

Proiect sisteme de operare - Sistemul de fisiere -

Bondalici Razvan - Sistemul de fisiere Linux. Drive pentru sisteme de fisiere în Windows.

Baluta Dan - Sistemul de fisiere Linux. Drive pentru sisteme de fisiere în Windows.

Bastea Mihai - Introducere în problematica generala a stocării în siguranță a datelor pe un dispozitiv. Drive pentru sisteme de fisiere în Windows. Functii Win32 API

Mantarau George - Introducere în problematica generala a stocării în siguranță a datelor pe un dispozitiv. Drive pentru sisteme de fisiere în Windows. Functii Win32 API

Ionescu Andrei - sistemul de fisiere FAT, Functii Win32 API

Popa Ionut - sistemul de fisiere HFS, Noțiuni de criptografie

Gheorghisan Laurentiu - sistemul de fisiere NTFS

Stanescu Silviu (455A fosta grupa 445A) - Noțiuni de criptografie

Badescu Stefan(455A fosta grupa 445A) - Tehnologii RAID

Capitolul 1.

Introducere în problematica generală a stocării în siguranță a datelor pe un dispozitiv

" Nimeni nu va fi supus la imixtiuni arbitrare în viața sa personală, în familia sa, în domiciliul lui sau în corespondența sa, nici la atingeri aduse onoarei și reputației sale. Orice persoană are dreptul la protecția legii împotriva unor asemenea imixtiuni sau atingeri."

-- Articolul 12 din Declarația Universală a Drepturilor Omului

În era comunicațiilor prin Internet și a stocării electronice a documentelor, intimitatea și confidențialitatea sunt greu puse la încercare, mai mult ca niciodată. Ca să reușești să spargi un seif de oțel este necesar un efort considerabil, însă interceptarea comunicațiilor electronice și intrarea remote pe un calculator poate fi făcută fără prea mult efort și fără prea multe piedici impuse de lege.

Datele importante trebuie protejate, mai ales pe sistemele împărțite de mai mulți utilizatori. Și unde se găsește majoritatea datelor dacă nu pe dispozitivele de stocare (hard discuri, floppy, carduri de memorie, etc.). Pentru a interzice accesul la datele confidențiale este nevoie de un mecanism prin care aceste date să fie protejate.

Protecția datelor de pe un disc se poate face prin trei modalități. Criptarea întregului disc, criptarea unei partiții sau criptarea unui fișier. Problemă cu ținerea ascunsă a datelor prin criptare este că acestea au nevoie „să se miște”. Imaginativă criptarea ca un gard în jurul datelor. Cât timp datele sunt în interior sunt în siguranță. Dar însă pentru a fi în același timp și utile acestea trebuie să poată fi disponibile utilizatorului, de drept, de fiecare dată când are nevoie și fără să fie nevoit să parcurgă prea mulți pași pentru realizarea acestui lucru. Ideal ar fi, pentru utilizator, ca procesul de criptare, să fie invizibil și să nu-i ocupe din timp prin necesitatea autentificării de fiecare dată când are nevoie să acceseze datele.

O altă problemă o reprezintă algoritmi criptografici folosiți, care nu trebuie să fie ușor de spart. Se știe că atunci când cheia secretă are o dimensiune convenabilă și este suficient de frecvent schimbată, spargerea cifrului devine practic imposibilă, chiar dacă se cunoaște algoritmul de cifrare. De aceea e indicat ca dimensiunea cheii să fie mare, dar chiar și așa experiența a arătat că orice schemă criptografică are o viață limitată și că avansul tehnologic reduce, mai devreme sau mai târziu, securitatea furnizată de ea.

Ar fi bine de acoperit de un asemenea program și așa numitul „refuz plauzibil”. Aceasta reprezintă o vulnerabilitate tipic umană, prin care o persoană supusă la stres fizic sau de alta natură (de exemplu sub amenințarea cu o armă), poate divulga parola. O soluție la această problemă ar fi un mecanism care permite ca datele să fie decriptate cu mai mult de o parolă. Rezultatele fiind diferite, cu o parolă se pot decripta niște date neimportante, dar care lasă impresia unei cooperări, iar cu cealaltă parolă să se ascundă datele cu adevărat sensibile.

Firma Microsoft pune la dispoziție un mecanism, numit EFS, care se apropie de cerințele pentru asigurarea securității unui disc.

Ce înseamnă caracteristica EFS (Encrypted File System – Sistem de Fișiere Criptat) a sistemului de fișiere NTFS pe Windows 2000/XP?

EFS permite să criptați datele de pe disc. Salvatul și accesul la fișiere se face normal, dar când se face scrierea pe disc, datele sunt criptate.

Pentru a putea folosi EFS trebuie folosită calitatea de „user” al 2000/XP. Fiecare utilizator ar trebui să aibă un cont propriu protejat prin parolă. Când un utilizator criptează datele, acestea sunt accesibile doar din contul lui. Alți utilizatori nu au acces la datele criptate. Conturile fără parolă nu au nici un efect. Fișierele sistem nu pot fi criptate.

Pentru criptarea unui fișier sau director se apasă click dreapta pe acesta și apoi se selectează properties, se apasă butonul Advanced și se selectează „Encrypt Contents”. În funcție de setări numele fișierelor criptate e colorat diferit.

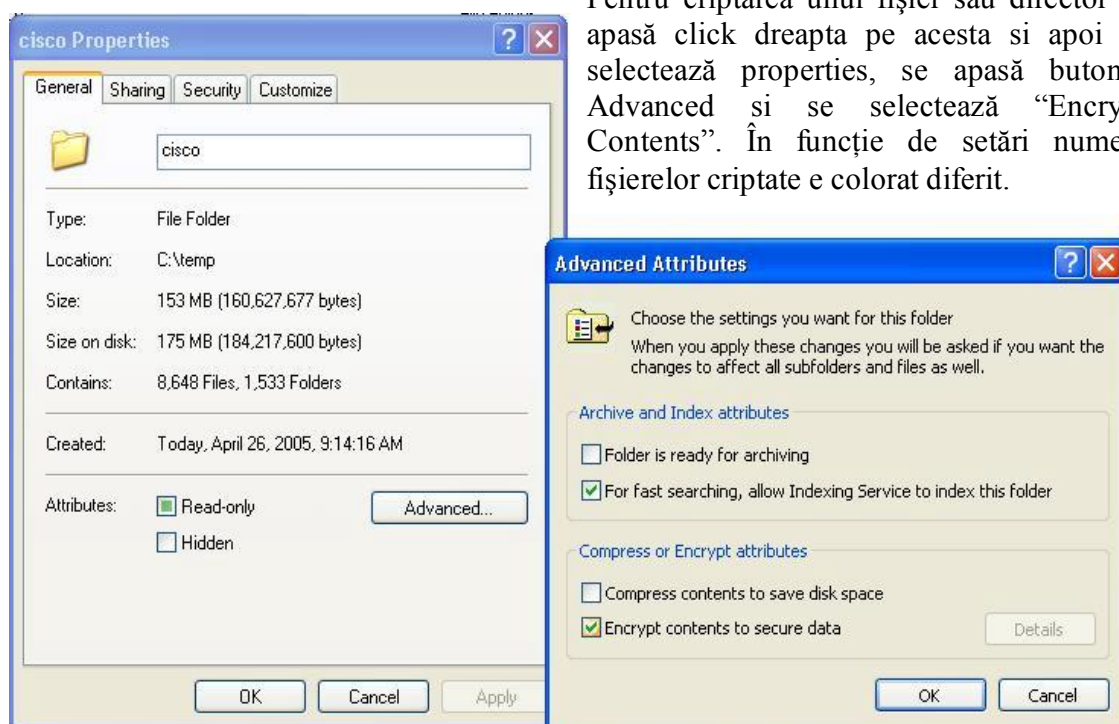


Figura 1.1 Cum se poate cripta un fișier sau director cu EFS.

O problema, des întâlnită, apare la re-formatarea sistemului și la re-instalarea Windows-ului în timp ce exista date criptate pe alt disc și nu există un back-up al “cheilor”. În acest caz nu se vor putea accesa datele criptate. Ce e de făcut? Răspunsul e ștergerea datelor, dacă nu pot fi recuperate folosind un Recover Agent sau dacă nu exista un backdoor (o metodă prin care pot fi sărite metodele obișnuite de autentificare).

Soluția: înainte de formatare trebuie întotdeauna făcut un backup al certificatului cu cheile publice, respectiv private. Cu toate că nu este menționat în timpul procesului de criptare, XP creează pentru fiecare utilizator o pereche de chei privată și publică. Acestea sunt necesare la folosirea EFS. Alți utilizatori nu au acces la chei, astfel neputând accesa datele criptate. Când se re-instalează sistemul, noul certificat asociat contului, nu va funcționa cu datele criptate înainte. De aceea e foarte important ca aceste chei să fie salvate înainte.

Pentru aceasta:

1. Run și comanda „mmc”.
2. În meniu click pe File->”Add/Remove Snap-in...” sau Ctrl+M
3. În tab-ul Standalone apăsați butonul Add
4. Selectați Certificates, apăsați add și apoi închideți și OK
5. Selectați “Certificates - Current User” din „Console Root”, apoi Personal și apoi Certificates
6. Se selectează certificatul care afișează cuvintele Encrypting File System în coloana Intended Purposes . Si apoi se exportă folosind click dreapta All tasks->Export

Folosind apoi Wizard-ul pentru export se salvează cheia privată într-un fișier .pfx. Se va cere și protejarea cu parola a acestor chei. Acum avem cheia exportată și în cazul în care vrem acces la datele criptate după formatare se dă un dublu click pe fișierul exportat și apoi urmărind pașii Wizard-ului.

În afară de acest neajuns EFS mai are și alte dezavantaje.

Unul din cele mai importante dezavantaje ale EFS este că este dependent de sistemul de fișiere NTFS. Un disc formatat FAT sau FAT32 astfel nu poate fi protejat în nici un fel.

Un alt dezavantaj al EFS este că la realizarea criptării acesta creează o copie a fișierului original, criptează copia, șterge originalul și apoi redenumesc copia cu numele original. Aceasta face ca, în primul rând, la criptarea inițială a unui fișier să fie nevoie de spațiu liber pe disc și în al doilea rând, prin copierea și apoi ștergerea fișierului necriptat, folosind un program pentru recuperarea fișierelor șterse, se poate afla conținutul acestuia.

După cum se vede soluția EFS nu este întotdeauna cea mai bună. Astfel apare nevoia de un mecanism care să permită o bună protecție și să lucreze în fundal, rezident în memorie, care să se comporte ca un traductor între containerul datelor (care poate să fie un fișier de pe disc, o partiție a discului sau alte dispozitive de stocare precum un floppy disc) și restul calculatorului, fără programul rezident în memorie și fără parolă datele arătând ca un șir de numere aleatoare.

De asemenea în cazul în care programul rezident în memorie este oprit, hard drive-ul este deconectat subit sau sistemul este oprit, volumul să fie sigur în continuare, bineînțeles aceasta depinzând în mare parte de sistemul de fișiere pe care îl are la bază.

După cum se știe operațiile de read/write de pe disc necesare aplicațiilor sunt realizate de sistemul de operare, în cazul nostru Windows 2000/Xp, cu ajutorul unui driver de disc (care de obicei face parte din sistemul de operare).

În figura 2.1 se poate vedea cum se fac în mod normal operațiile de read/write pe un disc.

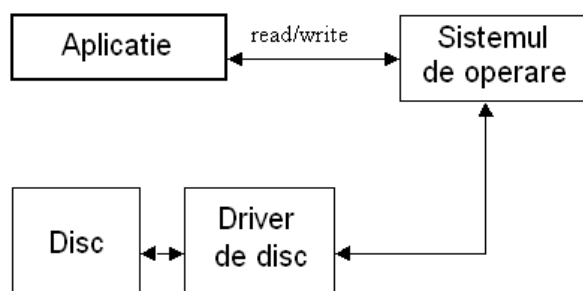


Figura 1.2 Operațiile de read/write pe un disc.

O modalitate prin care se poate interveni pentru criptarea/decriptarea datelor pe ciclurile de read/write este prin construirea unui driver filtru care să prindă datele din zbor (on-the-fly).

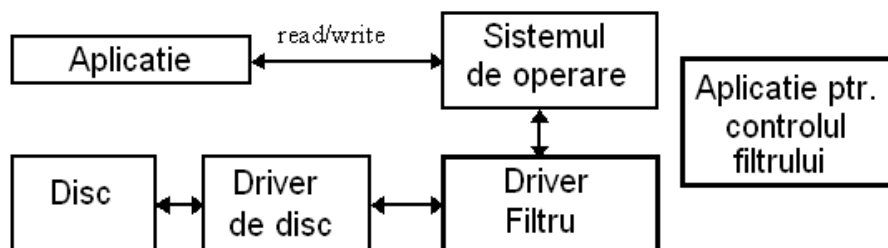


Figura 1.3 Operațiile de read/write pe un disc filtrate de un driver filtru.

Acest filtru nu trebuie să prelucreze toate cererile de I/O. Pentru ca procesul de criptare/decriptare să fie invizibil utilizatorului driverul filtru va crea propriile drive-uri virtuale. Doar operațiile pentru aceste drive-uri vor fi procesate de filtru. Aceste drive-uri virtuale pot fi văzute ca oricare alt disc având literele de drive corespunzătoare (de ex. E:, X: sau oricare litera care nu este deja folosită de un dispozitiv).

În continuare sunt descrise pe larg mecanismele folosite pentru construirea unei aplicații care să permită criptarea unui sistem de fișiere (FAT sau NTFS). În capitolul 2, informații despre cum se poate construi și instala un driver filtru , în capitolul 3, informații despre algoritmi criptografici folosiți, precum și alte detalii privind securitatea.

În capitolul 4 este tratat sistemul de fișiere Linux, în capitolul 5 sistemul de fișiere FAT, în capitolul 6 sistemul NTFS iar în capitolul 7 sistemul HFS al MAC OS.

În capitolul 8 sunt prezentate tehnologii RAID de securizare și stocare a datelor.

Capitolul 2.

Driveri pentru sisteme de fișiere în Windows

Judecând sistemele de operare prin prisma numărului de utilizatori și a domeniilor de întrebuințare, sistemele de operare Windows 2000/ XP sunt considerate ca fiind unele dintre cele mai complete sisteme de operare. Prețul plătit pentru această calitate este complexitatea arhitecturii sistemului, aceasta fiind prezentată mai mult sau mai puțin detaliată în nenumărate lucrări.

O utilizare deosebită a acestui sistem este și aceea de control a dispozitivelor hardware, facilitate pusă la dispoziție cu ajutorul driverelor de dispozitiv, în cadrul acestui capitol se va încerca, în limitele admise, o abordare a tehnologiei de construcție a driverelor de dispozitiv sub sistemele de operare Windows 2000/XP, arhitectura internă a acestor două sisteme fiind asemănătoare.

2.1. Arhitectura Windows 2000/XP

O reprezentare din punct de vedere funcțional a sistemului de operare Windows 2000 este prezentată în figura 2.1. Pentru a înțelege mai bine relațiile dintre diferitele componente, reprezentarea este simplificată în mod deliberat.

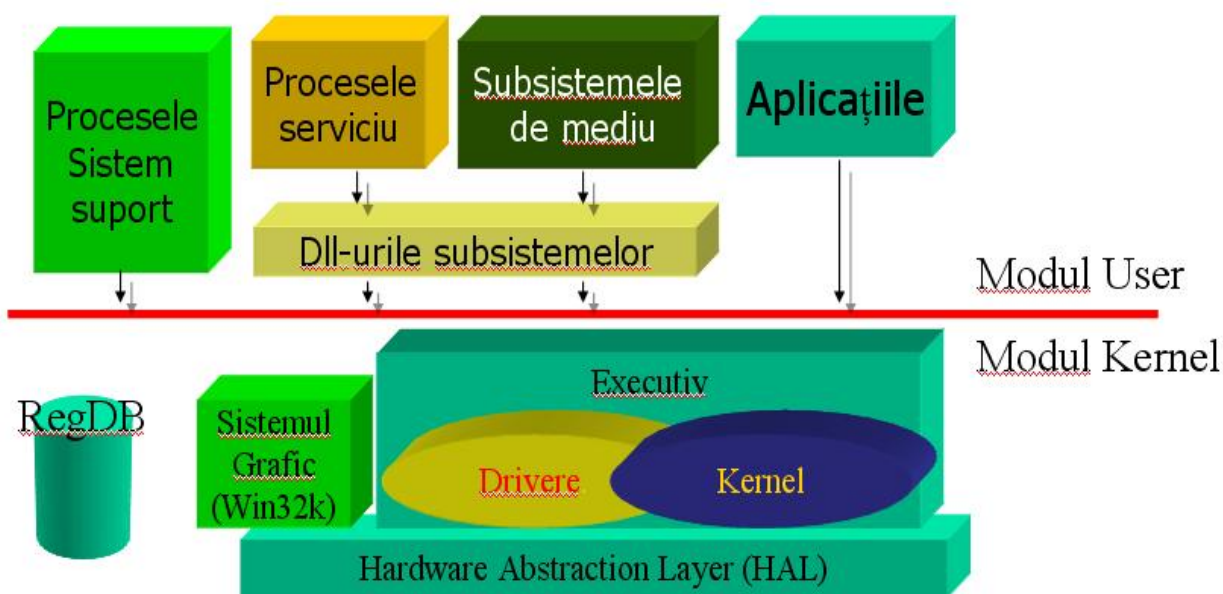


Figura 2.1 Arhitectura sistemului de operare Windows 2000

Înainte de a prezenta elementele care apar în figura 2.1, trebuie spus ce reprezintă modul user și modul kernel, care sunt separate în figură prin linia punctată.

Pentru a îndeplini cerințele impuse sistemului relativ la robustețea și siguranța funcționării, proiectanții tehnologiei NT căreia îi aparține și sistemul de operare Windows 2000 au ales o arhitectură de tip client-server pentru implementarea nucleului sistemului. Astfel, o aplicație utilizator funcționează în calitate de client a serviciilor puse la dispoziție de sistemul de operare. Această arhitectură presupune că orice instrucțiune care se execută pe o mașină pe care este instalat sistemul de operare Windows 2000 se execută în modul user sau în modul kernel. Această împărțire a fost realizată în scopul protecției sistemului de operare considerat parte de încredere din punct de vedere al operațiilor pe care le execută pe o mașină față de aplicațiile utilizatorilor sistemului considerate parte de neîncredere. Având ca scop robustețea și siguranța funcționării, s-a considerat ca părțile de neîncredere să funcționeze în modul user iar părțile de încredere în modul kernel. Diferența de funcționare în cele 2 moduri este aceea că în modul user nu sunt permise o serie de operații precum accesul la tot spațiul de memorie, accesul la porturi, execuția anumitor instrucțiuni ale CPU considerate privilegiate, iar în modul kernel este permisă orice operație. Pentru implementarea acestei arhitecturi de lucru, în general, se apelează la suportul hardware oferit de CPU. De exemplu, procesoarele Intel implementează

hardware aceasta facilitate utilizând nivelele de privilegiu 0 și 3 asociate modului kernel și respectiv user.

Implementarea practică a celor 2 moduri de funcționare ale sistemului de operare Windows 2000 este rezumată în continuare. Prin utilizarea unor structuri de date aflate în memoria calculatorului sistemul de operare gestionează spațiile de memorie virtuală pe care aplicațiile care rulează la un moment dat îl văd(utilizează). Spațiul maxim vizibil este de 4GB. Din acest spațiu o aplicație poate avea acces la cel mult primii 2 GB. Maparea acestui spațiu virtual, care la sistemele de operare Windows este același cu cel liniar, la memoria fizică are loc prin utilizarea altor structuri de date gestionate de sistemul de operare(tabelele de pagini) realizând astfel protecția spațiilor de adrese între aplicații. Ceilalți 2 GB reprezintă spațiul virtual(liniar) unde se găsește sistemul de operare, maparea în memoria fizică urmând același mecanism ca pentru aplicații. Deosebirea, totuși, între cele 2 mapări este aceea că spațiul de maxim 2GB al fiecărei aplicații este mapat în altă zonă a memoriei fizice, pe când cel al sistemului de operare este mapat pe tot timpul funcționării în aceeași zonă astfel că la comutarea de task-uri se schimbă doar maparea primilor 2 GB. acest lucru având o importanță deosebită pentru sistemul de operare(driver) care întotdeauna trebuie să știe spațiul cărei aplicații este vizibil în primii 2GB. Se consideră că. codul dispus în primii 2 GB(aplicație) se execută în modul user putând deci realiza operații restrânse, iar codul dispus în următorii 2 GB se execută în modul kernel putând realiza orice operație. Pentru implementarea optimă a mecanismelor enunțate mai sus Windows 2000 folosește facilitățile hardware oferite de diferitele arhitecturi(ex. Intel).

2.1.1. Procesele utilizator

Revenind la figura 2.1, cele 4 dreptunghiuri din partea superioară a figurii reprezintă cele 4 tipuri primare de procese utilizator care sunt:

1. *procesele sistem suport* - acestea sunt procese fixe care lucrează pentru sistemul de operare în scopul îndeplinirii unor sarcini care pot fi realizate din modul user; exemple de astfel de procese sunt: procesul logon, managerul de sesiune, etc. (a nu fi confundate cu procesele serviciu);
2. *procesele serviciu* - acestea găzduiesc serviciile Win32 precum Task Scheduler sau Spooler; multe aplicații server Windows 2000, precum Microsoft SQL Server, Microsoft Exchange Server includ componente care funcționează ca servicii; o caracteristică a acestor procese este că nu au interfață grafică;
3. *aplicațiile utilizator* - pot fi una din cele 5 tipuri posibile:Win32, Windows 3.1, MS-DOS, POSIX sau OS/2 1.2;

4. *subsistemele de mediu* - sunt acele procese care fac publice aplicațiilor utilizator serviciile native ale sistemului de operare prin intermediul unor funcții apelabile, creând astfel un mediu de funcționare, o "personalitate"; Windows 2000 cuprinde 3 subsisteme de mediu: Win32, POSIX, și OS/2;

În figura 2.1, sub procesele serviciu și aplicațiile utilizator se găsesc dll-urile subsistemelor care au rolul de a intermedia apelul serviciilor sistemului de operare din modul user. Acestea sunt niște librării dinamice care transformă o funcție documentată, apelabilă din modul user, în funcția serviciu corespunzătoare, internă sistemului de operare și nedocumentată; exemple de astfel de librării sunt: ntdll.dll, kernel32.dll, user32.dll, etc.

2.1.2. Componentele kernel ale sistemului de operare

Sub linia punctată din figura 2.1 sunt ilustrate componentele kernel ale sistemului de operare. Aceste componente trebuie înțelese mai degrabă ca niște servicii (în sensul de deținere a capacității de realizare a unei sarcini cerute de un anumit client care poate fi o aplicație utilizator sau o altă componentă a sistemului de operare) pe care sistemul de operare le pune la dispoziția solicitanților. Ca observație, trebuie subliniat faptul că aceste componente nu trebuie înțelese nicidecum ca fiind niște procese independente ca în cazul altor sisteme de operare.

