

Pasul 6. Stabilirea cotelor

Odată ce toate părțile sistemului de cote au fost instalate în mod corespunzător, administratorul poate începe să configureze cotele pentru utilizatori și grupuri. Sistemul funcționează folosind două principii de limite,

numite limită „soft” și limită „hard”. Limita „soft” se referă la o limită temporară pe care o poate depăși pentru o perioadă fixă de timp (implicit 7 zile), pe când limita „hard” se referă la limită absolută pe care un utilizator nu o poate depăși în niciun fel.

Pentru stabilirea limitelor individuale administratorul poate să folosească comanda: **quotatool -u (nume_utilizator) -bq (limita soft) -l (limita hard) (mount point)**

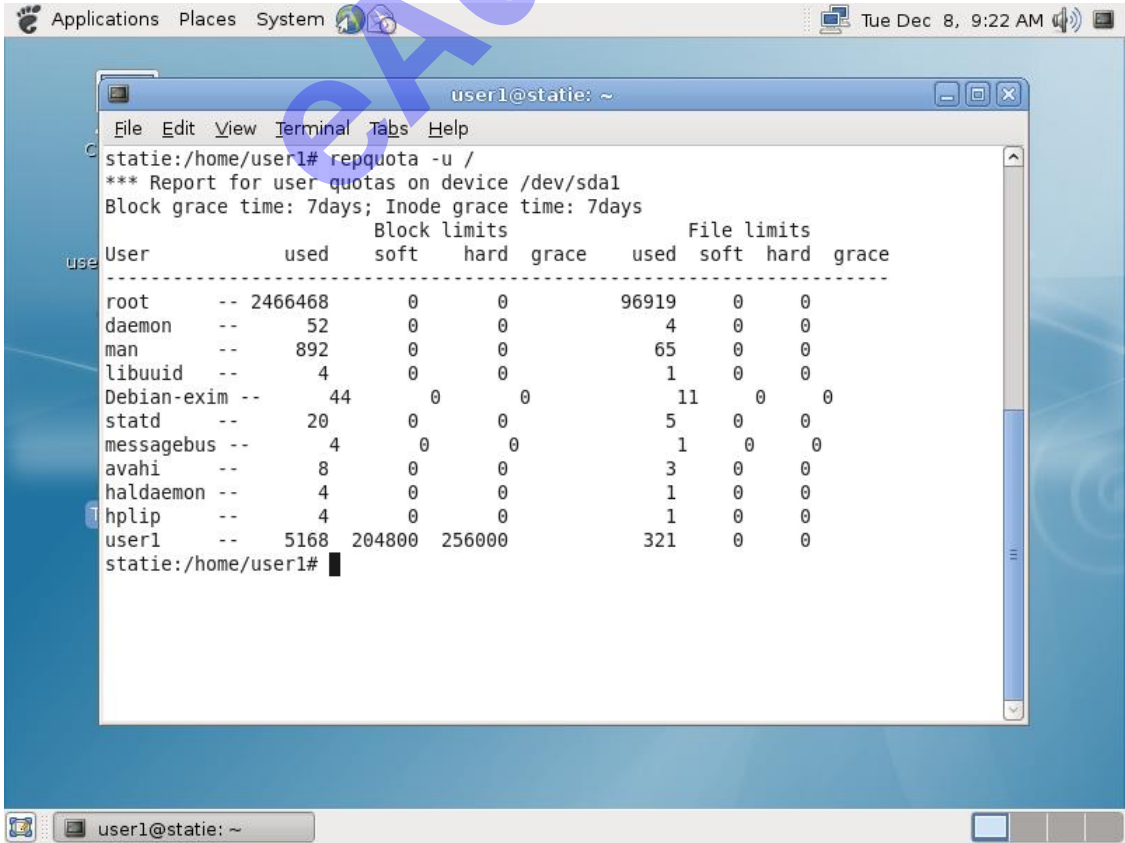
Comanda pentru stabilirea unei cote soft de 200Mb, respectiv 250 limita hard pentru utilizatorul **user1** pe partiția principală a sistemului: **quotatool -u user1 -bq 200M -l 250M /**

Pentru stabilirea limitelor de grup administratorul poate folosi comanda anterioară, însă trebuie schimbat parametrul **-u** (user) cu **-g** (grup). Pentru identificarea utilizatorilor sau a grupurilor, administratorul poate folosi atât numele conturilor cât și identificatoarele numerice (**uid** respectiv **gid**).

Pasul 7. Verificarea sistemului de cote

În cazul în care sistemul a fost implementat corect, în acest moment utilizatorul **user1** ar trebui să fie limitat la 250Mb spațiu de stocare. Administratorul poate oricând să verifice dacă sistemul de cote funcționează, prin apelarea utilitarului **quotacheck -vagum**.