Componentele sunt următoarele:

1. *executivul* - implementează serviciile de bază ale sistemului de operare precum managementul memoriei, al proceselor, al thread-urilor, securității, I/O. sau al comunicației interprocese;
2. *kernelul* - realizează funcții de nivel scăzut precum planificarea thread-urilor, tratarea întreruperilor și a excepțiilor, sau sincronizarea interprocesor;
3. *drivere de dispozitiv* - includ atât driverele pentru dispozitive hardware care translatează cererile de operații I/O emise de aplicațiile utilizator în cereri I/O specifice dispozitivului hardware controlat, cât și driverele sistemului de fișiere și de rețea;
4. *hardware abstraction layer (HAL)* - este un nivel de cod (subrutine) care izolează tot sistemul de operare de hardware-ul mașinii pe care acesta rulează în scopul evitării problemelor datorate din cauza diferențelor între platformele de lucru (de exemplu diferențele între plăcile de bază);

5. *sistemul grafic* - implementează funcțiile de realizare a interfeței grafice(GUI) deosebite oferită de sistemele Windows;

Din punctul de vedere al fișierelor care alcătuiesc nucleul sistemului de operare Windows 2000, prezentat mai sus din perspectiva rolului în funcționare, acesta se compune din fișierele din tabelul 2.1

Nume fișier	Componente incluse
ntoskrnl.exe	executivul și kernelul
hal.dll	HAL (<i>hardware abstraction layer</i>)
Win32k.sys	Partea subsistemului Win32 care rulează în modul kernel
ntdll.dll	Conține "puntea" de trecere din modul user în modul kernel
Kernel32.dll advapi32.dll user32.dll gdi32.dll	Dll-urile nucleu ale subsistemelor

Tabelul 2.1 Fișierele nucleului sistemului de operare Windows 2000

2.2. Driverelor în arhitectura Windows 2000/XP

Relativ la drivere, în figura 2.1 este prezentat locul driverelor kernel în cadrul arhitecturii sistemului Windows 2000 însă aceste sunt de mai multe tipuri, în funcție de criteriul de clasificare. În figura 2.2 sunt prezentate tipurile de drivere sub forma unui arbore, fiecare nivel reprezentând un criteriu iar nodurile reprezentând tipurile de drivere corespunzătoare unui anumit criteriu(nivel).

Pe primul nivel sunt situate driverele privite din perspectiva modului de execuție a codului obiect al acestora:

- *drivere user-mode* - sunt acele drivere care funcționează în modul user; la fel ca și aplicațiile utilizator, rolul lor este acela de a realiza operații care nu necesită nivel de privilegiu ridicat: exemple de astfel de drivere sunt driverele virtuale de dispozitiv(VDD), utilizate de mașina virtuală MS-DOS pentru asigurarea mediului de execuție al aplicațiilor MS-DOS prin translatarea operațiilor privilegiate pe care aceste le execută în operații Win32 care folosesc driverele care rulează în modul kernel;
- *drivere kernel-mode* - sunt acele drivere ale căror funcționare necesită privilegii sporite asupra resurselor hardware și software, cerință

asigurată prin funcționarea în modul kernel, alături de sistemul de operare după cum reiese și din figura 2.1; aceste drivere sunt cele de interes și în cazul acestei prezentări;

O altă clasificare este realizată din perspectiva tehnologiei în care sunt realizate driverele kernel. Unele sunt realizate utilizând o tehnologie moștenită de la versiunea anterioară de Windows NT 4.0 iar altele sunt construite conform modelului WDM. Cele 2 tipuri sunt destul de asemănătoare (modificarea unuia moștenit într-unul WDM este ușoară), însă cele WDM reprezintă ultima tehnologie promovată de firma Microsoft, incluzând un suport puternic pentru tehnologia Plug and Play (PNP), pentru controlul consumului de energie și totodată fiind compatibile din punct de vedere al codului sursă (pentru a fi și binar trebuie recompileate utilizând DDK-ul corespunzător) cu sistemul de operare Windows 98. Plug and Play reprezintă o tehnologie care implică atât hardware cât și software care conlucra pentru implementarea operațiilor de autodetecție și autoconfigurare specifice dispozitivelor moderne.

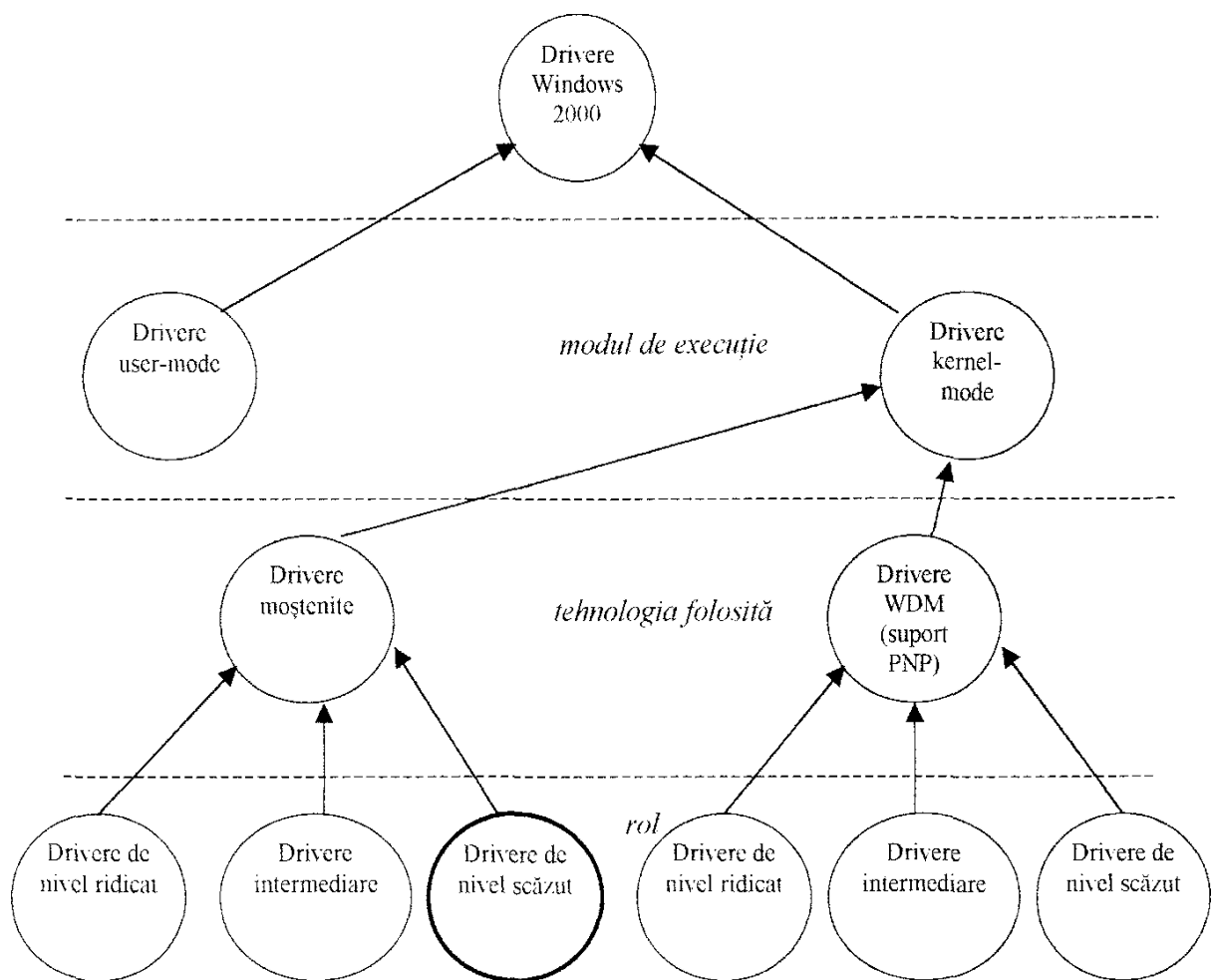


Figura 2.2 Clasificarea driverelor Windows 2000/XP

O caracterizare scurtă a driverelor kernel ar putea fi sintetizată astfel:

- au rolul de intermediar între aplicațiile utilizator sau alte componente ale sistemului de operare, și dispozitivul fizic asociat; ele fac publice tuturor clienților permisi serviciile pe care respectivul dispozitiv le implementează (ex. transport de date, operații grafice, multimedia); altfel zis se poate spune că gestionează funcționarea unui dispozitiv în contextul oferit de sistemul de operare gazdă;
- lucrează alături de sistemul de operare fiind considerate parte de încredere; reversul acestei facilități este că cerințele asupra corectitudinii codului sunt foarte ridicate, orice eroare care în cazul aplicațiilor ar fi prelucrată într-un mod riguros de sistemul de operare, în cazul driverelor poate duce la prăbușirea întregului sistem; toate operațiile de genul alocare/deallocare de memorie, referențiere/dereferențiere handle-uri. etc. trebuie realizate explicit de driver;

- driverul beneficiază de facilitatea apelului direct al serviciilor sistemului de operare însă nu poate apela funcțiile subsistemului Win32 ca în cazul aplicațiilor utilizator (operație fără logică de altfel);
- accesul la driver (apelul subrutinelor acestuia) are loc printr-o interfață standard pe care acesta și-o definește;
- din punctul de vedere al statului său în cadrul sistemului, driverul poate fi văzut ca o bibliotecă (colecție) de funcții încărcată în spațiul de memorie virtual (liniar) al sistemului de operare, funcții care sunt apelate la inițiativa explicită a altcuiva (aplicație, sistemul de operare) și nu precum thread-urile care la momente de timp bine definite li se dă controlul CPU; altfel zis "nu are viață"; ca observație trebuie spus că totuși un driver poate crea thread-uri care au statutul celorlalte thread-uri ale sistemului de operare;
- contextele de execuție ale unui driver sunt:
 - *contextul de trap sau excepție* - este cel mai comun și apare atunci când o anumită aplicație cere o operație I/O care apelează funcțiile Win32 care fac trecerea execuției din modul user în modul kernel utilizând, în cazul arhitecturii Intel de exemplu, o poartă de întrerupere(int 2Eh), după care sunt apelate diferite servicii ale sistemului de operare care în final apelează o anumită subrutină din driver-ul care tratează tipul de cereri I/O emis de aplicație; rezultă că, contextul de execuție (memoria din primii 2GB, etc.) a driver-ului este cel al aplicației (thread-ului) care a cerut operația I/O; acest scenariu este valabil în condițiile în care între timp dintre momentul cererii I/O și apelul codului driver-ului) nu este planificat pentru execuție alt thread, operație prin care se schimbă contextul (primii 2 GB văzuți de driver, etc.), însă această operației planificare) este restricționată prin utilizarea unor mecanisme specifice, bazate pe priorități software de execuție(IRQL) care vor fi prezentate mai târziu;
 - *contextul de întrerupere* - apare în cazul în care un anumit driver tratează o anumită întrerupere care poate apărea la orice moment de execuție, la care nu se cunoaște contextul de lucru;
 - *contextul thread-unlor kernel* - apare în cazul în care un driver creează un thread kernel care din punctul de vedere al planificării seamănă cu thread-urile user, astfel că și în cazul acesta nu se poate face nici o bănuială asupra contextului de execuție iar thread-ul respectiv trebuie să țină cont de acest aspect.

2.3. Dezvoltarea unui Driver Filtru al Sistemului de Fișiere pentru Windows 2000/XP

Majoritatea aplicațiilor moderne necesită o mare performanță și depind de rata de transfer a datelor pe și de pe disc. Pentru aceasta sistemul de operare are nevoie de un driver de dispozitiv scris special pentru dispozitiv.

Pentru înțelegerea modului de construcție a driverelor pentru sistemul de operare Windows 2000, în acest subcapitol se vor prezenta, pașii mai importanți care trebuie urmați pentru dezvoltarea de drivere, respectiv a unui driver filtru pentru un sistem de fișiere.

2.3.1. Caracteristici generale

Resursele software minime necesare pentru dezvoltarea driverelor kernel constau din Device Development Kit (DDK) (indicată fiind versiunea pentru WindowsXP) și produsul Visual C++ .NET din suita Visual Studio .NET 2003, însă se recomandă instalarea și a suportului oferit de Service Development Kit (SDK). Pentru editarea fișierelor sursă se poate folosi orice editor text. Obținerea driverelor (fișiere cu extensia .sys) din codul sursă poate fi realizată fie folosind utilitarul build din DDK (metodă recomandată versiunilor care vor intra în producție) fie se poate utiliza mediul integrat Visual C++ pentru care trebuie setate corespunzător opțiunile referitor la formatul fișierului de ieșire (*.sys), metodă folosită la dezvoltare. Ambele metode fac apel la compilatorul de C și celelalte utilitare necesare pentru obținerea driverului, livrate de firma Microsoft.

Înainte de a prezenta structura driverului, este necesară prezentare generală unor concepte și structuri la care se va face referință pe parcursul prezentării.

IRQL

Implementarea sistemului de operare Windows 2000 pe mai multe platforme hardware(Intel, PowerPC, SPARC) a dus la apariția, printre altele, a problemelor legate de diferitele metode prin care fiecare tip de CPU implementează mecanismul de întreruperi în funcție de priorități. Soluția la această problemă a fost abstractizarea acestui mecanism prin software, definind *nivelele, cererilor de întrerupere* -- interrupt request level (IRQL). În figura 2.3 sunt prezentate toate nivelele cererilor de întrerupere definite de Windows 2000/XP, scopul acestora și sursele care le generează.

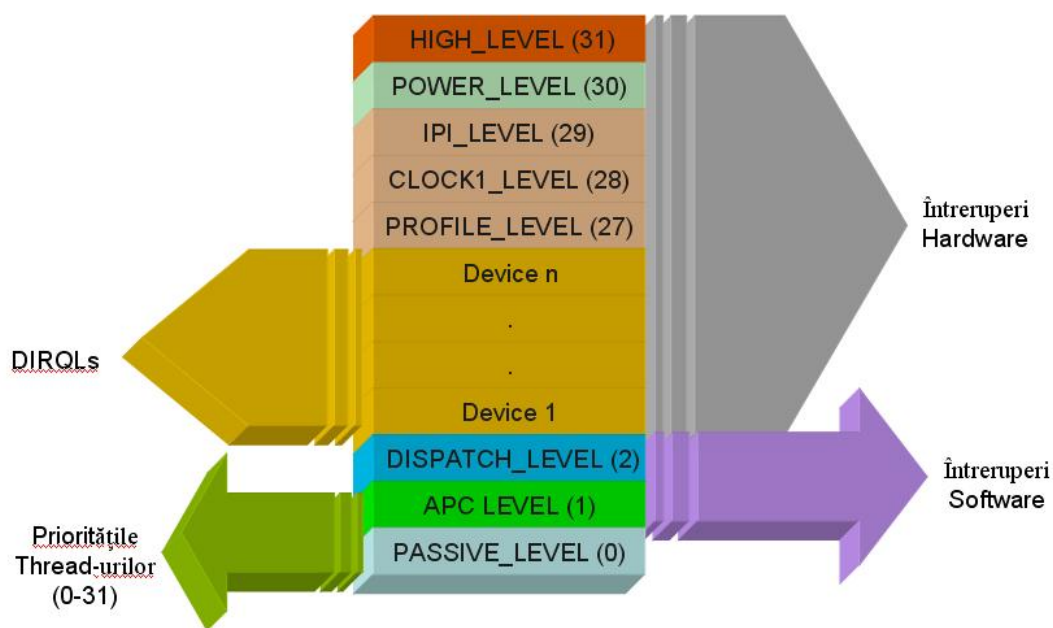


Figura 2.3 IRQL-urile.

Nivelele subliniate în tabelul 2.2 sunt cele mai utilizate în cazul driverelor. Acestea sunt folosite în cazul subrutinelor de tratare a întreruperilor (DIRQLs). În diferite subrutine care trebuie să se execute fără a fi întrerupte software (planificator sau alte thread-uri speciale - DISPATCH LEVEL), în thread-urile create de driver sau alte subrutine ale driverului a căror execuție poate fi întreruptă de altă sursă software (PASSIVE LEVEL).

Numele IRQL	Scopul IRQL	Sursa
HIGHEST_LEVEL	folosit în cazul erorilor de magistrală și a verificărilor hardware	hardware
POWER_LEVEL	folosit în cazul erorilor de alimentare	hardware
IPI_LEVEL	folosit în cazul sistemele multiprocesor	hardware
CLOCK2_LEVEL	folosit pentru clock 2	hardware
CLOCK 1_LEVEL	folosit pentru clock 1	hardware
PROFILE_LEVEL	profiling timer	hardware
DIROLS	folosite pentru maparea priorităților întreruperilor dispozitivelor externe	hardware

DISPATCH LEVEL	folosit la planificarea thread-urilor și în cazul apelurilor procedurilor întârziate (DPC)	software
APC_LEVEL	folosit pentru apelul asincron al procedurilor	software
PASSIVE LEVEL	folosit în cazul execuției normale a thread-urilor	software

Tabelul 2.2 Nivelele cererilor de întrerupere

Trebuie subliniat faptul de a nu se crea confuzii între contextul de execuție și nivelele de prioritate ale CPU (IRQL), în contextul unei aplicații codul putându-se executa la diferite nivele de prioritate (IRQL). Modificarea IRQL curent este realizată explicit prin software din modul kernel utilizând serviciile KeRaiseIrql(...) și KeLowerIrql(...), însă politica de gestiune a nivelului curent de execuție al CPU permite ridicarea nivelului și apoi coborârea acestuia și nu invers, întreruperea execuției unui anumit cod poate fi realizată numai dacă codul care urmează a se executa este mai prioritar din perspectiva IRQL.

Acest mecanism software de organizare a priorităților de execuție a codului pe un anumit procesor introduce anumite restricții în ceea ce privește tipul codului și nivelul de prioritate la care este executat. Motivul provine de la ideea că operații care nu sunt critice funcționării sistemului nu trebuie executate la nivele de prioritate ridicată (și pentru mult timp), care ar priva execuția operațiilor critice. De exemplu este incorect apelul funcției KeWaitForSingleObject la un nivel de prioritate mai mare decât PASSIVE_LEVEL ca de exemplu DISPATCH_LEVEL, din moment ce această funcție nu va face de obicei decât să aștepte apariția unui eveniment. Acest lucru ar putea să priveze, de exemplu, Scheduler-ul (planificatorul) care rulează la nivelul de prioritate DISPATCH_LEVEL să preia controlul la momentul potrivit ducând astfel la disfuncționalități în sistem.

IOCTL

Input/Output Control (IOCTL) reprezintă o structură pe 32 de biți utilizată pentru transmiterea dintr-o aplicație (apelând funcția DeviceIoControl) sau dintr-un driver (apelând IoCallDriver) unui alt driver a unor coduri de control. Prin acest mecanism se poate transporta și date, în afară de codul de control.

I/O Request Packet (IRP)

Aproximativ tot ceea ce reprezintă operații I/O în Windows 2000/XP se desfășoară utilizând ca formă de transfer a informației așa numitele pachete de cerere de intrare/ieșire (I/O Request Packets - IRP). Acestea sunt practic niște structuri de date pe care I/O Manager (componentă a Executivului) le schimbă cu driverele și acestea între ele în cadrul unei operații de intrare/ieșire cu un

dispozitiv. Din punct de vedere funcțional aceste structuri se comportă ca niște *ordine de lucru* transmise de la o componentă la alta. Dintre elementele acestei structuri cele mai importante sunt pointerii la zona de memorie care conține datele schimbate între componente, tipul operației cerute (read, write, deviceiocontrol, etc.), codul de eroare, etc.

Obiecte driver

Implementarea sistemului de operare se bazează în mare măsură pe tehnologia orientată obiect. Fiecare entitate importantă din structura sistemului de operare este reprezentată prin obiecte (thread-uri, drivere, întreruperi, etc.) care reprezintă acea entitate în operațiile în care este implicată.

Obiectele driver (driver objects) au rolul de a reprezenta un driver din punct de vedere al resurselor software cerute de acesta. Membri importanți ai acestui obiect sunt pointerii către subrutinele pe care driverul le face publice și care de altfel reprezintă interfața driverului către sistemul de operare (I/O Manager).

Singura subrutină exportată de un driver este subrutina cu numele (fixat) `DriverEntry` care este apelată la încărcarea driverului în memorie. Rolul principal al acesteia este de a crea interfața driverului prin inițializarea pointerilor obiectului driver cu adresele subrutinelor driverului care devin astfel cunoscute I/O Managerului care până în acest moment nu cunoștea driverul decât prin prisma obiectului driver inițializat la încărcarea driverului în memorie.

Obiecte dispozitiv

Obiectele dispozitiv (device objects) au rolul de a reprezenta din punct de vedere funcțional un anumit dispozitiv pe care driverul asociat îl gestionează. Există un obiect dispozitiv pentru fiecare dispozitiv fizic, virtual sau logic din sistem. Aceste obiecte sunt identificate prin nume unice în sistem (ex. `"\\DosDevices\\VEFS"`). Prin intermediul acestui nume, celelalte componente ale sistemului de operare (nu aplicații) pot interacționa cu driverul asociat. Pentru a fi accesibil din aplicații (user mode) acestui obiect i se creează o legătură simbolică (symbolic link) care din punctul de vedere al aplicațiilor este un nume asemănător cu cel cunoscut kernel-ului (ex. `"\\??\\VEFS"`). Prin această legătură o aplicație poate comunica cu driverul.

Un membru important al obiectului dispozitiv este legătura(pointerul) la lista coadă, pe care sistemul de operare (I/O Manager) o gestionează, prin care IRP-urile sunt semnalizate înainte de a fi prelucrate de driverul pe care acest obiect îl reprezintă. Aceasta este una din facilitățile pe care sistemul de operare le pune la dispoziția dezvoltatorilor de drivere degrevându-i pe aceștia de această

grijă. Totuși, dacă mecanismul gestionat de sistemul de operare nu convine, driverul își poate implementa propriile mecanisme de acces la dispozitiv.

Extensiile dispozitivului

Aceste extensii sunt practic niște blocuri de memorie nepaginabilă alocată în zona sistem (protejată față de aplicații) care au rolul de a menține informații despre un driver pe toată funcționarea dispozitivului. Aceasta este soluția la restricția impusă driverelor, comparativ cu aplicațiile, de a nu folosi variabile globale sau statice.

Spre exemplu structura folosită de driverul filtru.

```
typedef struct EXTENSION
{
    BOOL bRootDevice;
    /* Este acesta dispozitivul rădăcina ? cu care aplicațiile user-mode
    comunică */

    ULONG IMagicNumber;
    /* Pentru a fi sigur că rutina de terminare nu trimite un IRP greșit */

    int nDosDriveNo;
    /* Numarul de Drive pe care această extensie o administrează */
    BOOL bShuttingDown;          /* Se oprește driverul ? */
    BOOL bThreadShouldQuit;      /* Instruiește threadul de dispozitiv
                                să renunțe */
    PETHREAD peThread;           /* Thread handle */
    KEVENT keCreateEvent; /*Evenimentul de creare a dispozitivului */
    KSPIN_LOCK ListSpinLock;     /* IRP spinlock */
    LIST_ENTRY ListEntry;        /* IRP lista de intrare */
    KSEMAPHORE RequestSemaphore; /* IRP lista cu cerere Semafor */

#ifdef USE_KERNEL_MUTEX
    KMUTEX KernelMutex; /* Mutex de Sync. pentru întregul thread */
#endif

    HANDLE hDeviceFile;          /* Handle-ul de dispozitiv pentru
                                dispozitiv */
    PFILE_OBJECT pfoDeviceFile; /* FileObject-ul dispozitivului */
    PDEVICE_OBJECT pFsdDevice;   /* Handle-ul dispozitivului de cel
                                mai de jos nivel */
}
```

```

CRYPTO_INFO *cryptoInfo;    /* Informații Cryptografice */

__int64 DiskLength;        /* Dimensiunea discului menționat de dispozitiv */
__int64 NumberOfCylinders; /* Informații partiție*/
ULONG TracksPerCylinder;  /* Informații partiție*/
ULONG SectorsPerTrack;    /* Informații partiție*/
ULONG BytesPerSector;      /* Informații partiție*/
UCHAR PartitionType;      /* Informații partiție*/

KEVENT keVolumeEvent;     /* Structura eveniment folosită la
                           instalarea unui dispozitiv */

BOOL bReadOnly;            /* Este dispozitivul read-only ? */
BOOL bRemovable;          /* Este dispozitivul removable media ? */
BOOL bRawDevice;          /* Este aceasta o partiție-raw sau un
dispozitiv floppy-raw ? */
BOOL bMountManager;       /* Dacă Mount manager-ul știe despre
                           volum */

WCHAR wszVolume[64];      /* Pentru aplicații din user-mode, cum
                           apare în tree view calea și numele
                           volumului montat, aici sunt stocate doar
                           64 de caractere în loc de MAX_PATH =
                           260 pentru a menține dimensiunea acestei
                           structuri cât mai mică. A nu se schimba
                           dimensiunea aceasta fără a se schimba și
                           dimensiunea din structura
                           MOUNT_LIST_STRUCT! */

long mountTime;           /* Timpul când acest volum a fost montat ultima
                           dată, pentru aplicația din user-mode */

// Data/timpul fișierului container (folosite pentru a reseta data și timpul
// containerelor găzduite, după o încercare de montare nereușită sau demontare).
LARGE_INTEGER fileCreationTime;
LARGE_INTEGER fileLastAccessTime;
LARGE_INTEGER fileLastWriteTime;
LARGE_INTEGER fileLastChangeTime;

} EXTENSION, *PEXTENSION;

```

Obiecte întrerupere

Rolul acestor obiecte este acela de a oferi kernel-ului o modalitate de localizare a subrutinei de tratare a întreruperii.

Între toate obiectele prezentate mai sus există legături foarte strânse prin intermediul pointerilor pentru a permite accesul la toate informațiile acestora pe diverse căi.

În afara de obiectele prezentate mai sus mai există și alte obiecte care însă nu sunt utilizate la fel de des ca cele prezentate.

Înainte de a trece mai departe este necesar a se face o descriere sumară a subrutinelor pe care diferite componente kernel-mode ale sistemului de operare le oferă ca suport de dezvoltare driverelor, în condițiile în care acestea nu pot beneficia de accesul la subrutinele utilizate de aplicații (Win32) din modul user.

În tabelul 2.3 sunt prezentate componentele sistemului de operare, operațiile pentru care oferă suport și nomenclatura subrutinelor pe care acestea le exportă.

Componetă	Oferă suport pentru...	Nomenclatura
Executiv	-alocarea memoriei -gestiunea cozilor -gestiunea listelor lookaside -thread-urile de lucru sistem	ExXxx()
HAL	-accesul la regiștrii dispozitivelor fizice accesul la magistrală	HalXxx()
I/O Manager	-suport general pentru drivere	ToXxx()
Kernel	-sincronizare -deferred procedure call(DPC)	KeXxx()
Memory Manager	-maparea spațiilor virtual și fizic -alocarea memoriei	MmXxx()
Object Manager	-gestiunea handle-urilor	ObXxx()
Proces Manager	-gestiunea thread-urilor sistem	PsXxx()
Runtime Library	-manipularea șirurilor -operații cu numere pe 8 octeți(large integer) -accesul la regiștrii sistemului de operare -funcții de securitate -funcții pentru gestiunea timpului și a datei -suport pentru manipularea cozilor și a listelor	RtlXxx() (marea majoritate)
Security Monitor	-verificarea privilegiilor -funcții pentru manipularea descriptorilor de securitate	SeXxx()

<i>Amestecate</i>	-servicii sistem interne	ZwXxx()
-------------------	--------------------------	---------

Tabelul 2.2 Subrutine suport pentru drivere

2.3.2. Driverul de sistem de fișiere.

Pentru a înțelege modul de lucru al driverelor filtru pentru sisteme de fișiere trebuie înțeles ce sunt driverele de sisteme de fișiere.

Un driver pentru un sistem de fișiere este o componentă a subsistemului de management al spațiului de stocare. Acesta furnizează mijloacele utilizatorilor pentru a stoca și recupera informații de pe mediile nevolatile, precum hard-discurile. Aceste drivere sunt strâns legate cu subsistemele de management al memoriei și al cache-ului (Memory Manager și Cache Manager)

Driverele pentru sisteme de fișiere diferă de celelalte drivere prin câteva caracteristici care le dau un rol special în structura sistemului de operare. Aceste caracteristici sunt:

- Driverele pentru Sisteme de Fișiere garantează să fie apelate în contextual unui thread cerut.
- Driverele pentru Sisteme de Fișiere sunt strâns legate de subsistemele de management al memoriei și al cache-ului.
- Driverele pentru Sisteme de Fișiere sunt interconectate cu managerele de I/O și de Obiecte (I/O și Object Manager).
- Numai Driverele pentru Sisteme de Fișiere implementează puncte de lansare în execuție Fast I/O pentru operațiile de citire și scriere.

Planul modelului Driver pentru Sisteme de Fișiere.

Driverele pentru Sisteme de Fișiere primesc cereri de deschidere, închidere, creare, citire și scriere a fișierelor de pe disc. Aceste cereri de obicei încep de la procesul utilizator și sunt expediate Sistemului de Fișiere prin intermediul subsistemului de management al I/O. Fig. 2.4 descrie cum un Driver pentru Sisteme de Fișiere furnizează servicii unui fir de execuție (thread) utilizator.

Când un fir de execuție (thread) utilizator emite un apel către o funcție de I/O, subsistemul Win32 invocă apelul serviciului corespunzător care cere operația în numele apelantului. Acum procesorul schimbă nivelul de privilegiu la modul kernel. Manager-ul de I/O creează un pachet de cerere I/O (Irp) care descrie cererea de I/O și care apelează Driver-ul pentru Sistemul de Fișiere la punctul de lansare în execuție (entry point) corespunzător. Driver-ul realizează procesarea corespunzătoare și returnează rezultatul manager-ului de I/O, care îl returnează la rândul lui subsistemului Win32 (moment în care se schimbă din

nou nivelul de privilegiu la modul user), și subsistemul Win32 furnizează eventual rezultatul procesului care a realizat cererea.

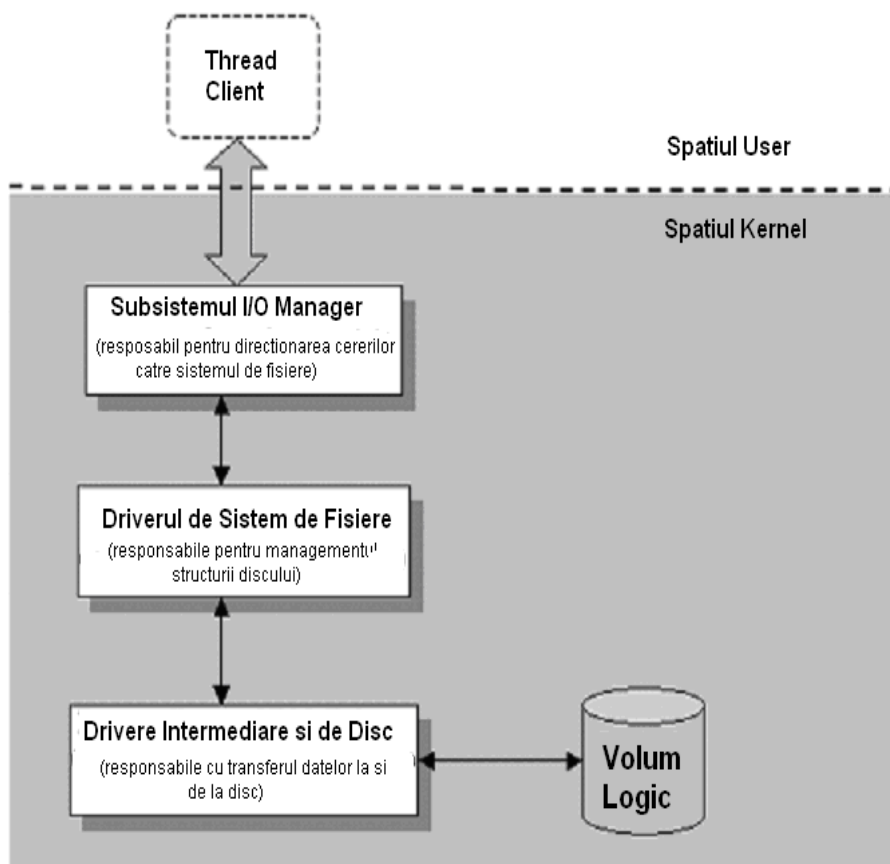


Figura 2.4 Planul modelului Driver pentru Sisteme de Fișiere

De obicei, o cerere de I/O pentru un fișier este transmisă de managerul de I/O driver-ului de sistem de fișiere folosind IRP-uri (Pachete de cereri de I/O). Oricum, cheltuielile asociate cu crearea, completarea și distrugerea Irp-urilor este de obicei un inhibitor al bunei performanțe. De asemenea, dacă Managerul de Cache, depozitează datele, este posibil ca acele date să fie obținute direct de la Cache-ul sistemului redirectând cererea către Managerul de Cache, în loc de a fi obținută prin Irp-uri. I/O rapide (Fast I/O) sunt realizate dacă fișierul este găsit în Cache și este întotdeauna o operație sincronă. Interesant este că dacă operația specifică cu fișierul nu poate fi realizată prin Fast I/O, atunci managerul de I/O folosește metoda Irp standard pentru realizarea operației.

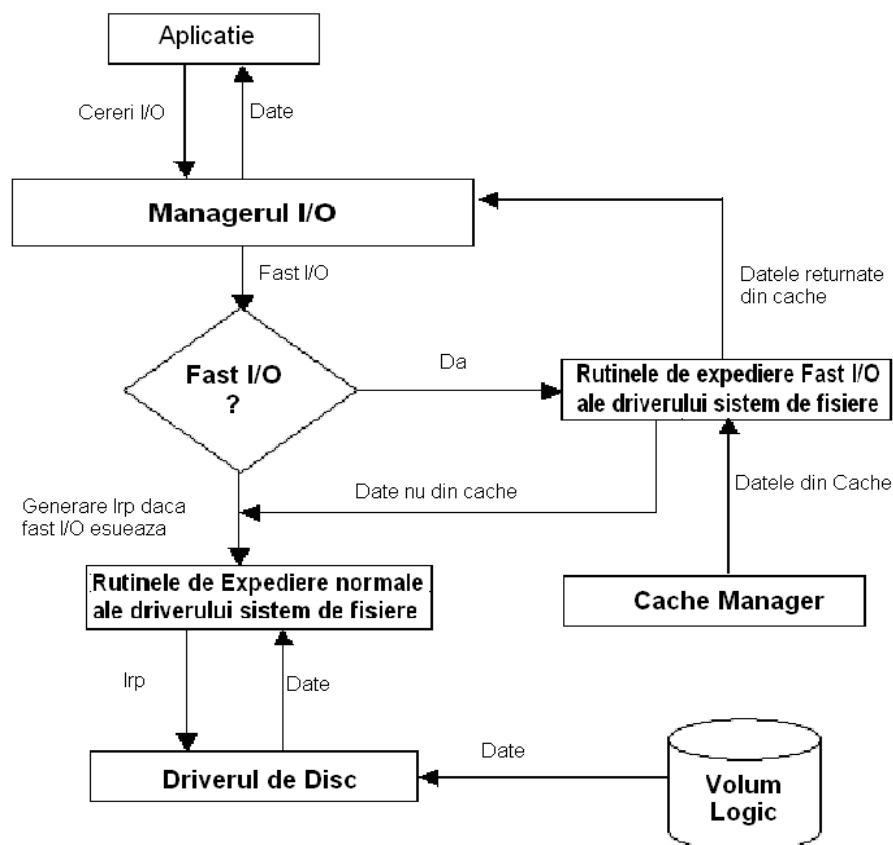


Figura 2.5 Controlul fluxului in modelul driver.

De fiecare data când managerul I/O primește o cerere utilizator pentru a accesa un fișier deschis, acesta invocă punctul de lansare în execuție al Fast I/O, care este de tip Boolean. Aceasta permite driverului să informeze managerul I/O dacă se poate procesa cererea. Dacă da, atunci driver-ul de sistem de fișiere returnează TRUE ca stare a rutinei Fast I/O. Dacă nu, returnează FALSE ca stare a rutinei Fast I/O, în acest caz managerul de I/O creând un IRP care descrie tipul cererii I/O și care apelează driver-ul de sistem de fișiere la entry point-ul de expediere (dispatch) corespunzător. Apoi driver-ul recepționează datele de la disk.

Managementul contextului unui thread

Driverul de sistem de fișiere este întotdeauna logic situat în capul stivei de drivere. Astfel driverul va fi întotdeauna apelat în contextul unui thread care face cererea. Aceasta garantează ca driverele de sistem de fișiere folosesc Neither I/O pentru descrierea cererilor. Implementarea Neither I/O permite unui driver de sistem de fișiere să manipuleze datele folosind adresa virtuală a procesului care face cererea. Astfel, cererile de I/O către driver trebuie pasate driverului în

contextul thread-ului care a inițiat cererea. Dacă un driver filtru schimbă contextul cererii, expediind cererea ca să fie procesată de un thread, driverul de sistem de fișiere va înceta să funcționeze deoarece adresa de referință din spațiul utilizator nu va mai fi validă în contextul noului thread.

Aceasta va duce la eroarea „page fault”, deoarece driverul va referi o zonă de memorie inexistentă. Pentru evitarea acestei catastrofe, driverul filtru trebuie să se asigure că adresa virtuală pentru bufferul, celui care face cererea, este utilizabilă în contextul thread-ului.

Manager-ul I/O construiește automat o MDL (Memory Descriptor List – Lista de Descriere a Memoriei) care descrie buffer-ul apelantului. MDL este o structură care urmărește paginile fizice asociate unui buffer virtual. O MDL constă dintr-un antet care descrie buffer-ul virtual, urmat de un array care listează paginile fizice asociate buffer-ului. Având, adresa virtuală în interiorul bufferului, este posibilă determinarea paginii fizice corespunzătoare.

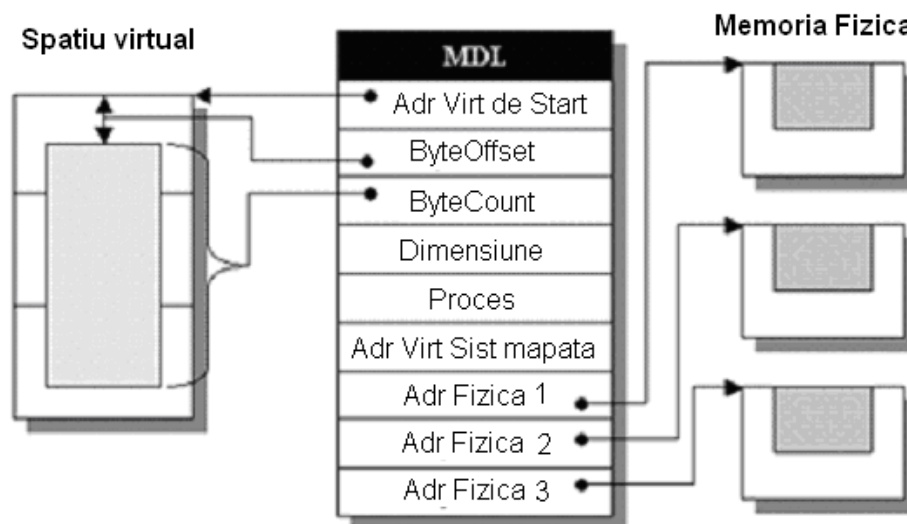


Figura 2.6 Structura unei MDL

De aceea dacă un filtru decide ca are nevoie ca buffer-ul procesului care a făcut cererea sa fie folosibil în contextul threadului, este necesară realizarea următorilor pași:

- crearea unei MDL pentru descrierea bufferului procesului care a făcut cererea
- verificarea accesibilității paginilor care conțin bufferul și blocarea lor în memoria fizică.

-maparea bufferului procesului care a făcut cererea descris de MDL în spațiul virtual de adrese al kernelului.

Un pointer la MDL este furnizat driver-ului într-un IRP (la IrpMdlAddress).

2.3.3. Driverul filtru de sistem de fișiere.

Mecanismul ideal de furnizare a unui I/O sigur este de a implementa un driver filtru deasupra unui driver de sistem de fișiere. Manager-ul I/O permite unui driver din modul kernel să atașeze unul din Obiectele Dispozitiv la un Obiect dispozitiv creat de un driver diferit. Rezultatul este că IRP-ul destinat driverului asociat Obiectului Dispozitiv original, va fi trimis driver-ului asociat cu Obiectul Dispozitiv asociat. Acest driver asociat este Driverul Filtru. Acest driver filtru poate examina, modifica, completa și pasa mai departe IRP-urile trimise driverului original.

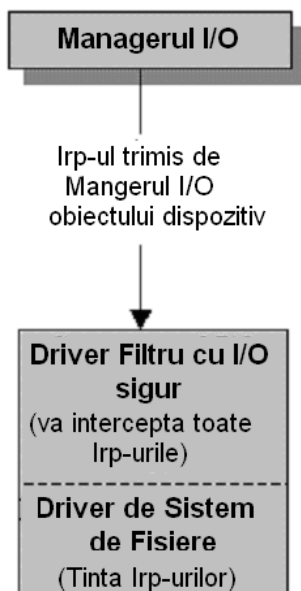


Figura 2.7 Driver filtru atașat

De aceea un driver atașat la un driver de sistem de fișiere va intercepta toate cererile I/O făcute către driverul de sistem de fișiere.

Driverule Filtru

Un driver filtru este un tip special la nivelul driverelor. Este diferit de celelalte drivere prin faptul că este „invisibil”. Acestea se atașează la oricare alt driver și interceptează cererile direcționate către obiectele dispozitiv ale driverelor de mai jos ca ierarhie. Principalul motiv pentru care sunt dezvoltate

este că aduce noi funcționalități față de cel disponibil. Driverul filtru poate folosi, de asemenea, serviciile țintei originale a cererilor de I/O sau chiar poate folosi serviciile altor drivere care lucrează în modul kernel.

Driverele filtru sunt folosite pentru a adăuga caracteristici unui dispozitiv fără să-l modifice driverul de dispozitiv principal sau programele care folosesc acel dispozitiv. Filtrele permit modificarea unor aspecte ale unui driver fără să-l rescrie. Necesitatea folosirii unui filtru driver poate fi ușor explicată.

De exemplu pentru dezvoltarea și implementarea unei criptări/decriptări, directe și în timp ce funcționează, a unui sistem de fișiere, este cea mai simplă soluție. În prezent sistemele de fișiere nu furnizează această capabilitate. Avantajul acestui filtru este că se vor putea folosi în continuare sistemele de fișiere native ale sistemului de operare Windows (de ex. FAT sau NTFS).

Cum preia controlul filtrul

De fiecare dată când este primită o cerere I/O utilizator, de către Manager-ul I/O pentru un fișier de pe un volum logic montat, Manager-ul I/O, în mod normal, înaintează cererea sistemului de fișiere care se ocupă de administrarea volumului montat.

Înainte de a înainta cererea, Managerul I/O verifică dacă alt obiect dispozitiv s-a suprapus peste obiectul reprezentând volumul logic montat și redirecționează cererea către acel obiect dispozitiv, care este în capul listei de obiecte dispozitiv. Astfel driver-ul filtru interceptează I/O înainte de a ajunge la sistemul de fișiere.

De aici, filtrul odată atașat sistemului de fișiere, trebuie să fie garantat, filtrul interceptând toate IRP-urile trimise driverului de sistem de fișiere. Filtrul poate manipula IRP-urile apoi expediindu-le driverului de pe nivelul inferior. Driverul de sistem de fișiere nu are nici o idee că un driver filtru este deasupra, el se comportă ca și cum I/O Managerul trimite cererea utilizatorului direct la el.

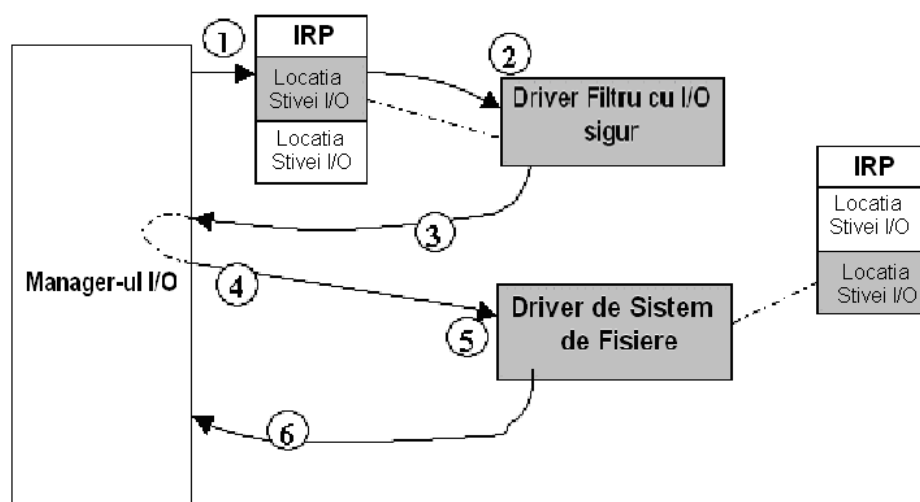


Figura 2.8 Fluxul Irp-urilor

Fig. 2.8 arată cum se desfășoară fluxul IRP-urilor prin sistem după ce un driver filtru garantat a fost atașat unui driver de sistem de fișiere.

Pasul 1. I/O Manager-ul creează IRP-ul pentru operația ce trebuie executată și completează locația din stivă corespunzătoare driver-ului din nivelul cel mai de sus (în cazul nostru – filtrul).

Pasul 2. Filtrul la primirea IRP-ului face operațiile de procesare necesare.

Pasul 3. Filtrul returnează IRP-ul Manager-ului I/O copiindu-și stiva la cea mai de jos locație a stivelor driverelor.

Pasul 4. Managerul I/O pasează IRP-ul următorului driver din listă (în cazul nostru – driver-ul de sistem de fișiere).

Pasul 5. Driverul de la cel mai jos nivel procesează IRP-ul, folosind propria stiva pentru procesare.

Pasul 6. Driverul de sistem de fișiere returnează IRP-ul, I/O Managerului. Acesta eliberează atunci resursele alocate IRP-ului.

Gestionarea lățimii de bandă

Filtrul determină lățimea de bandă totală a discului, prin generarea unor cereri de citire și estimând timpul luat până la servirea acestor IRP-uri.

Driverul filtru interceptează IRP-urile trimise driverului de sistem de fișiere, aceste IRP-uri fiind expediate conform lățimii de bandă necesară aplicațiilor. Lățimea de bandă necesară fiecărei aplicații este specificată printr-o interfață user. Lățimea de bandă reprezintă numărul de Octeți care trebuie accesați, de aplicație, în număr de secunde.

2.3.4. Structura driverului filtru de sistem de fișiere.

Structura unui driver este compusă dintr-o serie de funcții asemănător unei biblioteci. Absolut toate variabilele folosite sunt locale funcțiilor, alocarea lor fiind realizată pe stiva kernel, declararea de variabile globale sau statice fiind interzisă în cazul driverelor. O parte din funcțiile driverului sunt făcute publice componentei IO Manager prin intermediul obiectului driver care reprezintă driverul în relația sa cu I/O Manager-ul. Restul de funcții care nu sunt publice sunt utilizate strict de către driver pentru realizarea anumitor sarcini care necesită o abordare structurată. Un caz special îl reprezintă driverele care își implementează mecanisme proprii prin care alte componente (drivere) pot accesa funcțiile sale. În continuare vor fi prezentate pe rând acele rutine pe care driverul filtru, le pune la dispoziția I/O Manager-ului creând astfel interfața de acces la driver.

Valorile de retur ale funcțiilor apelate de I/O Manager trebuie să fie niște constante definite într-un fișier header inclus în DDK, în funcție de care Managerul I/O ia anumite decizii iar în cazul în care trebuie să dea un răspuns aplicației, setează corespunzător codul de eroare Win32 care poate fi captat din aplicația din modul user cu funcția GetLastError și setează ca valoare returnată de funcțiile apelante (CreateFile, ReadFile, etc.) o valoare diferită de cea pentru succes. Intre codul de eroare captat din aplicație cu GetLastError și membrul IoStatus.Status al structurii IRP, există o mapare gestionată de I/O Manager. Modul de mapare, valorile codurilor de eroare Win32, valorile returnate de funcțiile driverului și semnificația acestora sunt explicate în documentația MSDN livrată de firma Microsoft.