Pentru vizualizarea unui raport privind cotele și spațiul de stocare folosit de fiecare utilizator (și/sau grup) utilizatorul poate folosi comanda **repquota (mount point)**. Raportat la exemplul curent, pentru a vedea cota utilizatorului **user1** pe partiția principală cu punctul de montare **/**, administratorul ar trebui să folosească comanda: **repquota -u /**



```
statie:/home/user1# repquota -u /
*** Report for user quotas on device /dev/sda1
Block grace time: 7days; Inode grace time: 7days

```

User	used	Block limits			grace	File limits		
		soft	hard	grace		used	soft	hard
root	-- 2466468	0	0		96919	0	0	
daemon	-- 52	0	0		4	0	0	
man	-- 892	0	0		65	0	0	
libuuid	-- 4	0	0		1	0	0	
Debian-exim	-- 44	0	0		11	0	0	
statd	-- 20	0	0		5	0	0	
messagebus	-- 4	0	0		1	0	0	
avahi	-- 8	0	0		3	0	0	
haldaemon	-- 4	0	0		1	0	0	
hplip	-- 4	0	0		1	0	0	
user1	-- 5168	204800	256000		321	0	0	

```
statie:/home/user1#
```

În poza anterioară, se poate observa o listă cu cotele impuse pentru fiecare utilizator, cât și un raport cu spațiul ocupat de fiecare cont. Timpul de grație (grace period) este setat în mod implicit la 7 zile, semnificând că utilizatorul poate păstra datele peste limita sa soft timp de 7 zile. În cazul în care utilizatorul și-a atins cota, orice tentativă de scriere nouă vă fi oprită de către sistem, iar utilizatorul va fi atenționat printr-un mesaj de genul „**Disk quota exceeded**”.

OBSERVAȚII: În cazul în care timpul de grație al unui utilizator expiră, sistemul nu va șterge automat nimic, însă utilizatorul respectiv nu va mai avea dreptul să efectueze nici un fel de operație de scriere până când spațiul pe care îl ocupă nu va fi redus sub limita soft. De asemenea, în cazul în care administratorul dorește să elimine cota unui utilizator, acesta va trebui să configureze cota utilizatorului respectiv la valoarea zero.

Configurarea server-ului Xorg

Server-ul **Xorg** face posibilă funcționarea mediului grafic în cadrul unui sistem Linux, și configurarea sa corectă reprezintă un punct foarte important pentru toți ce doresc funcționalitatea unui mediu grafic - deși instalarea ghidată a sistemului a dus la configurarea automată a server-ului **Xorg**.

În majoritatea cazurilor nu va fi necesară intervenția manuală asupra configurării server-ului **X**, însă există cazuri în care s-ar putea ca utilizatorul să-și dorească activarea unor moduri funcționale speciale sau a unor rezoluții care nu sunt disponibile în mod normal. Fișierul de configurare asociat server-ului **X** este **/etc/X11/xorg.conf**.

Fișierul de configurare este structurat în secțiuni ce urmăresc următoarea sintaxă:

```
Section      [nume_secțiune]
Opțiuni...
...
EndSection
```

Secțiunile sunt folosite pentru delimitarea opțiunilor și restricționarea efectului la un punct de interes, și pot lua următoarele valori:

- **Files**
- **Module**
- **InputDevice**
- **Device**
- **VideoAdaptor**
- **Monitor**
- **Modes**
- **Screen**
- **ServerLayout**

La rândul lor, secțiunile pot fi împărțite în sub-sectiuni prin intermediul directivei **SubSection**, în general pentru definirea unui set de caracteristici legate de secțiunea de care aparțin. Un astfel de exemplu, este următoarea

secvență folosită pentru configurarea rezoluției monitorului la una mai puțin convențională; pentru aceasta toate modificările trebuie făcute în fișierul **xorg.conf** în cadrul secțiunii „**Screen**”. În cadrul acesteia se va defini o subsecțiune intitulată „**Display**”:

```
Section „Screen”  
Identifier „Default Screen”  
Monitor „Configured Monitor”  
SubSection „Display”  
    Depth 16  
    Modes „800x600” „1440x900”  
EndSubSection  
EndSection
```

În secvența de mai sus, am definit două moduri de funcționare asociate monitorului descris prin directiva **Monitor** – mai exact am definit rezoluțiile **800x600** respectiv **1440x900**. Important de menționat este că o bună parte a răspunderii pentru modul de funcționare a interfeii grafice îl are și display manager-ul – fie că este vorba de **gdm**, **kdm**, **xdm** sau alte astfel de programe.

eAdmin