Rutina DriverEntry

Aceasta este singura subrutină exportată de driver la încărcarea sa în memorie și în consecință numele său este fixat, toate celelalte rutine de acces la driver putând avea orice nume.

Această rutină este apelată automat de I/O Manager imediat după încărcarea driver-ului în memorie, parametrul principal transmis acesteia fiind un pointer la obiectul driver pe care I/O Managerul l-a instanțiat și care reprezintă driverul.

Rolul principal al acestei rutine în cazul driverelor non-PNP este acela de a face cunoscute I/O Managerului celelalte rutine de care acesta are nevoie pentru a avea acces la serviciile oferite de driver. În cazul driverelor PNP această sarcină este lăsată în seama altei rutine, CreateDevice, pe care o face publică. Crearea interfeței de acces la driver este realizată efectiv prin inițializarea membrilor obiectului driver cu adresele subrutinelor pe care driverul le publică.

Pentru a realiza acest lucru este necesar ca semnătura acestor funcții să urmeze un șablon fix.

O altă sarcină este aceea de instanțiere a obiectelor dispozitiv, fiecare dintre acestea reprezentând un dispozitiv fizic, logic sau virtual pe care driverul își propune să îl gestioneze. În cazul driverului filtru implementat, deoarece pot fi mai multe dispozitive referite, este necesară crearea unui management riguros al thread-urilor sistem pentru fiecare dispozitiv, astfel cererile care apar, sunt trimise spre tratare rutinei `DispatchQueueIRP` care pune într-o coadă orice IRP ca acesta să poată fi procesat mai târziu de către thread sau în unele cazuri imediat.

Această rutină rulează la `IRQL = PASSIVE_LEVEL`.

Rutina `CreateRootDeviceObject`

Această rutină este chemată imediat după ce rutina `DriverEntry` se termină și inițiază un Obiect Dispozitiv Radăcină. Aici se creează legătura simbolică prin care dispozitivele sunt montate din modul user.

Nivelul de privilegiu al acestei rutine precum și a rutinei `CreateDeviceObject` este `PASSIVE_LEVEL`.

Rutina `CreateDeviceObject`

În mare inițializează Obiectul Dispozitiv, returnând un pointer la Obiectul Dispozitiv, folosit și când nu mai este nevoie de el, pentru a putea fi șters, și setează extensiile acestuia. Este setat câmpul *AlignmentRequirement* la `FILE_WORD_ALIGNMENT` ceea ce înseamnă că un buffer este aliniat în limitele unui cuvânt și deoarece se va face un transfer cu cantități mari de date la un moment dat modul de transfer este setat la Direct I/O. Acest lucru îmbunătățește performanța driverului, reducând atât suprasolicitarea intreruperilor, cât și prin eliminarea alocării memoriei și operațiilor de copiere inerente la Buffered I/O.

Această rutină este apelată de fiecare dată când este primit un IRQ de montare a unui dispozitiv.

Rutina `UnloadDriver`

Este apelată de I/O Manager exact înainte de descărcarea din memorie a driverului. În linii mari face exact operațiile contrare celor realizate în `DriverEntry` în ordinea inversă: înainte de toate demontează toate dispozitivele montate, apoi distruge thread-ul sistem, se "deconectează" de la alte drivere, șterge obiectul dispozitiv și legătura simbolică către aplicații (devenind astfel "invizibil" modului user), reface starea inițială a dispozitivului fizic controlat în caz că aceasta a fost alterată pe parcursul funcționării, etc.

Este apelată la nivelul de privilegiu `PASSIVE_LEVEL`.

Rutina DeleteDeviceObject

Este rutina care șterge legătura simbolică dintre dispozitiv și aplicație sau dacă Obiectul Dispozitiv nu este rădăcină este terminat thread-ul Obiectului Dispozitiv. Este apelată când apare vreo eroare de sistem la montarea unui dispozitiv sau când este demontat un dispozitiv.

Rutinele StartThread, StopThread și ThreadIRP

Prin intermediul acestora se realizează managementul thread-urilor (Crearea, Distrugerea și Procesarea unui thread). Când este creat un Thread - acesta prin intermediul rutinei ThreadIRP face toată munca de procesare a IRP-urilor expediindu-le fie funcției **ReadWrite**, fie funcției **DeviceControl**.

Rutina ReadWrite

Rolul acestei rutine este acela de a efectua transferul de date cu dispozitivul cerut de funcția Win32 apelată (citire de date sau scriere de date la dispozitiv), criptarea și decriptarea datelor pe funcțiile de write și read.

Implementare acestei rutine este diferită de un ReadWrite normal, astfel pentru cererea făcută de I/O Manager (IRP) ca fiind în lucru (pending), la citire, controlul Irp-ului este redat prin apelul rutinei **Completion**.

IRP-ul construit de I/O Manager este trimis driverului ca ordin de lucru pentru realizarea unei operații de intrare, ieșire conține ca informație de bază un buffer, în care se află datele trimise din aplicație sau în care se vor depune datele pentru aplicație și lungimea bufferului specificată în apelul Win32.

Pentru început face o verificare ca offset-ul de început plus lungimea să nu depășească sfârșitul bufferului și lungimea să fie multiplul mărimii unui sector. În cazul apariției unei erori este apelată rutina *IoCompleteRequest*. Această rutină este apelată de obicei când un driver termină procesarea unui IRP, Managerul I/O returnând starea operației apelantului original.

La cererea făcută pentru scriere la dispozitiv se face criptarea buffer-ului folosind EncryptSectors (Buffer, Numarul Sectorului, Numărul de Sectoare, Cheia Secretă, Vectorul Inițial, Algoritmul de criptare).

DispatchReadWrite este apelată la nivelul de privilegiu PASSIVE_LEVEL

Rutina Completion

La întoarcerea din starea în lucru (pending) se realizează decriptarea pe funcția de read folosind DecryptSectors (Buffer, Numarul Sectorului, Numărul de Sectoare, Cheia Secretă, Vectorul Inițial, Algoritmul de criptare), după care se inițializează terminarea procesării IRP-ului.

Această rutină este apelată la un nivel de privilegiu mai mic sau egal cu DISPATCH_LEVEL.

Rutina DeviceIoControl

După ce este creată legătura între apelant și driverul de nivel mai jos este preluat pointerul obiectului dispozitiv, care își inițializează obiectele dispozitiv, precum și un pointer către obiectul fișier corespunzător. Astfel driverul filtru se atașează deasupra driverului de nivel mai jos.

Apoi este apelată funcția *IoBuildDeviceIoControlRequest* care setează IRP-urile pentru cererile de control al dispozitivului ca să fie trimis sincronizat driverului de nivel mai jos, adică driverului de sistem de fișiere. IRP-urile sunt puse într-o coadă de IRP-uri, specifică thread-ului curent. Dacă thread-ul este nevoit să se termine, Managerul I/O anulează IRP-urile.

Parametrii primiți de funcția *DeviceIoControl* sunt:

- handle-ul către driver
- codul de control (IOCTL) transmis de driverul filtru
- adresa buffer-ului care va fi transmis ca dată de intrare driverului; dacă nu sunt transmise date de intrare, este setat la NULL;
- lungimea buffer-ului de intrare; în caz că nu sunt date de intrare este setată la 0;
- adresa buffer-ului de ieșire în care vor fi depuse datele de ieșire ale driverului în cazul în care nu există date de ieșire această adresă este NULL;
- lungimea buffer-ului de ieșire; în caz că nu sunt date de ieșire este setată la 0;

Rutina DeviceControl

Această rutină manipulează anumite cereri venite de la sistemul de operare, necesare sistemului pentru a recunoaște un driver, și de asemenea se ocupă și de funcțiile specifice dispozitivului nostru precum montarea și demontarea.

Este apelată de I/O Manager atunci când o aplicație face apel la funcția *DeviceIoControl*. Aceasta reprezintă o modalitate prin care, o anumită aplicație care încearcă să schimbe date cu un driver de sistem fișiere, este interceptată, anumite cereri fiind manipulate sau chiar fiind adăugate noi funcționalități.

2.3.5. Accesul aplicațiilor la driverul filtru de sistem de fișiere.

Anterior a fost prezentată interfața pe care driverul o oferă sistemului de operare (I/O Manager-ului) pentru a avea acces la serviciile oferite de acesta. În continuare va fi prezentată interfața pe care sistemul de operare (I/O Manager-ul) o oferă aplicațiilor pentru ca acestea să poată beneficia de serviciile oferite de driverul filtru.

Accesul la driver este realizat utilizând o interfață standard pe care Windows 2000 o pune la dispoziție aplicațiilor user-mode pentru accesul la o gamă largă de surse de date: fișiere, pipe-uri, drivere, etc. Cele mai utilizate funcții ale acestei interfețe, sunt:

1. *CreateFile* - are rolul de a livra aplicației un handle în sursa de date în funcție de numele pe care aceasta îl face vizibil aplicațiilor (cale fișier, legătură simbolică, etc); acest handle este utilizat de toate celelalte funcții Win32 ale interfeței;
2. *ReadFile*, *WriteFile* - au rolul de a transfera date între aplicație și sursa de date la care s-a obținut un handle cu *CreateFile*;
3. *DeviceIoControl* - are rolul de a transmite coduri de control (IOCTL);
4. *CloseHandle* - șterge handle-ul obținut cu *CreateFile* din tabela de handle-uri a aplicației, făcând astfel inaccesibilă sursa de date până la un nou apel al funcției *CreateFile*;

În continuare sunt prezentate codurile de control (IOCTL) tratate de către driverul filtru. Acestea sunt împărțite în două categorii, cele din prima parte sunt cele care sunt manipulate, iar în a doua parte controalele private. Pentru fiecare IOCTL vor fi descriși parametrii de intrare, cei de ieșire și codul de eroare returnat de funcția Win32 *GetLastError*.

IOCTL_MOUNTDEV_QUERY_DEVICE_NAME

Este obligatoriu suportul pentru acest control pentru realizarea accesului clientului care administrează montarea. La primirea acestui IOCTL driverul trebuie să furnizeze numele dispozitivului pentru volum, care va fi folosit de către aplicația client ca țintă pentru legătura simbolică. De exemplu numele dispozitivului ar putea fi "*Device\HarddiskVolume1*".

Ca intrare primește *Parameters.DeviceIoControl.OutputBufferLength* aceasta indicând dimensiunea în Octeți a buffer-ului de ieșire. La ieșire clientul care administrează montarea returnează o structură de tipul *MOUNTDEV_NAME*, în partea de început a buffer-ului *Irp->AssociatedIrp.SystemBuffer*. Numele dispozitivului este introdus la adresa indicată de membrul *Name* al acestei structuri. În cazul în care buffer-ul de ieșire este prea mic ca să încapă numele dispozitivului va fi returnată starea de *STATUS_BUFFER_OVERFLOW*.

IOCTL_MOUNTDEV_QUERY_UNIQUE_ID

Este obligatoriu suportul pentru acest control pentru realizarea accesului al clientului care administrează montarea. La primirea acestui IOCTL driverul trebuie să furnizeze un identificator unic pentru volum sau dispozitiv.

Ca intrare primește `Parameters.DeviceIoControl.OutputBufferLength` aceasta indicând dimensiunea în Octeți a buffer-ului de ieșire. La ieșire este returnată o structură de tipul `MOUNTDEV_UNIUE_ID`, în partea de început a buffer-ului `Irp->AssociatedIrp.SystemBuffer`. În cazul în care buffer-ul de ieșire este prea mic ca să încapă numele dispozitivului va fi returnată starea de `STATUS_BUFFER_OVERFLOW`.

IOCTL_MOUNTDEV_QUERY_SUGGESTED_LINK_NAME

Uneori clientul care administrează montarea poate păstra litera de drive după ce sistemul a fost restartat, fără ajutorul managerului de mount. Ca răspuns la acest IOCTL este sugerată o literă de drive managerului de mount. Dacă nu există deja în baza de date a managerului de mount, atunci este folosită litera sugerată, altfel este ignorată. Litera de drive include calea completă a legăturii simbolice ssi are sintaxa clasică MS-DOS. Astfel, de exemplu, litera de drive „E” este reprezentată ca `"\DosDevices\E:"`

Ca intrare primește `Parameters.DeviceIoControl.OutputBufferLength` aceasta indicând dimensiunea în Octeți a buffer-ului de ieșire. La ieșire clientul care administrează montarea returnează o structură de tipul `MOUNTDEV_SUGGESTED_LINK_NAME`, în partea de început a buffer-ului `Irp->AssociatedIrp.SystemBuffer`. Numele dispozitivului este introdus la adresa indicată de membrul `Name` al acestei structuri. În cazul în care buffer-ul de ieșire este prea mic ca să încapă numele dispozitivului va fi returnată starea de `STATUS_BUFFER_OVERFLOW`.

IOCTL_DISK_GET_DRIVE_GEOMETRY

Returnează geometria drive-ului pentru disc. Valorile returnate sunt făcute să se potrivească dimensiunii discului. Acestea sunt tipul, numărul cilindrilor, numărul sectoarelor pe o pistă și numărul Octeților dintr-un sector.

Ca intrare primește `Parameters.DeviceIoControl.OutputBufferLength` aceasta indicând dimensiunea în Octeți a buffer-ului de ieșire. La ieșire driverul returnează o structură de tipul `DISK_GEOMETRY`, în buffer-ul de la `Irp->AssociatedIrp.SystemBuffer`. Returnează `STATUS_SUCCESS` sau în cazul că apar erori următoarele stări

`STATUS_UNRECOGNIZED_MEDIA,`
`STATUS_INVALID_PARAMETER,`
`STATUS_INVALID_DEVICE_REQUEST,`
`STATUS_INFO_LENGTH_MISMATCH,`
`STATUS_INSUFFICIENT_RESOURCES,`
`STATUS_BUFFER_TOO_SMALL.`

IOCTL_DISK_GET_PARTITION_INFO

Returnează informații despre tipul, dimensiunea și natura unei partiții.

Ca intrare primește `Parameters.DeviceIoControl.OutputBufferLength` aceasta indicând dimensiunea în Octeți a buffer-ului de ieșire. La ieșire driverul returnează o structură de tipul `PARTITION_INFORMATION`, în buffer-ul de la `Irp->AssociatedIrp.SystemBuffer`.

IOCTL_DISK_GET_DRIVE_LAYOUT

Returnează informații despre numărul de partiții, semnătura unui disc și caracteristicile fiecărei partiții de pe disc. La ieșire driverul returnează o structură de tipul `DRIVE_LAYOUT_INFORMATION`.

IOCTL_DISK_IS_WRITABLE

Setează dacă tipul poate fi scris sau nu.

IOCTL-urile private.

OPEN_TEST:

Aceasta are un nume reprezentativ și ceea ce face este deschiderea unui dispozitiv pentru a verifica buna funcționare a driverului.

WIPE_CACHE:

Șterge complet memoria cache a driverului. În aceasta sunt stocate maximum 4 parole.

CACHE_STATUS:

Returnează starea cache-ului driverului care poate fi `STATUS_SUCCESS` dacă există informații despre parole în memoria cache.

MOUNT_LIST:

Returnează o listă cu volumele criptate montate, în structura `MOUNT_LIST_STRUCT`. Aceasta are următoarele proprietăți:

```
unsigned long ulMountedDrives;      /* Zonă de biți cu literele tuturor
                                     drive-urilor montate */
short wszVolume[26][64];           /* Numele Volumelor montate */
unsigned __int64 diskLength[26];    /* Dimensiune disc (drive montat) */
int ea[26];                         /* Algoritm criptografic */
BOOL hiddenVol[26];                /* Dacă volumul este de tipul ascuns */
```

VOLUME_PROPERTIES:

Returnează proprietățile volumului montat, pentru aceasta se folosește structura VOLUME_PROPERTIES_STRUCT cu următoarele proprietăți:

```
int driveNo; /*numarul driveului*/
short wszVolume[64]; /* numele volumului*/
unsigned __int64 diskLength; /*dimesiunea volumului*/
int ea; /*o structură care reprezintă algoritmul de criptare */
int pkcs5; /*funcția de hash folosită */
BOOL hiddenVolume; /*dacă volumul este ascuns */
BOOL readOnly; /*dacă volumul este deschis doar pentru citire */

unsigned __int64 volumeCreationTime; /*data creării volumului*/
unsigned __int64 headerCreationTime; /*data creării antetului*/
```

RESOLVE_SYMLINK:

Realizează o legătură simbolică cu ținta. Structura folosită RESOLVE_SYMLINK_STRUCT are următoarele proprietăți:

```
WCHAR symLinkName[MAX_PATH]; /*numele legăturii simbolice*/
WCHAR targetName[MAX_PATH]; /*numele țintei*/
```

MOUNT:

Comanda pentru montarea unui volum. Aceasta folosește structura MOUNT_STRUCT cu proprietățile:

```
int nReturnCode; /*Codul de reîntoarcere de la driver */
short wszVolume[MAX_PATH]; /* Numele volumului ce va fi montat[max 260 caractere] */
char szPassword[MAX_PASSWORD + 1]; /* Parola [max 64+1] */
int nPasswordLen; /* Lungimea parolei */
BOOL bCache; /* Daca există parole în cache-ul driverului */
int nDosDriveNo; /* Numărul de drive care va fi montat */
BOOL bMountReadOnly; /* Daca montarea volumului va fi în modul read-only */
BOOL bMountRemovable; /* Daca montarea volumului va fi ca removable media */
BOOL bExclusiveAccess; /* Deschide fișierul/dispozitivul în modul de acces exclusive */
BOOL bMountManager; /* Anunță volumul managerului de mount */
long time; /* Timpul când a fost montat volumul */
```

Când este cerut controlul mount este apelată rutina **MountDevice**, care face o verificare dacă este introdusă o literă de drive bună, creează Obiectul dispozitiv pe baza structurii mount, initializează Extensiile, creează un thread, apoi este anunțat Managerul de mount de venirea noului volum ssi, în plus, este creată o legatura simbolica cu Managerul de mount.

Pot apărea următoarele erori:

„ERR_BAD_DRIVE_LETTER” – când litera de drive nu este corectă

„CREATE_DEVICE_ERROR” – când nu poate fi creat Obiectul dispozitiv

„STATUS_INSUFFICIENT_RESOURCES” – când nu pot fi alocate resursele pentru începerea unui thread.

UNMOUNT:

Demontează un volum. Aceasta folosește structura UNMOUNT_STRUCT cu proprietățile:

```
int nDosDriveNo;           /* Litera de drive care va fi demontată */
BOOL ignoreOpenFiles;
int nReturnCode;           /* Codul de reîntoarcere de la driver */
```

Când este cerut controlul mount este apelată rutina **UnmountDevice**, după deschiderea și validarea handle-ului volumului, întâi blochează volumul, iar apoi trece la demontarea acestuia și la ștergerea legăturii simbolice.

UNMOUNT_ALL:

Demontează toate volumele, montate prin intermediul acestui driver. Este folosită rutina **UnmountAllDevices**, care conform listei de drive-uri montate face unmount pentru fiecare în parte.

Capitolul 3.

Noțiuni de criptografie

3.1. Introducere în criptografie

Criptografia are o istorie lungă și pitorească. În această secțiune, voi schița doar câteva dintre aspecte, ca informații de bază pentru ceea ce urmează.

Din punct de vedere istoric, la arta criptografiei au contribuit patru grupuri de oameni: armata, corpurile diplomatice, cei ce au ținut jurnale și îndrăgostiții. Dintre acestea, armata a avut rolul cel mai important și a structurat domeniul. În interiorul organizațiilor militare, mesajele ce trebuiau criptate erau de obicei date unor funcționari codori prost plătiți, pentru codare și transmitere. Volumul de mesaje nu permitea ca această muncă să fie făcută doar de câțiva specialiști de elită.

Până la apariția calculatoarelor, una din marile constrângeri ale criptografiei a fost capacitatea funcționarilor codori de a realiza transformările necesare, adeseori pe câmpul de luptă, cu echipament redus. O constrângere suplimentară a fost dificultatea de comutare rapidă de la o metodă la alta, deoarece acesta implica reantrenarea unui număr mare de oameni. Totuși pericolul ca un cod să fie capturat de către inamic a făcut să devină esențială posibilitatea de a schimba metoda criptografică imediat, în caz de nevoie.

Mesajele ce trebuie criptate, cunoscute sub numele de text clar (plain text), sunt transformate printr-o funcție parametrizată de o cheie (key). Ieșirea procesului de criptare, cunoscută sub numele de text cifrat (ciphertex), este apoi transmisă, adeseori prin curier sau radio. Presupunem că intrusul ascultă și copiază cu acuratețe întreg textul cifrat. Totuși, el nu știe care este cheia de criptare și astfel el nu poate decipta prea ușor textul cifrat. Uneori intrusul poate nu numai să asculte canalul de comunicație (intrus pasiv), ci și să înregistreze mesajele și să le retransmită mai târziu, să injecteze propriile sale mesaje sau să modifice mesajele legitime înainte ca ele să fie preluate de receptor (intrus activ). Artă de a sparge cifrul se numește criptanaliză (cryptanalysis). Artă de a concepe cifruri (criptografia) și cea de a le sparge (criptanaliză) sunt cunoscute sub numele de colectiv de criptologie (cryptology).

Criptografia poate fi considerată ca studiul de cum să amesteci datele în așa fel încât:

- oricine fură sau interceptează datele amestecate să nu poată să le refacă

- dar persoanele cărora le sunt destinate acele date să poată să le refacă cu ușurință

3.2. Găuri de securitate care pot fi exploatare.

Nici un sistem de securizare nu este complet sigur. Întotdeauna există metode de infiltrare a unui atacator în sistem. Se pot aminti câteva:

1. În primul rând este bine de avut grijă la aplicațiile care au scurgeri de date. Pot exista bug-uri în program prin care anumite date importante să fie salvate pe disc.

2. Atenție mare și la datele care nu au fost șterse sigur. Orice fișier șters poate fi recuperat destul de ușor. Chiar și fișierele care au fost șterse “sigur” este posibil de a fi recuperate prin scanarea suprafeței discului. Fișierele șterse ar trebui securizate, de exemplu după ștergere zona unde era fișierul de șters să fie umplută cu date complet aleatoare.

3. Fișierele temporare și cele de swap pot fi de asemenea o sursă de scurgere de date. Acestea trebuiesc de asemenea șterse sigur.

4. O metodă mai neobișnuită ar fi atacurile de tip “Vijelie”. Ce înseamnă aceasta? Teoretic, emisiile electromagnetice ale monitorului, hard discului și chiar ale tastaturii pot fi detectate și înregistrate de la distanță. Aceasta permite unui ascultător să vadă ce este pe ecran sau să detecteze parola în timp ce este tastată.

5. *Brute Force* (Forța Brută). Există două metode de forță brută (în mare înseamnă generarea tuturor sau a cât mai multor variante posibile până la reușirea decriptării) : cea care se încearcă asupra parolei și încercarea forței brute asupra algoritmului. Pentru contrarea acestui atac este importantă folosirea unor parole care să fie greu de ghicit, de asemenea ajută dacă sunt folosite atât caracterele mari cât și cele mici, caracterele speciale și cifrele. Este o muncă dificilă care va necesita multe operații.

6. Unii algoritmi pot fi sensibili la atacuri necunoscute publicului larg. Agențiile de securitate **ar putea** deține mecanisme mai rapide decât aplicarea forței brute, sau care exploatează slăbiciunile algoritmilor.

7. Un atac credibil și probabil eficient, este utilizarea forței brute asupra persoanei care deține parola, până ce aceasta este revelată. Factorul uman este, de obicei, cea mai mare gaură de securitate. În același scop pot fi folosite de asemenea și așa numitele „seruri ale adevărului”. Și aceste metode pot fi contracarate printr-un mecanism care să asigure așa numita “negare plauzibilă”.

3.3. Detalii tehnice necesare securizării

3.3.1. Notății

C	Blocul cifrat
$D_k()$	Algoritmul de decriptare folosind cheia de K
$E_k()$	Algoritmul de criptare folosind cheia de K
$H()$	Funcția Hash (ex. RIPEMD-160)
i	Index-ul blocului pentru blocuri de n -biți; n este dependent de context
K	Cheia de criptare/decriptare
m	Index-ul blocului pentru blocuri de 64-biți
n	Index-ul blocului pentru blocuri de 128-biți
P	Bloc cu text în clar
R	Random
W	Valoare ptr XOR
$\wedge$	Sau Exclusiv (XOR)
$\oplus$	Adunare modulo 2^n , unde n este mărimea în biți a celui mai din stânga operand și a valorii rezultate (ex., dacă operandul stâng este o valoare pe 1-bit, iar cel drept este o valoare pe 2-biți, atunci: $1 \oplus 0 = 1$; $1 \oplus 1 = 0$; $1 \oplus 2 = 1$; $1 \oplus 3 = 0$; $0 \oplus 0 = 0$; $0 \oplus 1 = 1$; $0 \oplus 2 = 0$; $0 \oplus 3 = 1$)
$\parallel$	Concatenare

3.3.2. Schema de Criptare

La montarea unui volum se execută următorii pași (aceasta dacă nu sunt parole în cache):

1. Primii 512 Octeți (de ex. antetul unui volum standard) sunt citați în RAM, din care primii 64 Octeți reprezintă valoarea de *salt* (Mai multe detalii la Specificații privind formatul unui volum).
2. 512 Octeți la offset-ul 1536 de la sfârșitul unui volum sunt citați în RAM. Dacă acolo se află un volum ascuns, în acest punct se știe antetul acestuia (indiferent dacă acolo este un volum ascuns trebuie determinat prin decriptarea acestor date).

3. Acum se încearcă decriptarea antetului volumului standard și posibilul antet al unui volum ascuns. Toate datele folosite și generate în timpul procesului de decriptare sunt ținute în RAM. Următorii parametrii nu se cunosc și pot fi determinați prin intermediul metodei aproximărilor succesive (adică să fie testate toate combinațiile posibile ale următoarelor):
 - a. Funcția pseudo-aleatoare folosită pentru derivarea cheii de antet, care poate să fie HMAC-SHA-1 sau AC-RIPEMD-160. Parola introdusă de utilizator și valoarea de *salt* aflată în primii pași sunt pasate funcției de derivare a cheii de antet, care produce o secvență de valori din care sunt derivate cheia de criptare a antetului și vectorul inițial IV-ul (valori folosite pentru decriptarea antetului volumului).
 - b. Algoritmii de criptare AES-256, Blowfish, CAST5, Serpent, Triple DES, Twofish.
 - c. Tipul operației: CBC sau outer-CBC
 - d. Dimensiunea unui bloc.
 - e. Dimensiunea cheilor.
4. Decriptarea este cu succes dacă primii 4 Octeți din datele decriptate conține șirul ASCII "TRUE" și dacă suma de control CRC-32 a ultimilor 256 de Octeți din datele decriptate ale antetului volumului este egală cu valoarea situată la în al 8-lea Octet din datele decriptate. Dacă aceste condiții nu sunt îndeplinite montarea eșuează (principalele motive fiind o parolă incorectă, un volum corupt sau un tip de volum necorespunzător formatat).
5. Acum se presupune că avem parola corectă, algoritmul de criptare, modul, dimensiunea cheii și a blocului, și funcția de derivare corectă a cheii de antet. De asemenea se știe dacă se montează un volum ascuns sau nu. Mai există și verificarea versiunii.
6. Rutina de criptare este reinițializată cu cheia principală și cu vectorul inițial obținut din antetul volumului decriptat. Această cheie poate fi folosită pentru decriptarea oricărui sector al volumului, mai puțin zona cu antetul volumului (care a fost criptat folosind cheia de antet). Apoi volumul este montat.

3.3.3. Modul de cifrare

Există două moduri principale de utilizare în practică a algoritmilor simetrici; cifrarea bloc și cifrarea șir de caractere (secvențială). Cifrarea bloc operează cu blocuri de text clar și cifrat - de regulă de 64 de biți dar, uneori, și

mai mari. Cifrarea secvențială operează cu secvențe de text clar și cifrat de un bit sau octet (uneori chiar și cu cuvinte de 32 de biți), în cazul cifrării bloc, același bloc de text clar va fi cifrat de fiecare dată în același bloc de text cifrat, folosind aceeași cheie, în cazul cifrării secvențiale, secvențe similare de text clar vor fi cifrate diferit, în cazul unor cifrări repetate.

Pentru această aplicație a fost aleasă cifrarea pe blocuri CBC (Cipher Block Chaining – Cifrare bloc cu Înlănțuire) deoarece aduce un surplus de securitate și este perfect pentru criptarea datelor de mari dimensiune precum fișierele sau partițiile. De asemenea are o viteză ridicată.

Modul CBC adaugă mecanismului de criptare un bloc de reacție. Rezultatul rotării unui bloc anterior revine prin buclă și intervine în criptarea blocului curent. În felul acesta, textul cifrat nu mai depinde doar de textul clar, ci și de modul de cifrare a blocului anterior.

Figura 3.1 prezintă modul de criptare CBC. După ce blocul de text clar este criptat, textul cifrat rezultat este stocat într-un registru al buclei de reacție, înainte ca următorul text clar să fie criptat, el este făcut XOR cu blocul din registrul de reacție și devine următoarea intrare în rutina de criptare. După criptare, conținutul registrului este înlocuit cu blocul criptat. În acest fel criptarea blocului i depinde de toate cele $i-1$ blocuri anterioare.

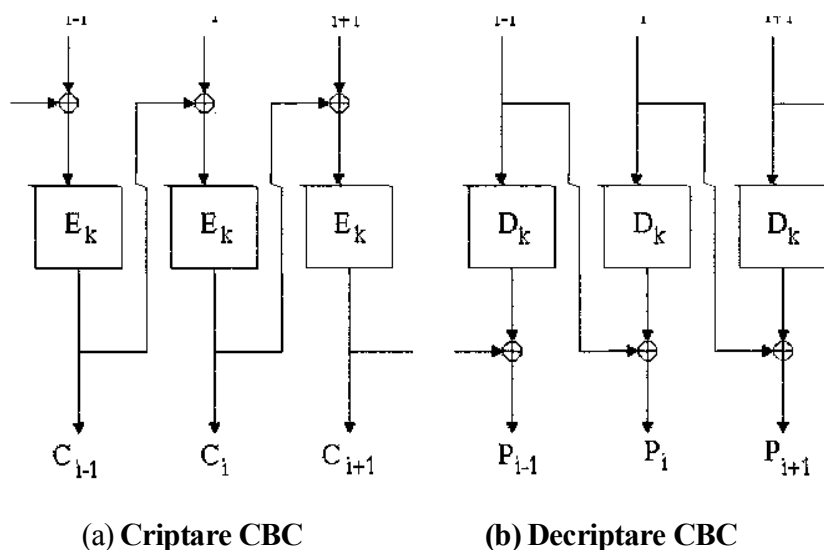


Fig. 3.1 Modul *Cipher Block Chaining*

De asemenea este folosit și un mod de cifrare multiplă, outer CBC.

În următorul tabel pot fi văzuți algoritmi de criptare și modul de cifrare corespunzător fiecăruia.

Algoritmul de Criptare	Modul de Cifrare	Detalii despre Modul de Cifrare
AES-256	CBC	$C_i = E_k(P_i \wedge C_{i-1})$; C_0 este IV.
Blowfish	CBC	$C_i = E_k(P_i \wedge C_{i-1})$; C_0 este IV.
CAST5	CBC	$C_i = E_k(P_i \wedge C_{i-1})$; C_0 este IV.
Serpent	CBC	$C_i = E_k(P_i \wedge C_{i-1})$; C_0 este IV.
Triple DES	Outer-CBC	$C_i = E_{k3}(D_{k2}(E_{k1}(P_i \wedge C_{i-1})))$; C_0 este IV.
Twofish	CBC	$C_i = E_k(P_i \wedge C_{i-1})$; C_0 este IV.

Tabelul 3.1 Algoritmi și modurile de cifrare

3.3.3. IV (Vectorul Inițial)

Este întotdeauna un număr aleator, unic pentru fiecare sector (dimensiunea unui sector este 512 Octeți) și volum. Această valoare este generată astfel:

1. Octeții 256-263 (pentru un cifru pe bloc de 128 –biți este de la 256-271) din antetul decriptat al volumului sunt citați. Dacă cifrurile în cascadă folosesc mai mult de un IV, următorii octeți reprezintă vectorii inițiali suplimentari (de exemplu pentru 3 cifruri pe 128 -biți în cascadă cu modul inner-CBC, primul IV primește Octeții 256-271, al doilea 272-287 și 288-303 pentru al treilea IV).
2. Datele recuperate la pasul 1. sunt făcute XOR cu numărul de sector pe 64-biți (fiecare sector are 512 Octeți, și sunt numerotate de la 0). În cazul unui cifru cu bloc de 128 –biți, valoarea recuperată pe 128 -biți este împărțită în două cuvinte de 64 –biți care sunt făcute XOR cu o valoare identică. Rezultatul pe 64, respective 128 –biți reprezintă IV.

De exemplu:

Ptr. un cifru pe bloc de 128 –biți

T1 = primii 64-biți ai valorii recuperate la pasul (1)

T2 = următorii 64-biți ai valorii recuperate la pasul

S = numărul sectorului (întreg pe 64-biți)

IV = $(T1 \wedge S) \parallel (T2 \wedge S)$

Ptr. un cifru pe bloc de 64 –biți

T = valoarea pe 64-biți recuperată la pasul (1)

$S = \text{numărul sectorului (întreg pe 64-biți)}$

$IV = T \wedge S$

Pasul (1) este realizat doar o dată, imediat după montare, datele recuperate rămânând în RAM.

3.3.4. Whitening (înălbirea)

Pentru a fi și mai dificilă obținerea din text cifrat, textul în clar, este folosită tehnica, care se mai numește și “whitening”. Fiecare 8 octeți ai unui sector (după ce acesta a fost criptat) sunt făcuți XOR cu o valoare pe 64 de biți, unica fiecărui sector și volum. Valoarea este generată după cum urmează:

1. Octeții 264-271 (pentru un cifru pe 128 –biți octeții 272-279) sunt recuperați din antetul decriptat al volumului.
2. Octeții 272-279 (pentru un cifru pe 128 –biți octeții 280-287) sunt recuperați din antetul decriptat al volumului.
3. Datele recuperate la pasul (1) sunt făcute XOR cu numărul de sector pe 64 –biți.
4. Datele recuperate la pasul (2) sunt făcute XOR cu numărul de sector pe 64 –biți.
5. Este calculată valoarea pe CRC-32 pe 32 –biți a primilor 8 octeți ai valorii rezultate la pasul (3).
6. Este calculată valoarea pe CRC-32 pe 32 –biți a următorilor 8 octeți ai valorii rezultate la pasul (3).
7. Este calculată valoarea pe CRC-32 pe 32 –biți a primilor 8 octeți ai valorii rezultate la pasul (4).
8. Este calculată valoarea pe CRC-32 pe 32 –biți a următorilor 8 octeți ai valorii rezultate la pasul (4).
9. Valoarea calculată la pasul (5) este făcută XOR cu cea de la pasul (8).
10. Valoarea calculată la pasul (6) este făcută XOR cu cea de la pasul (7).
11. Valoarea pe 32 –biți calculată la pasul (9) reprezintă primii 32 de biți din valoarea de whitening, iar valoarea pe 32 –biți calculată la pasul (10) reprezintă următorii 32 de biți din valoarea de whitening.

Pașii 5-8 sunt parcurși pentru a crește distanța Hamming (numărul de biți prin care diferă două șiruri binare, se poate scrie ca fiind A și B două șiruri binare distanța este $\sum |A_i - B_i|$) dintre două valori de înălbire.

Pe scurt:

$T1 = \text{valoarea pe 64-biți recuperată la pasul (1)}$
 $T2 = \text{valoarea pe 64-biți recuperată la pasul (2)}$
 $S = \text{numărul sectorului (întreg pe 64-biți)}$
 $T1 = T1 \wedge S$
 $T2 = T2 \wedge S$
 $Q1 = \text{primii 32-biți ai } T1$
 $Q2 = \text{următorii 32-biți ai } T1$
 $Q3 = \text{primii 32-biți ai } T2$
 $Q4 = \text{următorii 32-biți ai } T2$
 $W = (\text{CRC32}(Q1) \wedge \text{CRC32}(Q4)) \parallel (\text{CRC32}(Q2) \wedge \text{CRC32}(Q3))$

Înălbirea este aplicată astfel:

Pentru un cifru pe 128-biți

$T1 = \text{primii 64-biți ai } C$

$T2 = \text{următorii 64-biți ai } C$

$C' = (T1 \wedge W) \parallel (T2 \wedge W)$

Pentru un cifru pe 64-biți

$C' = C \wedge W$

3.3.5. Algoritmii de criptare

Pentru criptarea datelor au fost aleși un număr de 6 algoritmi simetrici.

AES

AES (Advanced Encryption Standard sau Standardul Evoluat de Criptare (numit și Rijndael, după creatorii lui Joan Daemen și Vincent Rijmen, care l-au publicat în 1998) specifică un algoritm criptografic aprobat de FIPS (Federal Information Processing Standards) și poate fi folosit de către departamentele și agențiile federale din SUA pentru protecția criptografică a informațiilor sensibile, dar nesecrete. Acest algoritm poate avea cheile ssi blocurile de dimensiune variabilă.

Folosit este AES cu cheia pe 256 –biți și 14 runde. Acesta este anunțat în iunie 2003, după ce NSA a făcut revizuirea și o analiză asupra acestuia, că este suficient de puternic pentru a proteja informațiile secrete catalogate ca TOP-SECRETE. Astfel produsele care incorporează AES încep să fie considerate satisfăcătoare privind cerințele de Securitate a Informațiilor asociate cu protecția sistemelor de securitate și/sau a informațiilor de securitate națională.

Blowfish

Este realizat de către Bruce Schneier în 1993. Nu este patentat, codul fiind liber pentru toți utilizatorii. Are cheia pe 448 –biți și 16 runde. Acest cifru este cel mai rapid cifru dintre cei implementați.

CAST5

CAST5 sau CAST-128 a fost creat de către Carlisle Adams și Stafford Tavares în 1997. Are cheia pe 128 –biți, iar blocul de 64 –biți. Acest cifru este patentat, dar este liber a fi folosit atât pentru aplicațiile necomerciale cât și pentru cele comerciale. Este de asemenea algoritmul oficial folosit de Guvernul canadian pentru protecția datelor nesecrete.

Serpent

Creat de către Ross Anderson, Eli Biham și Lars Knudsen. A fost publicat în 1998. Are cheia pe 256 –biți, iar blocul de 128 –biți. Serpent a fost unul dintre finaliștii pentru AES. Acesta nu a fost propus ca să fie noul algoritm standard cu toate că părea că are o margine de securitate mai mare decât al lui Rijndael, care a câștigat. Algoritmul Rijndael chiar a fost criticat, sugerându-se că, în viitor, structura matematică poate fi atacată.

Echipa care a creat algoritmul Twofish a făcut un tabel cu coeficienții de siguranță ai algoritmilor finaliști. Coeficientul de siguranță este definit ca numărul de runde a unui cifru de împărțit la cel mai mare număr de runde care au fost sparte. Astfel, dacă un cifru spart are cel mai mic coeficient de siguranță care este egal cu 1, Serpent a avut coeficientul cel mai mare, 3.56 pentru toate dimensiunile de chei suportate, iar Rijndael-256 a avut 1.56.

În ciuda acestui lucru, Rijndael a fost considerată alegerea cea mai bună avându-se în vedere combinațiile dintre performanță, eficiență, implementabilitate și flexibilitate. La sfârșit Rijndael a primit 86 de voturi, Serpent 59, Twofish 31, RC6 23 și MARS 13.

Triplu DES

Publicat în 1978. Este un algoritm cu un mod de criptare tripla (Outer-CBC), are 3 iterații a algoritmului DES (criptare-decriptare-criptare). Folosește 3 chei diferite de 56 –biți.

De remarcat că acest algoritm este mai încet.

Twofish

Creat de Bruce Schneier, David Wagner, John Kelsey, Niels Ferguson, Doug Whiting, și Chris Hall a fost publicat în 1998. Are cheia pe 256 –biți și blocul pe 128 –biți. Acest cifru folosește cutii S dependente de cheie. Acest

algoritm poate fi privit ca o colecție de 2^{128} de sisteme criptografice diferite, unde 128 de biți derivați dintr-o cheie de 256 –biți, controlează sistemul criptografic selectat. Acest lucru poate constitui o formă de securitate împotriva unor atacuri necunoscute.

Algoritm	Dezvoltați de către	Dimensiunea Cheii (Biți)	Dimensiunea Blocului (Biți)	Mod
AES	J. Daemen, V. Rijmen	256	128	CBC
Blowfish	B. Schneier	448	64	CBC
CAST5	C. Adams, S. Tavares	128	64	CBC
Serpent	R. Anderson, E. Biham, L. Knudsen	256	128	CBC
Triple DES	IBM, NSA	3*56	64	Outer-CBC
Twofish	B. Schneier, J. Kelsey, D. Whiting, Wagner, C. Hall, N. Ferguson	256	128	CBC

Tabelul 3.2 Algoritmii de cifrare

3.3.6. Algoritmii de hash

Acești algoritmi sunt folosiți pentru derivarea cheii de antet.

SHA-1

NIST împreună cu NSA au proiectat Secure Hash Algorithm - Algoritm de Dispersie Sigură, standardul numindu-se SHS. Algoritmul procesează un mesaj de lungime maximă $2^{64}-1$ biți și produce un rezumat de 160 -biți.

Lucrează pe blocuri de 512 –biți, folosește 4 funcții matematice și 4 constante aditive.

RIPEMD-160

Algoritmul RIPEMD-160 a fost elaborat de un grup de cercetători de la European RACE Integrity Primitives Evaluation. Algoritmul RIPEMD-160 primește la intrare un mesaj de lungime arbitrară și produce la ieșire un rezumat de 160 de biți.

Lucrează pe blocuri de 512 –biți, folosește 5 funcții matematice și 9 constante aditive.

3.4. Derivarea cheii de antet, a valori Salt și numărul de iterații folosit pentru generarea cheii.

Cheia de antet este folosită la criptarea și decriptarea antetului unui volum. Pentru aceasta este folosită tehnica PBKDF2 specificată în PKCS #5 v2.0.

Este folosită o valoare Salt pe 512 –biți (generată de Generatorul de numere aleatoare la crearea unui volum), ceea ce înseamnă că sunt 2^{512} chei pentru fiecare parolă. Aceasta scade semnificativ vulnerabilitatea atacurilor de tip dicționar (precalcularea tuturor cheilor pentru un dicționar de parole este foarte greu de făcut când este folosită o valoare Salt). Cele 2000 de iterații ale funcției de derivare a cheii de antet fac ca timpul necesar pentru încercarea aflării cheii prin *Brute Force* (Forță Brută), în special atacul de tip Dicționar, să crească foarte mult. Derivarea cheii de antet se face pe baza unei funcții de hash HMAC-SHA-1 sau HMAC-RIPEMD-160. Calitatea și lungimea cheii derivate nu sunt afectate de dimensiunea rezultatului funcției de Hash (care este pentru ambele funcții pe 160 –biți).

3.5. Generatorul de Numere Aleatoare

Este folosit pentru generarea cheii de criptare principale, a Salt-ului și pentru crearea valorilor IV și de whitening.

Acest generator creează, în memoria RAM, o zonă de numere aleatoare de dimensiune 256 –biți. Sursele din care sunt luate datele necesare generării valorilor sunt:

- Mișcările mouse-ului (este făcut un rezumat CRC32 al coordonatelor mouse-ului și a timpului acestui eveniment):
 $\text{CRC32}(\text{MouseCoordinates}) \parallel \text{CRC32}(\text{EventDeltaTime} \parallel \text{Evențime})$
- Apăsatul butoanelor mouse-ului: (este făcut un rezumat CRC32 al ID-ului butonului și a timpului acestui eveniment):
 $\text{CRC32}(\text{MouseButonID}) \parallel \text{CRC32}(\text{EventDeltaTime} \parallel \text{Evențime})$

- Apăsarea tastelor atâta timp cât mouse-ul este situat pe fereastra programului (este făcut un rezumat CRC32 al codului tastei și a timpului acestui eveniment):

$$\text{CRC32}(\text{KeyID}) \parallel \text{CRC32}(\text{EventDeltaTime} \parallel \text{AbsoluteEventTime})$$
- Statistici despre performanța discurilor.
- Librăria MS Windows CryptoAPI.
- Variabile de timp și contoare ale sistemului de operare, colectate la intervale de 250 ms.

Înainte de scrierea unei valori obținute mai sus să fie scrisă în aceea zonă din RAM, valorarea este împărțită în octeți (de ex. rezultatul unui CRC-32 de 32-biți este împărțita în patru octeți). Acești octeți sunt apoi scriși în zonă, dar nu prin înlocuirea vechilor valori, ci prin adunarea modulo 2^8 cu ce se află acolo. La scrierea unui Octet poziția cursorului este crescută cu un Octet. Când cursorul ajunge la sfârșitul zonei, poziția sa este resetată la începutul zonei. În plus după ce o valoare este adusă în zonă, întreaga zonă este rezumată folosind o funcție de Hash (SHA-1 sau RIPEMD-160).

Implementarea și modelul generatorului de numere aleatoare se bazează pe:

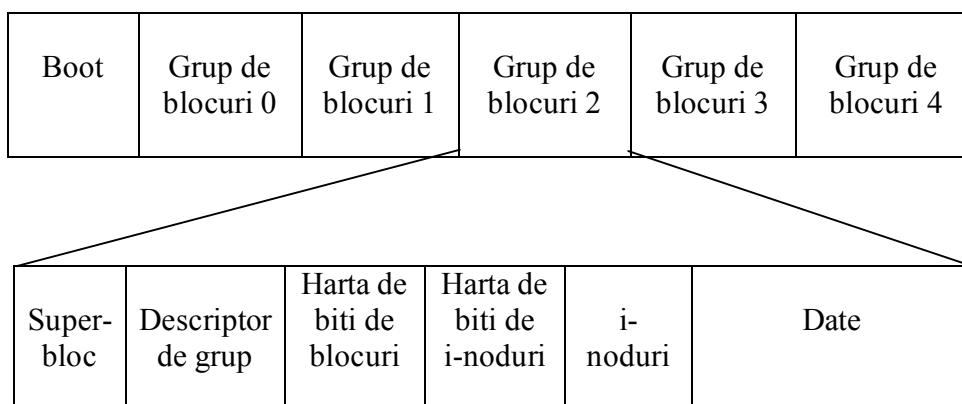
- *Software Generation of Practically Strong Random Numbers* de Peter Gutmann
- *Cryptographic Random Numbers* de Carl Ellison

Capitolul 4

Sistemul de fișiere Linux

Sistemul inițial de fișiere al Linuxului a fost sistemul de fișiere MINIX. Totuși datorită faptului că numele fișierelor erau limitate la 14 caractere și că mărimea maximă a unui fișier era de 64 MB, aproape de la început a existat un interes pentru un sistem mai bun de fișiere. Prima îmbunătățire a fost sistemul de fișiere Ext care permitea nume cu până la 255 de caractere și fișiere de 2 GB, dar era mai lent decât sistemul de fișiere MINIX, deci căutarea a mai continuat un pic. Până la urmă, a fost inventat sistemul de fișiere Ext2 cu nume lungi de fișiere, fișiere mari și o performanță îmbunătățită și acesta a devenit principalul sistem de fișiere. Totuși Linux-ul suportă în continuare o mulțime de sisteme de fișiere folosind sistemul de fișiere NFS. Atunci când Linux-ul este legat, se oferă o alegere de ce sisteme de fișiere să fie legate în nucleu. Altele pot fi încărcate dinamic ca module în timpul execuției, dacă este cazul.

Ext2 este foarte asemănător cu sistemul Berkeley Fast File cu mici diferențe. Structura Ext2 este descrisă mai jos. În loc de a folosi grupuri de cilindre, care aproape nu mai înseamnă nimic cu geometriile virtuale de astăzi, el împarte discul în grupuri de blocuri fără a ține cont de unde cad marginile cilindrelor. Fiecare grup de blocuri începe cu un superbloc, care specifică câte blocuri și i-noduri sunt, mărimea unui bloc, etc. Apoi urmează descriptorul de grup, care conține informații despre poziționarea hărților de biți, numărul de blocuri și i-noduri libere din grup și numărul de cataloage din grup. Această informație este importantă deoarece Ext2 încearcă să împrăștie cataloagele în mod egal pe disc.



Două hărți de biți sunt folosite pentru a ține evidența blocurilor libere și a i-nodurilor libere, respectiv. Fiecare hartă are lungimea de un bloc. Cu blocuri de 1KB, această proiectare limitează un grup de blocuri la 8192 de blocuri și 8192 de i-noduri. Primul număr este o limitare, dar al doilea nu este o limitare în practică. După aceea urmează chiar i-nodurile. Fiecare i-nod este de 128 de octeți de două ori mărimea unui i-nod unix standard. Spațiul suplimentar este folosit în modul următor. În loc de 10 adrese de bloc directe și 3 adrese de bloc indirecte Linux-ul permite 12 directe și 3 adrese indirecte. Mai mult adresele au fost extinse de la 3 la 4 octeți pentru a suporta partiții mai mare decât 224 blocuri (16 GB), ceea ce era deja o problemă în UNIX. În plus, sunt rezervate câmpuri pentru pointeri la liste de acces pentru un control mai fin al drepturilor, dar acestea sunt încă pe planșa de proiectare. Restul i-nodului este rezervat pentru utilizări ulterioare. Istoria a arătat ca biții nefolosiți nu rămân așa pentru mult timp.

Operarea sistemului de fișiere este similară cu a sistemului de fișiere Berkeley rapid. Totuși o diferență față de Berkeley este folosirea a 1KB standard peste tot. Sistemul rapid Berkeley folosește blocuri de 8KB și apoi le împarte în bucăți de 1KB dacă este nevoie. Ext2 face asta în modul simplu. Ca și sistemul Berkeley atunci când un fișier se mărește, Ext2 încearcă să pună blocurile în același grup de blocuri ca și restul fișierului, de preferat imediat după blocul precedent. De asemenea atunci când se adaugă un fișier, Ext2 încearcă să-l pună în același grup de blocuri ca și catalogul său. Cataloagele noi sunt împrăștiate uniform pe disc.

Alt sistem de fișiere Linux este sistemul de fișiere /proc (proces), o idee care a fost inventată inițial în a 8-a ediție a UNIX-ului de la Bell Labs și mai apoi copiată în BSD 4.4 și System V. Totuși, Linux-l extinde ideea în mai multe feluri. Ideea de bază este că pentru fiecare proces din sistem, este creat un catalog în proc. Numele catalogului este PID-ul procesului exprimat ca număr zecimal. De exemplu, /proc/619 este catalogul corespunzător procesului cu PID-ul 619. În acest catalog există fișiere care par să conțină informații despre proces, cum ar fi linia de comandă șirurile de mediu și măștile de semnale. De fapt, aceste fișiere nu există pe disc. Atunci când sunt citite, un apel de sistem obține, pe măsură ce este nevoie, informația de la proces și o întoarce într-un format standard.

Multe dintre extensiile Linux-ului sunt legate de alte fișiere și cataloage din proc. Ele conțin o gamă largă de informații despre UCP, partițiile discului,

dispozitive, vectori de întreruperi, contoare nucleu, sisteme de fișiere, module încărcate și multe altele. Programele utilizator fără privilegii pot citi mare parte din aceste informații și afla comportamentul sistemului într-un mod sigur. În unele dintre aceste fișiere se poate scrie pentru a modifica parametrii de sistem.

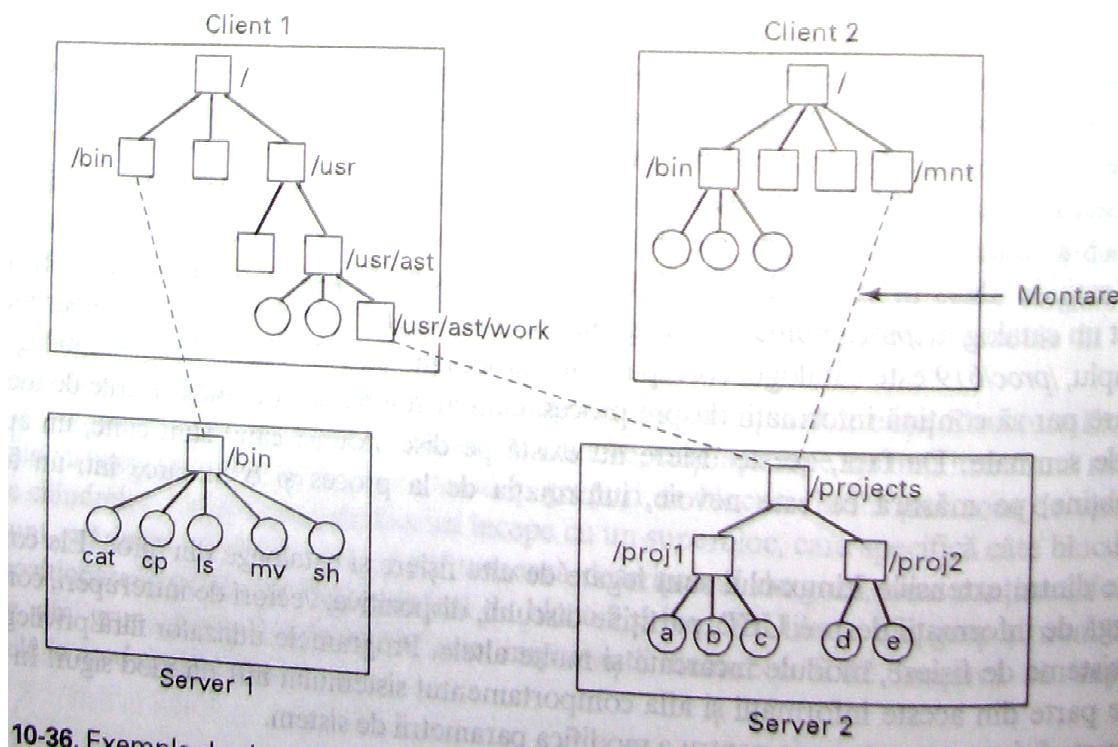
NFS: The Network File System (Sistemul de fișiere de rețea)

Rețelele au jucat un rol important în UNIX încă de la început (prima rețea UNIX a fost construită pentru a muta nucleele noi de pe PDP-11/70 pe Interdata 8/32 în timpul portării pe aceasta din urmă). În această secțiune vom studia sistemul NFS (Network File System) al Sun Microsystems care este folosit de către toate sistemele UNIX moderne (și de către unele sisteme ne-UNIX) pentru a uni sisteme de fișiere de pe calculatoare distincte într-o entitate logică. Sunt interesante trei aspecte ale NFS-ului: arhitectura, protocolul și implementarea.

Arhitectura NFS

Ideea de bază a NFS-ului este să permită unei mulțimi oarecare de clienți și servere să partajeze un sistem de fișiere comun. În multe cazuri toți clienți și serverele se află în același LAN dar aceasta nu este necesar. Este de asemenea posibil să se ruleze NFS peste o rețea de mare întindere dacă serverul este departe de client. Pentru simplitate vom vorbi despre clienți și servere ca și cum ar fi mașini distincte, deși în realitate NFS permite fiecărei mașini să fie și client și server în același timp.

Fiecare server de NFS exportă unul sau mai multe din cataloagele sale pentru a fi accesate de către clienți de la distanță. Atunci când un catalog este făcut disponibil, sunt făcute disponibile și toate subcataloagele, deci de fapt, un întreg arbore de cataloage este exportat unitar. Lista cataloagelor exportate este menținută într-un fisier, deseori `/etc/exports`, astfel încât aceste cataloage să poată fi exportate automat în momentul în care serverul pornește. Clienții accesează cataloagele exportate montându-le. Atunci când un client montează un catalog (la distanță), acesta devine parte a ierarhiei de cataloage, după cum se arată în fig.



Exemple de sisteme de fișiere de la distanță montate. Cataloagele sunt reprezentate ca pătrate și fișierele sunt reprezentate ca cercuri.

În acest exemplu, clientul 1 a montat catalogul `bin` de pe serverul 1 în propriul catalog `bin`, astfel că acum poate referi programul interpretorului ca `/bin/sh` și obține programul de pe serverul 1. Stațiile fără disk au deseori doar un sistem de fișiere schelet (în RAM) și obțin toate fișierele în acest fel de la servere de la distanță. Similar, clientul 1 a montat catalogul `/projects` al serverului 2 în catalogul său `/usr/ast/work` deci poate accesa acum fișierul `a` ca `/usr/ast/work/proj/a`. În final, clientul 2 a montat de asemenea catalogul `projects` și poate de asemenea să acceseze fișierul `a`, ca `mnt/proj1/a`. După cum s-a văzut aici, același fișier poate avea nume diferite pe clienți diferiți datorită faptului că poate fi montat în locuri diferite în respectivii arbori. Punctul de montare este în întregime local pentru clienți; serverul nu știe locația sa de montare pe nici unul din clienți.

Protocoalele NFS

Având în vedere faptul că unul din scopurile NFS este să suporte sisteme eterogene, cu clienți și servere rulând eventual sisteme de operare diferite pe hard diferit, este esențial ca interfața dintre clienți și server să fie bine definită.

Numai atunci este posibil ca cineva să poată scrie o nouă implementare de client și să se aștepte ca aceasta să funcționeze cu serverele existente și invers.

NFS realizează acest scop definind două protocoale client-server. Un protocol (protocol) este o mulțime de cereri trimise de la client server-ului, împreună cu replicile corespunzătoare trimise de către server înapoi la client.

Primul protocol NFS tratează montarea. Un client poate trimite o cale la un server și cere permisiunea de a monta catalogul undeva în ierarhia sa de cataloage. Locul unde va fi montat nu este conținut în mesaj, deoarece pe server nu îl interesează unde va fi montat. Dacă, calea este legală și catalogul specificat a fost exportat, serverul trimite clientului un identificator sau manipulator de fișier (file handle). Identificatorul de fișier conține câmpuri care identifică unic tipul sistemului de fișiere, discul, numărul i-nodului catalogului și informații de securitate. Apelurile următoare de read și write asupra fișierelor din catalogul montat sau orice alt subdirector al său folosesc identificatorul de fișier.

Atunci când UNIX-ul pornește, se rulează fișierul cu comenzi /etc/rc înainte de a intra în modul multiutilizator. Comenzile pentru montarea sistemelor de fișiere de la distanță se pot afla în acest fișier, montând astfel automat sistemele de fișiere de la distanță înainte de a permite orice logare. Alternativ, mare parte din sistemele UNIX suportă automontarea (automounting). Această caracteristică permite ca o mulțime de cataloage de la distanță să fie asociate cu un catalog local. Nici unul dintre aceste cataloage nu este montat (nici măcar serverele lor nu sunt contactate) atunci când clientul pornește. În loc de aceasta, prima oară când este deschis un fișier la distanță, sistemul de operare trimite un mesaj fiecărui server. Primul care răspunde câștigă și catalogul sau este montat.

Automontarea are două avantaje principale asupra montării statice prin fișierul /etc/rc. În primul rând, dacă se întâmplă ca unul din serverele din /etc/rc să fie inactiv este imposibil să se pornească clientul, sau cel puțin nu fără dificultăți, întârzieri și destul de multe mesaje de eroare. Dacă utilizatorul nici nu are nevoie de server în acel moment, toata munca este irosită. În al doilea rând permițând clientului să încerce o mulțime de servere în paralel, se obține un grad de toleranță la defecte (pentru că este necesar să funcționeze unul singur) și performanța poate fi îmbunătățită (fiind ales prunul care răspunde - probabil cel mai puțin încărcat).

Pe de altă parte se presupune că toate sistemele de fișiere specificate ca alternative pentru automontare sunt identice. Deoarece NFS nu pune la dispoziție nimic pentru replicarea fișierelor și a cataloagelor, este treaba utilizatorului să

aranjeze lucrurile astfel încât toate sistemele de fișiere să fie identice. În consecință automontarea este folosită cel mai des pentru sisteme de fișiere numai cu citire care conțin fișierele binare ale sistemului și alte fișiere care se modifică rar.

Al doilea protocol NFS este pentru accesul la fișiere și cataloage. Clienții pot trimite mesaje la server pentru a manipula cataloage și a citi și scrie din fișiere. De asemenea ei pot accesa atributele fișierelor, cum ar fi modul fișierului, mărimea și data ultimei modificări. Marea majoritate a apelurilor de sistem UNIX sunt suportate de NFS, poate cu excepția surprinzătoare a lui open și close.

Omiterea lui open și close nu este un accident. Este intenționată. Nu este necesar să se deschidă un fișier înainte de a fi citit, nici nu este necesar să se închidă când s-a terminat. În loc de aceasta, pentru a citi un fișier, un client trimite serverului un mesaj lookup care conține numele fișierului cu o cerere de a-l căuta și a întoarce identificatorul de fișier, care este o structură care identifica fișierul. Spre deosebire de un apel call acesta operație lookup nu copiează nici o informație în tabelele de sistem. Apelul read conține identificatorul de fișier al fișierului care trebuie citit, offset-ul în fișierul care va fi citit și numărul de octeți doriți. Fiecare mesaj de acest tip este independent. Avantajul acestei scheme este că serverul nu reține nimic despre conexiunile deschise între apeluri. Astfel, dacă un server devine nefuncțional și apoi își revine, nu se pierde nici o informație despre fișierele deschise, pentru că nu există nici o astfel de informație. Se spune despre un astfel de server care nu menține informații despre starea fișierelor deschise că este fără stare (stateless).

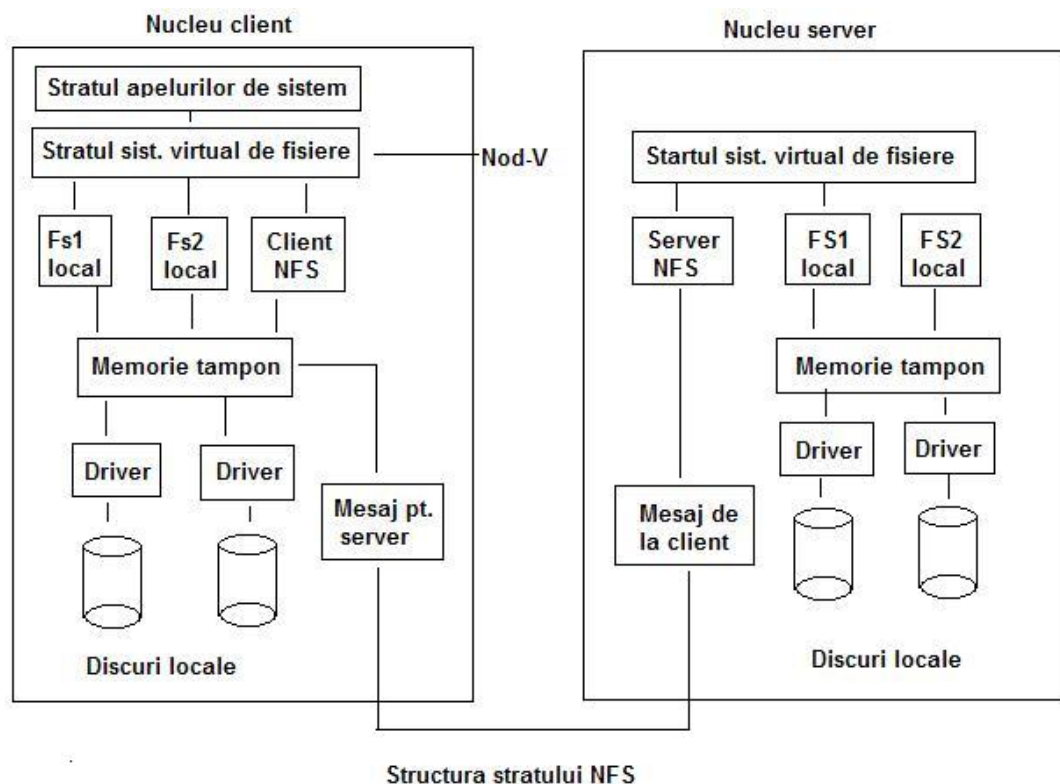
Din păcate, metoda NFS face dificilă obținerea semanticii exacte a UNIX privind fișierele. De exemplu, un fișier UNIX poate fi deschis și blocat astfel încât alte procese să nu îl poată accesa. Atunci când un fișier este închis, blocajele sunt eliberate. Într-un server fără stare ca NFS, blocajele nu pot fi asociate cu fișierele deschise, pentru că serverul nu știe care fișiere sunt deschise. Deși NFS necesită un mecanism adițional separat pentru a trata blocarea fișierelor.

NFS folosește mecanismul de protecție standard UNIX cu biții rwx pentru proprietar, grup și alții. La început, fiecare mesaj de cerere conținea id-urile de grup și utilizator ale apelantului, care erau folosite de serverul de NFS pentru a valida accesul. În realitate avea încredere că acel client nu trișa. Experiența mai multor ani a demonstrat din plin că o astfel de presupunerea era - cum să spunem - naivă. Acum, criptografia cu chei publice poate fi folosită pentru a stabili o

cheie sigură pentru a valida clientul și serverul la fiecare cerere și răspuns. Atunci când această facilitate este pornită, un client răuvoitor nu se poate da drept alt client pentru că nu cunoaște cheia secretă a acelui client.

Implementarea NFS

Deși implementarea clientului și a serverului este independentă de protocoalele NFS, marea majoritate a sistemelor UNIX folosesc o implementare cu trei straturi. Primul strat este stratul apelurilor de sistem. Acesta tratează apeluri cum ar fi open, read și close. După analizarea apelului și verificarea parametrilor, invocă al doilea strat, stratul sistemului virtual de fișiere (VFS). Treaba stratului VFS este să mențină un tablou cu o intrare pentru fiecare fișier deschis, analog cu tabloul de i-noduri pentru fișiere deschise din UNIX. În UNIX-ul obișnuit un i-nod este identificat unic printr-o pereche (dispozitiv, număr i-nod). În loc de aceasta stratul VFS are pentru fiecare fișier deschis o înregistrare numită v-nod (virtual i-node, i-nod virtual). V-nodurile sunt folosite pentru a determina dacă fișierul este local sau la distanță. Pentru fișierele de la distanță, este pusă la dispoziție suficientă informație pentru a putea fi accesate. Pentru fișierele locale, sunt reținute sistemul de fișiere și i-nodul deoarece sistemele UNIX moderne suportă sisteme de fișiere multiple (de ex V7, Berkeley FFS, ext2fs, -proc, FAT, etc.). Deși VFS a fost inventat pentru a suporta NFS, marea majoritate a sistemelor UNIX moderne îl suportă ca o parte din system, chiar dacă NFS-ul nu este folosit.



Pentru a vedea cum sunt folosite v-nodurile, să urmărim o secvență de apeluri de sistem mount, open și read. Pentru a monta un sistem de fișiere la distanță, administratorul de sistem (sau etc/rc) apelează programul mount specificând catalogul de la distanță, catalogul local în care va fi montat și alte informații. Programul mount analizează numele catalogului de la distanță care trebuie montat și descoperă numele serverului de NFS pe care se află catalogul de la distanță. Apoi contactează acea mașină cerând un identificator de fișier pentru catalogul de la distanță. Dacă catalogul există și este disponibil pentru montare la distanță, serverul returnează un identificator de fișier pentru catalog, în final, se realizează un apel de sistem mount, trimițându-se identificatorul nucleului.

Nucleul construiește apoi v-nodul pentru catalogul de la distanță și cere codului client NFS în figura de mai jos să creeze un r-node (remote i-node, i-nod la distanță) în tabelele sale interne pentru a stoca identificatorul de fișier. V-nodul indică către r-node. Fiecare v-node din stratul VFS va conține fie un pointer către un r-nod în codul client NFS, sau un pointer către un i-nod într-unul dintre sistemele locale de fișiere (reprezentați ca linii punctate în figura). Astfel dintr-un v-nod este posibil să se vadă dacă un fișier sau catalog este local sau la distanță. Dacă este local, pot fi localizate sistemul de fișiere corect și i-nodul. Dacă este la distanță, pot fi localizate mașina de la distanță și identificatorul de fișier.

Atunci când un fișier la distanță este deschis pe client, la un moment dat în timpul analizei căii, nucleul se lovește de catalogul în care este montat sistemul de fișiere de la distanță. Vede că acest catalog este la distanță și în v-nodul catalogului găsește un pointer către r-nod. Cere apoi codului client NFS să deschidă fișierul. Codul client NFS caută restul căii pe serverul de la distanță asociat cu catalogul montat și obține un identificator de fișier. Creează un r-nod pentru fișierul de la distanță în tablourile sale și raportează stratului VFS, care pune pentru fișier în tablourile sale un v-node care punctează către r-node. Încă o dată vedem aici că fiecare fișier sau catalog deschis are un v-nod care punctează fie către un i-nod sau un r-nod.

Apelantului îi este dat un identificator de fișier pentru fișierul la distanță. Identificatorul de fișier este pus în corespondență pe v-nod de către tablourile din stratul VFS. Este de observat că nu se realizează nici o înregistrare în tabelele de la server. Deși serverul este pregătit să pună la dispoziție identificatori de fișiere la cerere, el nu ține evidența a ce fișiere au identificatori de fișiere și care nu. Atunci când îi este trimis un identificator de fișier pentru a fi accesat, el verifică identificatorul și dacă este valid, îl folosește. Validarea poate include verificarea cheii de autentificare conținută în antetele RPC, dacă securitatea este activată.

Atunci când identificatorul de fișier este folosit pentru un alt apel de sistem, de exemplu, read, stratul VFS localizează v-nodul corespunzător și din acesta determină dacă este local sau la distanță și de asemenea care i-nod sau r-nod îl descrie. Trimite apoi un mesaj serverului conținând indentificatorul, offsetul în fișier (care este reținut la client, nu la server) și numărul de octeți. Din motive de eficiență, transferurile dintre client și server se realizează în bucăți mari, în mod normal 8192 octeți, chiar dacă sunt ceruți mai puțini octeți.

Atunci când cererea ajunge la server, este pasată stratului VFS, care determină care sistem local de fișiere are fișierul cerut. Stratul VFS face apoi un apel către acel sistem de fișiere pentru a citi și a returna octeții. Aceste date sunt apoi trimise înapoi clientului. După ce stratul VFS client a obținut bucata de 8-KB care a cerut-o, face automat o cerere pentru următoarea bucată, pentru a o avea dacă va fi cerută în scurt timp. Această facilitate, cunoscută ca citire înainte, îmbunătățește performanța considerabil.

Pentru scrieri este urmată o cale asemănătoare de la client la server. De asemenea transferurile se realizează tot în bucăți de 8-KB. Dacă un apel de sistem write pune la dispoziție mai puțin de 8-KB de date, datele sunt doar acumulate local. Numai când întreaga bucată de 8-KB este plină este trimisă serverului. Totuși, atunci când un fișier este închis datele sale sunt trimise la server imediat.

Altă tehnică folosită pentru îmbunătățirea performanțelor este folosirea memoriei intermediare ca și în UNIX-ul obișnuit. Serverul folosește memoria intermediară pentru date pentru a evita accesele la disc, dar aceasta este invizibil clientului. Clienții mențin două memorii intermediare (zone tampon), una pentru atributele fișierelor (i-noduri) și una pentru datele fișierelor. Atunci când este nevoie de un i-nod sau bloc de fișier, se face o verificare dacă cererea poate fi satisfăcută folosind zona tampon. Dacă da, se va evita traficul în rețea.

În timp ce zona tampon de la client ajută enorm la mărirea performanțelor, el și introduce niște probleme urâte. Să presupunem că doi clienți au în zona tampon același bloc dintr-un fișier și că unul din ei îl modifică. Atunci când celălalt îl va citi, va citi versiunea veche. Zona tampon nu este coerentă.

Având în vedere potențiala gravitate a acestor probleme, implementarea NFS face mai multe lucruri pentru a o atenua. Întâi, fiecărui bloc din zona tampon îi este asociat un cronometru. Atunci când cronometrul expiră, înregistrarea este îndepărtată. În mod normal cronometrul este setat la 3 secunde pentru blocuri de date și 30 de secunde pentru blocuri catalog. Realizând aceasta se reduce într-un fel riscul. În plus, ori de câte ori este deschis un fișier din zona tampon este trimis un mesaj la server pentru a afla când a fost modificat ultima oară fișierul. Dacă ultima modificare a survenit după ce copia locală a fost

intodusă în zona tampon, copia din zona tampon este îndepărtată și noua copie este preluată de la server. În final a fiecare 30 de secunde expiră un cronometru al memoriei intermediare (zonei tampon) și toate blocurile murdare (modificate) din zona tampon sunt trimise la server. Chiar dacă nu este perfect, aceste ajustări fac sistemul foarte utilizabil în marea majoritate a situațiilor practice.

Capitolul 5

Sistemul FAT

Fat e un sistem dezvoltat de Microsoft pentru MS DOS fiind primul sistem de fisiere al versiunilor de Microsoft Windows, pana la si incluzand Microsoft Vista. FAT-ul a fost standardizat ca ECMA-107 si iso/iec 9293 aplicandu-se la discuri floppy si discuri optice. Sistemul FAT nu este foarte complicat si este suportat din punct de vedere virtual de toate sistemele de operare pentru uz personal. Acest lucru il face un format ideal pentru floppy disc si carduri de memorie.

Cele mai commune implementari au un foarte mare neajuns, in asa fel incat cand fisierele sunt sterse si altele noi sunt create, fragmentele de directoare tind sa devina imprastiate , in acest fel citirea si scrierea devenind un proces incet. Defragmentarea e o solutie pentru a acest lucru, dar deseori este un process de durata si trebuie efectuat in mod regulat pentru a pastra fisierele sistemului FAT curate.

Sistemul FAT a fost creat pentru discurile Microsot Standalone Disc Basic. In august 1980, Tim Paterson a incorporat Fat in sistemul de operare 86-DOS pentru s-100 8086 CPU. Sistemul de fisiere era diferenta principala dintre 86-dos si predecesori.

Numele vine din folosirea unei tablete care centralizeaza informatia a caror zone apartin fisierelor, care sunt libere sau posibil neutilizabile si unde fiecare fisier este stocat pe disc. Pentru limitarea marimii acestei tablete si a numarului de biti necesari in intrarile directoarelor, spatial discului este alocat pentru fisiere in grupuri numite cluster. Numarul maxim posibil de clustere a crescut de-a lungul timpului in mod dramatic, si numarul de biti care este necesar pentru pentru a identifica un cluster e folosit pentru a denumi versiuni successive majore ale formatului. Standardul FAT a fost de asemenea extins si pe alte cai, in timp ce pastreaza in mod invers compatibilitatea cu software existente.

5.1 FAT 12

Versiunea initiala a FAT este acum cunoscuta ca FAT 12. construita ca un system de fisiere pentru floppy disk, si avea urmatoarele neajunsuri: adresele clusterelor aveau doar 12 biti, care nu limita doar numarul clusterelor pana la 4078, dar facea manipularea FAT putin mai inselatoare pentru un PC de 8 biti

sau 16 biti (registrii). Spatiul discului a fost stocat in sectoare de 16 biti care limiteaza marimea la 32 MB. FAT 12 era folosit de diferiti manufacturieri cu diferite formate fizice, dar un floppy disc tipic al momentului era 5.25-inch, 40 tracks cu 8 sectoare pe track, rezultand intr-o capacitate de numai 160 KB pentru ambele arii ale sistemului si fisierelor.

Prin conventie toate structurile de control au fost organizate sa se incadreze pe prima linie, evitandu-se miscarile importante in timp real si operatiunile de scriere., cu toate astea, aceasta varia depinzand de modul de fabricatie si formatul fizic al discului. De cand DOS ul a fost introdus pe FAT 12, nu avea support ierarhic al directoarelor; numarul maxim al fisierelor era limitat la cateva zeci.

O limitare care nu a fost adresata pana mai tarziu, a fost ca orice sector nefunctional in zona controlului structurii, track 0, ar putea preveni ca discheta sa fie folosita. Instrumentul de formatare DOS a respins complet acest tip de dischete. Sectoarelor nefunctionale li s-a permis numai accesul in zona fisierelor, unde au facut intregul sustinator de cluster deasemenea neutilizabil.

5.2 FAT 16

In 1984 IBM a lansat pe piata de PC AT, care avea hard disk de 20 MB. Microsoft a introdus MS-DOS 3.0 in paralel. Adresele clusterelor erau marite la 16 biti, permitand un numar mai mare de clustere (pana la 65,517) si in cele din urma cu marimea fisierelor mai mare. Oricum, numarul maxim posibil de sectoare si marimea (partitie, decat pe disc) maxima de 32 MB nu a fost schimbata. Cu toate acestea, deja numitul FAT 16, acest format nu era ceea ce azi se intelege sub acest nume. Un disk e 20 MB format sub MS-DOS 3.0 nu era accesibil in vechitului MS-DOS 2.0. bineinteles MS-Dos 3.0 poate inca accesa stilul

Partitii extinse si drive-uri logice

Pt a imbunatatii structura sistemului de fisiere Fat ,a fost dezvoltata o cercetare in paralel permitand o crestere a posibilitatilor maxime a spatiului de stocare FAT prin folosirea mai multor partitii FAT .Partitiile initiale trebuiau sa se foloseasca numai pentru impartirea discului intre sistemele de operare, DOS si Xenix pe vremea aceea, asa ca DOS a fost pregatit numai sa faca fata numai UNEI partitii FAT . Nu era posibil sa se creeze a partitie DOS multipla folosind

instrumentele Dos iar alte utilitare avertizau ca o astfel de operatie nu era compatibila cu DOS. Numai permitand mai multe partitii ideentice in DOS poate sa se ajunga la mai multe probleme: poate C: sa fie prima partitie FAT, pentru a fi mai simplu sau mai bine partittia marcata ca fiind activa in tabela partitilor, ca mai multe variante de DOS sa co-existe? Si care partitie trebuie sa fi C:daca sistemul a boot-at dupa dischete ?

Pentru a permite folosirea mai multor partitii FAT intr-un mod compatibil , o noua partitie a fost introdusa (in MS-DOS 3.2 ianuarie 1986), se numea partitia extinsa; care era de fapt doar un container pentru partitii aditionale numit "logical drives". Initial numai 1 drive logic era posibil,permitand sa se foloseasca pana la 64 MB din discurile hardului. In MS-DOS 3.3 (din august 1987) aceasta limita a fost marita la 24 de drive..Au fost intotdeauna stocate pe disk ca un rand pe blocuri separate intr-o singura cutie; despre aceste blocuri se spune adesea ca sunt legate impreuna de linkurile din sectoarele EBR.Numai o partie extinsa a fost permisa. Drive-urile logice nu erau bootabile si partitile extinse putea numai sa fie create dupa prima partie FAT, care elimina toata ambiguitatea, dar de asemenea si posibilitatea de a boot-a mai multe versiuni DOS din acelasi harddisk.

Un advers folositor a schemei partitiei extinse a fost sa creasca semnificativ numarul maxim de partitii posibile pe hard discul PC-ului peste 4 care putea sa fie descrise un singur MBR.

Inaintea introducerii partitilor extinse niste hard discuri(care la timpul acela erau optionale, deoarece IDE exista inca) puteau face hard discuri mai mari care puteau aparea ca doua discuri separate.

Ultimul FAT16

In sfarsit in noiembrie 1987 Compaq DOS 3.31 a introdus ce noi azim formatul FAT16 cu o extindere pe sectorul discului de 16 biti pana la 32 biti.Rezultatul s-a numit initial DOS 3.31 Large File System(sistemul cu fisiere mari). Cu toate ca pe disc schimbarile erau aproape minore ,intregul cod de disc DOS a fost schimbat sa foloseasca 32 de biti pe sectorul de numere, o sarcina complicata prin faptul ca era scris pe 16 biti in limbajul de programare.

In 1988 imbunattirea a devenit mai generala fiind disponibila si pe MS-DOS 4.0 si OS/2 1.1.Limita pe marimea partitiei era acum data pe 8 biti care a avea valoare maxima pe 64. Cu hardiscul obisnuit de 512 bitipe sector, care dadea 32k de clusters, a definit limita de 2 gigabiti pentru partitia de fat16. Pe media magneto-optica care poate avea 1 sau 2 kilobiti pe sector limita era proportional mai mare.

Mult mai tarziu WINDOWS NT a marit marimea maxim a clusterului pana la 64 KB . Oricum formatul final nu era compatibil cu nici un FAT implementat in timpul acela si genera massive fragmentary interne.Windows 98 suporta de asemea cititul si scrisul pe aceasta varianta dar utilitatile discului nu mergeau pe el.

Numarul directoarelor de intrare disponibile in timpul formatatului si stocate intr-un camp de 16 biti a obtinut o limita absoluta de 32767 intrari(32767 un multiplu de 32).Din motive istorice FAT12 si FAT16 foloseau 512 directoare de intrare de root pe mediile non-floppy si celelalte marimi erau necompatibile cu niste software sau dispozitive.

5.2 FAT32

FAT32 este o modalitate de formatare care permite o marime a fisierului de maxim 4 GB .Fisiere mai mari au nevoie de alt tip de formatare cum ar fi HFS+ sau NTFS.

Pentru a nu depasii volumul maxim de marime a lui FAT16, in timp ce inca ii permite DOS-ului in timp real sa faca fata formatului faca sa reduca neaparat memoria conventionala, Microsoft a decis sa implementeze a noua generatie de FAT32 cu numarul de clustere tinuta intr-un camp de 32 de biti, din care 28 de biti sunt folsiti pentru a tine numarul de clustere pentru un maxim de aproape 250 de milioane de clustere.

Aceasta poate permite ca marimea discului sa fie de 8 terabiti cu 32k de clustere, dar sectorul de boot foloseste un camp de 32 de biti pentru numaratoarea sectorului,limitand marimea la 2TB pe hrd disk cu pana la 512 biti pe sector.

Pe Windows 95/98 , datorita versiunii Microsoft ScanDisk care includea aceste sisteme de operare fiind o aplicatie 16-bit,structuraFAT ne permitea cresterea pana la 4 milioane de clustere, determinand o limita a volumului de 127.53 gibabiti. O limitare a Fdisc-ului in versiunile orginale Windows 98/98 SE cauza o raportare incorecta a volumului diskurilor de peste 64 GB. O corecta versiune este disponibila la Microsoft. Aceste limitary nu se aplica pe Windows 2000/XP cu exceptia Setup-ului in care este o limita de 32 de GB.Windows ME suporta sistemele FAT32 fara nici o limitare.

FAT32 a fost introdus cu Windows 95 OSR2, cu toate ca reformatarea necesita folosirea acesteia, si DriveSpace 3 nu a suportat-o niciodata. Windows 98 a introdus o utilitate de convertire a hard-disk-urilor din FAT16 in FAT32 fara nici o pierdere de date. In linia NT, suportul pt FAT32 a aparut in Windows 2000. Un driver gratuit FAT32 pentru Windows NT4.0 a fos pentru Winternals,

o companie achizitionata mai tarziu de catre Microsoft. De la achizitia driverului nu mai este oficial valabila.

Windows 2000 si XP pot citi sau scrie fisiere de sistem FAT32 de orice marime, dar programul de formatare inclus in Windows 2000 sau mai nou nu pot crea decat sisteme FAT32 de 32 de GB sau mai putini. Aceasta limitare a fost impusa prin design si datorita faptului ca multiplele sarcini deveneau incete si ineficiente pentru un system foarte mare ca FAT32. Aceasta limitare poate fi evitata cu utilitati de formatare third-party.

Maximul posibil a marimi fisierelor FAT32 este 4GB minus 1 byte. Captura video si editarea aplicatiilor si alte software-uri pot cu usurinta depasii aceasta limita. Pana la mid-2006, aceia care foloseau sisteme boot duale sau cine muta dirvere externe de date intre computere cu diferite sisteme de operare nu au avut de ales decat sa ramana la FAT32. De atunci, suportul complet pentru NTFS a devenit disponibil in Linux si multe alte sisteme de operare, instaland FUSE library impreuna cu aplicatia NTFS-3G. Schimbul de date este de asemenea posibil intre Windows si Linux folosind fisiere system ale Linux-native, ext2 sau ext3, prin intermediul folosirii driverelor externe pentru Windows cum ar fi ext2 IFS, cu toate acestea Windows nu poate boota din partiitiile ext2 sau ext3.

Capitolul 6

Sistemul NTFS

NTFS este sistemul standard de fisiere Windows NT, incluzand versiunile ulterioare WIndows 2000, Windows XP, Windows Server 2003, Windows Server 2008 si Windows Vista.

NTFS inlocuieste sistemul de fisiere FAT ca sistem preferat pentru Windows. NTFS prezinta avantaje fata de FAT si HPFS (High Performance File System) cum ar fi suportul imbunatatit pentru metadata, folosirea structurilor avansate pentru a imbunatati performatia, fiabilitatea si utilizarea spatiului pe disc si alte extensii cum ar fi cele de securitate: access control lists si file system journaling. Specificatiile exacte ale sistemului de securitate sunt confidentiale desi sistemul poate fi licentiat de Microsoft prin programul de licentiere a proprietatii intelectuale.

Versiuni

NTFS are 5 versiuni:

- v1.0
- v1.1
- v1.2 intalnita in NT 3.51 si NT 4
- v3.0 intalnita in WIndows 2000
- v3.1 intalnita in Windows XP, Windows Server 2003 si Windows Vista

Versiunile NTFS instalate pe sistemele de operare amintite mai sus sunt uneori denumite dupa sistemul de operare cu care sunt distribuite. Versiunea v1.2 este cunoscuta si ca 4.0 iar 3.1 este uneori cunoscuta(incorect) ca 5.0 si 5.1. Fiecare noua versiune adauga imbunatatiri ca de exemplu: Windows 2000 a introdus diviziuni (disk quotas) iar Windows Vista a introdus Transactional NTFS, legaturi simbolice NTFS si functionalitatea "self-healing".

Caracteristici

In NTFS toate datele - nume de fisiere, data creatiei, permisiile de acces si continut - sunt stocate ca meta-date (date despre date). Aceasta apropiere abstracta a permis adaugarea unor imbunatatiri ulterioare in timpul dezvoltarii

Windows NT - un exemplu interesant este adaugarea campurilor de indexare folosite de utilitarul "Active dosary".

NTFS permite orice secventa de valori pe 16 biti pentru codarea numelor (nume de fisiere, fluxuri de date, index-uri) ceea ce inseamna ca sistemul suporta codarea UTF-16 dar nu verifica daca o secventa respecta sau nu UTF-16 (permite orice secventa de valori scurte, nerestrictionate de standardul UNICODE).

Sistemul foloseste structuri de tipul "B+ trees" pentru a indexa datele de sistem. Desi aceasta structura presupune o implementare complexa ea permite o cautare mai rapida a fisierelor in majoritatea cazurilor. Un fisier tip "jurnal" este folosit pentru a garanta integritatea sistemului de fisiere - dar nu si integritatea continutului fisierelor. Sistemele care folosesc NTFS sunt cunoscute pentru fiabilitatea imbunatatita fata de sistemele de fisiere FAT.

Tabelul Master File (MFT) contine meta-date despre aproximativ fiecare fisier, dosar sau meta-fisier pe un volum NTFS - incluzand nume de fisier, locatii, dimensiuni si permisii. Structura sa suporta algoritmi care minimizeaza fragmetarea discului. O inregistrare de tip dosar consta intr-un nume de fisier si o cheie de identificare constand in numarul inregistrarii reprezentante in MFT. Identificatorul unui fisier contine deasemenea si un contor de reutilizare pentru a detecta vechimea acestuia.

Metafisiere

In Windows NT 5.0 (Windows 2000) si versiunile mai vechi comanda utilitara DIR ar lista informatii despre meta-fisier daca ar fi cerute explicit: dir /a c:\\$MFT, altfel acestea s-ar comporta ca fisier invizibile, fiind mai bine ascunse decat fiserele "hidden" sau "system". Aceast comportament a fost eliminat incepand cu Windows NT 5.1 (XP).

Fid	Nume fisier	Scop
0	\$MFT	Describe toate fisierele de pe acest volum
1	\$MFTMirr	Dubleaza primele inregistrari foarte importante ale \$MFT, de obicei 8 (4 kb)
2	\$LogFile	Jurnal al schimbarilor efectuate asupra fisierelor de sistem
3	\$Volume	Contine identificatorul volumului, eticheta sa, versiunea fisierului de system si modificatoare ale

		volumului: montat, verificare chkdsk ceruta, redimensionare \$LogFile ceruta, montat pe NT4, numarul de serie dupa actualizare, actualizare a structurii ceruta. (numarul de serie al volumului se afla in \$Boot)
4	\$AttrDef	Describe tipurile de inregistrari folosite in intrarile \$MFT
5	.	dosar radacina
6	\$Bitmap	harta alocarii gruparilor volumului
7	\$Boot	contains a <u>Volume boot record</u> including level 2 bootloader, a <u>BIOS parameter block</u> including <u>volume serial number</u> . This file is always placed at the beginning of the volume. It also contains the clusters where \$MFT and \$MFTMirr begin.
8	\$BadClus	Acest fisier contine toate gruparile volumului marcate ca deteriorate ("bad sectors")
9	\$Secure	Baza de date a listei de control a accesului
10	\$UpCase	Se speculeaza a fi o harta a folosirii literelor mari
11	\$Extend	Un dosar a fisierelor de system continand fisierele 24, 25, 26
12..23	Rezervat pentru extensii \$MFT	
24	\$Extend\$Quota	administrarea a spatelui
25	\$Extend\$ObjId	Identificator contextual de securitate
26	\$Extend\$Repars e	Baza de date a legaturilor simbolice
27..	pagefile.sys	Inceputul intrarilor ce contin fisiere

Intrarile pana in 24 sunt tratate special de NTFS si sunt nume de fisiere sau continuturi dificile: unelte special create pentru scop.

Fisiere permanente vs. fisiere temporare

Pentru a optimiza depozitarea cazului usual de fisiere de date mici, NTFS tinde sa plaseze continutul fisierului in tabelul fisierului master - daca se potriveste, in loc sa foloseasca acelasi spatiu MFT pentru a lista gramada ce contine datele. Prima precizata este numita "data permanenta" de programatorii juridici . Data actuala care se potriveste este foarte dependenta de caracteristicile fisierului, dar 700 caractere sunt obisnuite pentru nume de fisiere scurte fara ACL-uri.

Interoperabilitate

Detalii despre internii implementarii sunt inchisi, fapt ce este dificil pentru vanzatorii partii terte sa furnizeze ustensile pentru a manevra NTFS.

Linux

Citirea/scrierea totala si sigura este furnizata de driver-ul NTFS-3G. Este inclus in majoritatea distributiilor de Linux.

Alte solutii majoritar read-only si out-datate exista deasemenea:

- esenta Linux 2.2:Partitiile NTFS pot fi citite de esenta de la versiunea 2.2.0.
- esenta Linux 2.6: contine un driver scris de Anton Altaparmakov (Universitatea Cambridge) si Richard Russon. Suporta citire de fisiere , scriere peste si modificarea marimii,in cateva cazuri.
- NTFSMount: Un driver cu suport limitat de scriere/citire al fisierului si al dosarului este disponibil folosind NTFSMount.
- NTFS pentru Linux : Un driver comercial cu suport total de citire/scriere disponibil de la Paragon.
- NTFS captiv: Un driver "de ambalare" care foloseste propriul driver the Windows, ntfs.sys.

Este notabil faptul ca toate cele 3 drivere de spatiu ale utilizatorului, in principal NTFSMount, NTFS-3G si NTFS captiv, sunt construite pe sistemul de fisiere in spatiul pentru user (FUSE), o esenta a modulului de Linux insarcinat cu arcuirea spatiului de user si al codului de esenta pentru a salva si a recupera datele.

Aproape toate driverele listate mai sus (cu exceptia Paragonului NTFS pentru Linux) sunt surse deschise (GPL). Datorita complexitatii structurilor interne NTFS, atat driverul esenta built-in 2.6.14 cat si driverele FUSE nu permit schimbari ale volumului care nu este considerat sigur, pentru a evita coruptia.

Windows

In timp ce versiunile diferite de NTFS au un mare grad atat de compatibilitate de inaintare cat si de inapoiere, exista consideratii tehnice pentru a monta volume noi de NTFS in variante mai vechi de Windows. Acest lucru afecteaza butarea dubla, cat si hard disk-urile portabile externe.

De exemplu, "Versiunile anterioare" (Copia Umbra a Volumului) sunt pierdute deoarece OS-ul mai vechi nu stie cum sa pastreze noile caracteristici ale datelor update.

Altele

Driverul NTFS-3G asigura acces citire/scriere sigur si total pe FreeBSD, Mac OS X, NetBSD, Linux, Haiku si FreeDOS.

Versiunile FreeBSD, eComStation si Mac OS X 10.3 si urmatoarele ofera suport NTFS read-only (exista un driver NTFS beta care permite scrierea/stergerea pentru eComStation, dar in general acest procedeu este considerat nesigur). Un instrument gratis tert pentru BeOS, care e bazat pe NTFS-3G, permite scrierea si citirea totala a NTFS.

Exista deasemenea un driver scriere/citire comercial pentru DOS numit "NTFS4DOS".

Compatibilitatea cu FAT

Momentan Microsoft furnizeaza un instrument (convert.exe) pentru a converti HPFS (doar pe Windows NT 3), FAT16 si, pe Windows 2000 si variante mai curente, FAT32 in NTFS, dar nu invers. Instrumente tert variate sunt toate capabile sa redimensioneze partiile NTFS. Microsoft a adaugat abilitatea de a restrange sau a extinde o partiile cu Windows Vista .

Din motive istorice, versiunile de Windows care nu suporta NTFS ,toate tin timpul intern ca timp zonal local, si astfel asta fac toate celelalte sisteme de fisiere in afara de NTFS care sunt suportate de versiuni curente de Windows. Totusi, Windows NT si descendenti sai inregistreaza timp de tip UTC si fac conversiile cele mai potrivite pentru scopuri de display. Astfel, timpii NTFS sunt in UTC. Asta inseamna ca atunc cand fisierele sunt copiate sau mutate intre partiile NTFS si non-NTFS, OS-ul trebuie sa converteasca timpii

pe loc. Dar daca cateva fisiere sunt mutate cand timpul de salvare al zilei (DST) este activ, si alte fisiere sunt mutate cand timpul standard este activ, pot aparea cateva ambiguitati in conversii. Ca rezultat, in mod special la scurt timp dupa una din zile in care timpul zonal local se schimba, userii pot observa ca unele fisiere au timpi incorecti cu o ora. Datorita diferentelor in implementarea DST intre emisferele nordice si sudice, asta poate rezulta intr-o potentiala eroare de timp mai sus de 4 ore in oricare 12 luni.

Article I.Caracteristici

NTFS v3.0, a 3-a versiune a NTFS care este introdusa, include diverse noi caracteristici deasupra predecesorilor ei: partea folosirii discului, suportul fisierului din loc in loc, cautarea linkului si criptarea nivelului de fisier, cunoscut ca si Fisierul de Sistem Criptat (EFS)

- Fluxul alternat de date (ADS)

Fluxul alternat de date permite fisierelor sa fie asociate cu mai mult de un flux de date. De exemplu, un fisier ca text.txt poate avea un ADS cu numele de text.txt : secret.txt (a formei filename:ads) care poate fi accesata cunoscand numele de ADS sau navigand programe in dosare specializate. Fluxurile alternate nu sunt detectabile in marimea originala a fisierului, dar sunt pierdute cand fisierul original (ex: text.txt) este sters cu un InlocuireFisier sau cu un apel la InlocuireaFisier (sau un apel care foloseste acele apeluri) , sau cand fisierul este copiat sau mutat intr-o partitie care nu suporta ADS (ex: o partitie FAT, un floppy, sau o partajare de retea). In timp ce ADS este o caracteristica folositoare, poate usor ocupa spatiul hard discului daca a fost uitat sau n-a fost detectat.

- Cote

Cotele discului au fost introduse in NTFS v3. Ele permit administratorului computerului care ruleaza o versiune a Windowsului care suporta NTFS sa seteze o limita a spatiului de disc pe care userii o pot utiliza Permite de asemenea administratorilor sa tina urma a cat spatiu al discului foloseste fiecare utilizator. Un administrator poate specifica un nivel concret de spatiu pe care un user il poate folosi inainte sa ajunga la limita de spatiu. Cotele discului nu tin cont de compresia fisierului transparent NTFS , daca acesta este ativat.

- **Fisierele risipite**

Fisierele risipite sunt fisiere ce contin seturi de date risipite, date in principal umplute cu zerouri. Multe aplicatii stiintifice pot genera mari seturi de date risipite. Din aceasta cauza, Microsoft a implementat suport pentru fisierele risipite permitand o aplicatie sa specifice regiuni de date goale(zerouri). O aplicatie ce citeste un fisier risipit il citeste in mod normal cu systemul de fisier , calculand ce date trebuie returnate ,bazandu-se pe offsetul fisierului. O data cu fisierele compresate, marimea initiala a fisiereleor risipite nu sunt luate in considerare cand se determina limitele cotei.

- **Puncte de montare pentru volumelor**

Asemanator punctelor de montare Unix, unde radacina altui fisier de sistem este atasata unui dosar. In NTFS acesta permite fisiere aditionale sa fie montate fara a necesita o litera separata (ex C: sau D:) pentru fiecare.

- **Jonctiuni ale dosarelor**

Asemanator punctelor de montare pt volume, oricum jonctiunile dosarelor leaga alte dosare din sistemul de fisiere in loc de alte volume. De exemplu, dosarul C:\exampledir cu un atribut jonctiune dosar care contine o legatura catre D:\linkaddir va trimite automat spre dosarul D:\linkaddir cand este accesat de o aplicatie in mod utilizator. Acesta functionalitate este similara coceptual cu o legatura simbolica spre un dosar in UNIX exceptand faptul ca tinta in NTFS trebuie sa fie intotdeauna alt dosar. (Fisierele Unix permit in general legaturi simbolice spre orice tip de fisiere.)

- **Legaturi de tip „Hard”**

Incluse initial pentru a suporta subsistemul POSIX in Windows NT, legaturile tip „Hard” sunt similare cu jonctiunile dosar dar folosite pentru fisiere. Acestea pot fi aplicate numai pentru fisierele care apartin aceluiasi volum.

- **Managmentul Ierarhic al stocarii (HSM)**

Managmentul ierarhic al stocarii reprezinta o cale de a transfera fisiere care nu sunt folosite pentru o perioada de time spre medii mai ieftine.

- **Stocare structurata nativ (NSS)**

NSS a fost o tehnologie ActiveX de stocare a documentelor care de la inceput a fost oprita de Microsoft. Aceasta permitea documentelor ActiveX sa fie stocate in acelasi format multi-flux pe care ActiveX in folosea intern. Un filtru fisier de sistem NSS a fost incarcat si folosit pentru a procesa fluxurile multiple independent de aplicatie si cand fisierul a fost copiat catre un disk formatat Non-NTFS va permite transferul mai multor fluxuri de date intr-unul singur.

- **Compresia fisiereilor**

NTFS poate compresa fisiere folosind o varianta a algoritmului LZ77 – folosit si de celebrul utilitar WinZip. Desi accesul citire-scriere catre fisiere compresate este transparent, Microsoft recomanda compresia pe sisteme server deoarece aceasta incarca in mod considerabil procesorul.

- **Stocarea unei singure instante (SIS)**

Atunci cand mai multe dosare au unele fisiere identice SIS permite ca stocarea unui singur fisier si creeaza referinte catre acel fisier. SIS consta intr-un fisier de sistem tip filtru care se ocupa de copierea, modificarea si concatenarea fisiereilor, si un serviciu al spatilui ultizatorului care cauta fisierele identice.

- **Criptarea fisiereilor de sistem (EFS)**

EFS ofera modalitati puternice si transparente de criptare a oricarui fisier sau dosar al unui volum NTFS. EFS lucreaza in conjunctura cu serviciile EFS, CryotoAPI al Microsoft si libraria EFS File System Run-Time Library.

- **Legaturi simbolice**

Legaturile simbolice au fost introduse in Widows Vista. Legaturile simbolice (sau legaturi „Soft”) sunt rezolvate de catre client. Deci cand o

legatura simbolica este partajata, tinta este subiectul restrictiilor de acces ale clientului, nu ale serverului.

- **Transactional NTFS**

Ca si Windows Vista, aplicatiile pot folosi Transactional NTFS pentru a grupa schimbarile fisierelor impreuna intr-o sesiune. Sesiunea va garanta ca toate schimbarile se aplica, sau nici una dintre ele, si ca aplicatiile din afara sesiunii nu vor vedea schimbarile pana exact in secunda cand ele vor fi gata.

Article II.Limitari

Sistemul NTFS are urmatoarele limitari:

- **Nume de fisiere rezervate**

Desi sistemul de fisiere suporta cai de pana la 32,000 caractere Unicode cu fiecare component ai caii (dosar sau fisier) de pana la 255 caractere, unele nume sunt neutilizabile deoarece NTFS isi stocheaza meta-data in fisiere normale (desigur ascunse si in cea mai mare parte inaccesibile), deci utilizatorii nu pot folosi aceste nume. Aceste fisiere se afla toate in dosarul radacina al volumului (si sunt rezervate doar pentru acel dosar). Aceste nume sunt: \$MFT, \$MFTMirr, \$LogFile, \$Volume, \$AttrDef, . (punct), \$Bitmap, \$Boot, \$BadClus, \$Secure, \$Upcase, si \$Extend; . (punct) si \$Extend sunt ambele dosare, celelalte sunt fisiere.

- **Maximum Volume Size**

Teoretic, dimensiunea maxima a unui volum NTFS este $2^{64}-1$ grupari dar practic aceasta dimensiune a fost implementata in Windows XP ca $2^{32}-1$ grupari. De exemplu, folosind grupari de 64kb, dimensiunea maxima a unui volum NTFS este de 256 TB minus 64 KB. Folosind marimea implicita a unei grupari de 4 kb dimeansinea maxima a unui volum NTFS este de 16 Tb minus 4 Kb. Deoarece tabela de partitii de pe discul master boot record (MBR) suporta partitii de pana la 2 Tb, volumele dinamice sau de timp GPT (Globally Unique Identifier Partition Table) trebuiesc folosite pentru a crea partitii bootabile NTFS de peste 2 TB

- **Marimea maxima a unui fisier**

Teoretic aceasta este de 16 Eb minus 1 Kb ($2^{64} - 2^{10}$ bytes). Implementare: 16 Tb minus 64 kb ($2^{44} - 2^{16}$ bytes).

- **Fluxuri alternative de date**

Sistemul Windows poate – sau nu poate – folosi fluxuri alternative de date. Depinzand de sistemul de operare, fisiere de sistem utilitare si rezervate, un transfer de fisiere poate extrage fluxuri de date neobservat. O modalitate sigura de a copia sau muta fisiere este folosind functiile de sistem BackupRead si BackupWrite, care permit programelor sa enumere fluxurile, sa verifice daca fiecare trebuie scris pe volumul destinatie si sa sara fluxurile amenintatoare.

- **Lungimea maxima a unui cai**

O cale absoluta poate contine pana la 32767 caractere, iar una relativa pana la 255 caractere.

Capitolul 7

Sistemul HFS

Hierarchical File System (HFS) e un system de fisiere dezvoltat de Apple Inc. pentru a fi folosit in computerele ce ruleaza Mac OS. Conceput initial pentru folosirea pe hard diskuri sau pe floppy diskuri poate fi gasit si pe medii read-only cum sunt CD Rom-urile.

HFS a fost introdus de Apple in Septembrie 1985 ca inlocuitor pentru Macintosh File System (MFS), primul sistem de fisiere introdus un an mai devreme cu calculatorul Macintosh. Dezvoltat de Patrick Dirks si Bill Bruffey HFS impartaseste cateva caracteristici cu MFS, care nu erau disponibile la vremea respectiva in alte sisteme de fisiere cum era de exemplu FAT.

Fisierele puteau avea mai multe forkuri de informatie (de obicei de date si de resurse), cu ajutorul carora codul unui program putea fi stocat separat de resursele programului cum ar fi icoanele. Fisierele erau apelate dupa un ID unic decat doua nume, iar numele unui fisier putea contine 255 caractere.

Din cauza ca MFS era optimizat pentru a fi folosit pe unitati media mici si incete, anume floppy disk-uri HFS a fost introdus pentru a rezolva problemele de performanta aparute cu introducerea unitatilor mai mari anume a hard disk-urilor. Principala problema era timpul necesar pentru aratarea continutului unui folder. La MFS toate informatiile despre foldere si fisiere era stocate intr-un singur fisier in care sistemul trebuia sa caute pentru a construi o lista cu fisierele dintr-un anumit director. Acest procedeu dadea rezultate pe un sistem de cateva sute de kilobytes de stocare sau cateva sute de fisiere dar pe masura ce capacitatea de stocare a crescut la megabytes performantele au scazut vizibil.

Solutia a fost inlocuirea structurii de directoare a MFS cu una potrivita pentru sisteme de fisiere mai mari. HFS a inlocuit structura tabelara cu Catalog File care foloseste o structura de tip arbore B*, o structura in care se pot face cautari foarte rapide idiferent de dimensiuni. HFS are refacute diferite structuri pentru a putea stoca numere mai mari, integerii pe 16 biti fiind inlocuiti in totalitate de cei pe 32 de biti. Ciudat singurul loc unde aceasta marire nu a avut loc este directorul de fisiere, unde HFS este limitat la maxim 64k fisiere.

Deși HFS este un sistem de fisiere de sine statator, este bine documentat și sunt multe modalități de a accesa diskurile formate HFS din majoritatea sistemelor de operare moderne.

În 1998, Apple introduce HFS Plus pentru a adresa neajunsuri legate de alocarea ineficientă a spațiului pe disk în HFS și pentru a adăuga alte îmbunătățiri. HFS este compatibil cu versiunile curente de Mac OS, dar începând cu Mac OS X un volum HFS nu poate fi folosit pentru bootare.

Design

Sistemul de fisiere ierarhic (HFS) împarte un volum în blocuri logice de 512 bytes. Aceste blocuri sunt grupate în blocuri de alocare care pot conține unul sau mai multe blocuri logice, în funcție de dimensiunile volumului. HFS folosește o valoare pe 16 biți pentru a apela blocurile de alocare, limitând numărul acestora la 65 536.

Sunt 5 structuri care alcătuiesc un volum HFS:

1. Blocurile logice 0 și 1 ale volumului sunt blocuri de boot care conțin informațiile pentru start-up-ul sistemului.
2. Blocul logic 2 conține Master Directory Block (MDB). În acesta sunt definite o mare parte de date referitoare la volum în sine, de exemplu data și ora când volumul a fost creat, locația altor structuri ale volumului sau mărimea structurilor logice cum ar fi blocurile de alocare. Mai există un duplicat al MDB numit Alternative Master Directory Block localizat la capatul celălalt al volumului. Acesta este folosit în mare parte de utilitățile de disc.
3. Blocul logic 3 este blocul de început pentru Volume Bitmap, care ține cont care blocuri de alocare sunt folosite și care sunt libere. Fiecare bloc de alocare din volum este reprezentat printr-un bit în Volume Bitmap. Dacă bitul este setat blocul este folosit, dacă bitul este liber blocul este de asemenea liber.

4. Extent Overflow File e un arbore B* care contine spatiu suplimentar unde se inregistreaza ce blocuri de alocare sunt folosite pentru care fisiere, odata ce spatiul din Catalog File este folosit. Versiunile mai noi au adaugat abilitatea ca in Extent Overflow File sa poate stoca informatii despre blocuri bad, pentru a impiedica sistemul sa aloce un bloc bad unui fisier.
5. Catalog File este un alt arbore B* care contine registri despre toate fisierle si directoarele stocate in volum. Stocheaza 4 tipuri de informatii. Fiecare fisier contine un File Thread Record si un File Record in timp ce fiecare director contine un Directory Thread Record si un Directory Record. Fisierle si directoarele din Catalog File sunt apelate dupa un unic Catalog Node ID sau CNID.

Probleme specifice HFS

Catalog File care stocheaza toate informatiile despre fisiere si directoare intr-o singura structura de date sufera de probleme legate de performanta cand sistemul adminte multitasking, pentru ca un singur program poate scrie/citi din aceasta structura intr-un moment, astfel ingreunand sistemul.

Mai apare si o problema de siguranta pentru ca o eroare in acest fisier poate distruge intregul sistem. Acest lucru vine in contrast cu alte sisteme de fisiere care tin informatiile legate de fisiere si directoare in structuri separate (cum ar fi FAT sau Unix File System) unde stocarea structurilor pe tot discul inseamna ca stricarea unui singur director nu e fatala si datele pot fi recuperate folosind datele retinute in portiunile stricate.

Capitolul 8

Tehnologii RAID de stocare si securitate a datelor

Definite

RAID (Redundant Array of Inexpensive (or Independent) Drives (or Disks)) este o tehnica implementata hardware sau software folosind doua sau mai multe unitati HardDisk pentru a crea un singur Disk logic cu performante imbunatatite.

Istorie

Prima decriere a tehnicilor RAID a fost patentata in anul 1987 cu titlul "System for recovering data stored in failed memory unit." de catre Norman Ken Ouchi (IBM). Revendicarile facute in acest patent descriu ce vom numi in continuarea acestei prezentari termenul de RAID 5. De asemenea mentioneaza tehnici de „disk mirroring” si „disk duplexing” (RAID 1) si protectie cu paritate dedicata.

Abia in anul 1987 termenul RAID a fost definit de catre David A. Patterson, Garth A. Gibson si Randy Katz (Berkeley University of California) prin lucrarea "A Case for Redundant Arrays of Inexpensive drives (RAID)". Cu toate ca pana in prezent implementarile tehnicilor RAID difera substantial fata de cele din anul 1987 numele numerotate ale tehnicilor au ramas (RAID 1,2,3,4...).

Avantaje si Aplicabilitate

La baza oricarei inovatii tehnologice sta necesitatea. RAID a aparut si a fost dezvoltat ca urmare a incapacitatii unui simplu HDD de a face fata cerintelor tehnice din ce in ce mai mari. Crearea unor dispozitive cu viteza de acces si spatiul de stocare cat mai mari au pus in dificultate (sau imposibilitate) producatorii de HDD. De aceea tehnica RAID aduce urmatoarele avantaje:

Accesul in paralel a datelor ceea ce mareste gradul de acces la date decat pe un singur disc si in acelasi timp se mareste si spatiul de stocare. Pretul pentru o unitate (in cazul in care ar putea fi tehnologic realizabila) cu spatiu de stocare 10 TB cu o viteza de acces de 700MB pe secunda ar fi enorm.

Redundanta datelor pe mai multe discuri ofera protectie in cazul defectarilor. Spre exemplu, se pot folosi configuratii care, in cazul defectarii si inlocuirii unui HDD, datele sa fie recuperate fara erori.

Domeniul de folosire al acestuia este cel al serverelor, in care nevoia pentru viteza de acces si spatiul de stocare este foarte mare. Parasind domeniul enterprise se pot totusi intalni tehnici RAID si in statii de lucru, insa nivelul RAID folosit este mic (de obicei intalnit RAID 0).

Moduri de implementare a tehnicii RAID

Distributia de date intre multiplele unitati poate fi realizata atat pe cale software cat si pe cale hardware. Adicional exista si un mod hibrid, partial hardware, partial software.

Implementare software

Implementarile software sunt furnizate de majoritatea sistemelor de operare. Acest layer este pus deasupra driverelor discurilor si ofera un strat de abstractizare discurile logice si discurile fizice. In mod tipic RAID-ul software este limitat la RAID 0, 1 si 5. In sisteme de operare multi-thread (Linux, Mac OS X, Novell, Windows NT, 2000..) sistemul de operare poate face mai multe procese I/O (permite multiple scrieri si citiri). Aceasta capacitate permite, in plus, realizarea RAID 0/1 in implementare software cu toate ca nimeni nu foloseste aceasta practica. Motivul este cererea ridicata de putere de procesare pentru calcularea paritatii (datele redundante). Implementarile software necesita foarte putine resurse si timp de prelucrare mici. In anul 2007 calculatoarele personale, obisnuite, ofera puterea de calcul necesara crescand folosirea procesorului cu doar un procent in plus. Ele sunt eficiente ca si cost (pana intr-un punct) insa nu ofera acelasi grad de performanta ca si RAID-ul bazat pe hardware. Fiind software, el trebuie sa foloseasca resursele unui server gazda legat de spatiul de stocare, CPU-ul gazdei fiind o parte ocupat cu procesarea datelor. In cazul unei defectiuni a gazdei datele din spatiul de stocare nu pot fi accesate atat cat defectiunea persista. In hardware matricea de HDD-uri poate fi atasata usor altei gazde astfel timpul de recuperare se reduce simtitor. Totusi acest sistem este la fel de eficient ca si cel software in cazul in care nu exista un sistem de back-up al hostului pe care sa se faca trecerea rapid.

O alta facilitate adusa de implementarea software este ca permite crearea de matrici RAID si din partitiile ale unui HDD.

Implementare hardware

RAIDul hardware are ca cerinta minima un controler RAID. Intr-un calculator personal acesta poate fi un card PCI sau poate fi inclus in cip-set-ul placii de baza. In aplicatiile industriale controlerul se creeaza ca si unitate separata. HDD-urile pot fi IDE/ATA, SATA, SCSI, SSA, canal optic si combinatii intre acestea. Unitatea ce foloseste acest sistem de stocare poate fi legat direct la controler sau, mai des intalnit, printr-o retea SAN (Storage Area Network - o retea ce foloseste un protocol de comunicare asemanator cu SCSI ce, spre exemplu, se poate suprapune peste o retea Ethernet). Controlerul se ocupa de managementul discurilor si calculeaza datele redundante caracteristice nivelului RAID ales. Majoritatea acestor sisteme au implementat o memorie cache nevolatila, ce imbunatatesc performantele. Implementarea hardware prezinta performante garantate, fara utilizare a procesorului gazdei, controlerul prezentand sistemului de operare un simplu disc logic. De asemenea, sistemele industriale prezinta suport pentru "hot swapping" (schimbarea HDD-urilor defecte fara a fi necesara oprirea sistemului).

Implementare hibrida

Echipamentele RAID hibride au aparut odata cu introducerea controlerelor RAID ieftine, implementate in controlere HDD si ca extensii de BIOS reprezentate de diverse sisteme. Din nefericire aceste controlere fac toate calculele necesare in regim software, ele aduna majoritatea dezavantajelor implementarii software si hardware. Ca si implementariile HW, ele sunt proprietate unui fabricant de controlere RAID si nu se pot face combinarii intre mai multe controlere. Avantajele constau in abilitatea de a boota de pe discul logic si integrarea mai stransa cu driverul device-ului ce ofera o manipulare mai buna a erorilor.

Nivele RAID

Nivele standard

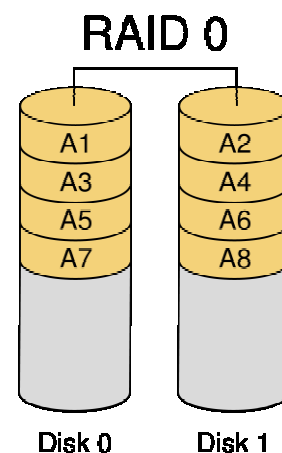
Nivelele RAID standard sunt un set de configuratii de baza ce se gasesc in orice implementare ce implica tehnicile : striping (partajarea unui bloc de date in mai multe "fasii" - stripes), mirroring (oglundirea, copierea datelor) sau parity (tehnica resopnsabila pentru generarea datelor de corectie si control a erorilor).

RAID 0

RAID 0 (sau "stripe set") inparte blocurile de date in mod egal pentru fiecare doua sau mai multe discuri fara date ce asigura redundanta. De mentionat este ca RAID 0 nu este inclus in primele documente ce definesc tehnicile RAID si nu ofera reundanta a datelor. RAID 0 este folosit pentru a imbunatati performantele si pentru a crea un numar mai mic de discuri virtuale.

Timpul de defectare a acestei metode este n ori mai mic decat timpul de defectare a unui singur disc ceea ce maresa vulnerabilitatea in timp al acestui sistem. Pierdere a unui singur HDD duce la o pierdere masiva de date. Cu toate ca exista posibilitatea de recuperare a datelor pierdute prin programe speciale (in cazul in care discul nu este in totalitate stricat), cel mai probabil ele vor fi incomplete si corupte. Astfel recuperarea de date in acest caz este foarte scumpa si ne este garantata.

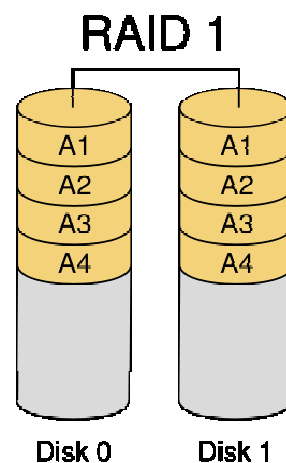
Timpul de acces este imbunatatit in special daca sistemul este folosit in aplicatii ce necesita citiri de date mai mici decat blocurile de date scrise (de obicei dimensiuni de 512B). Un exemplu bun ar fi bazele de date. La fisierele de dimensiuni mari timpul de acces este la fel ca al unui singur disc. Avantajul apare atunci cand datele sunt dispersate in mod egal pe discuri si cand zonele de interes sunt mai mici ca dimensiuni decat unitatea de stripe. Putin mai norocosi au fost cei ce au incercat performantele in divertisment (jocuri) inregistrandu-se performante minimale sau chiar negative.



RAID 1

Functionarea este simpla, RAID 1 creaza copii a discului, imagini in oglinda. Acesta este folositor cand este nevoie de performante bune la citire, fata de performante asupra capacitatii maxime, teoretic ajungandu-se la viteze duble de acces la citire. Astfel de matrice poate avea marimea maxima cat al celui mai mic disc din sistem. Pentru ca fiecare HDD contine o copie a informatiei creste (exponential) redundanta sistemului.

RAID 1 are multe avantaje administrative. Spre exemplu se poate crea un backup al sistemului fara intreruperea serviciilor prin simpla deconectare a discului si apoi copierea lui pe un sistem de backup. Apoi se introduce discul la loc in sistemul RAID si sistemul isi "reconditioneaza" imaginea dupa ultimele modificari.



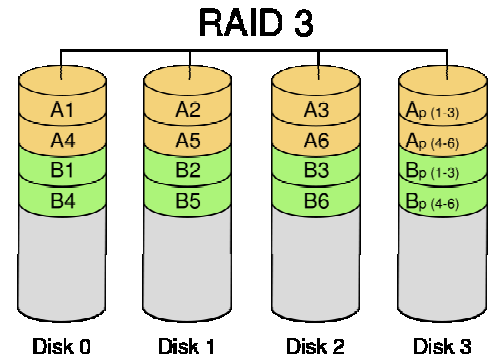
RAID 2

Cu toate ca este o metoda neintalnita in practica ea este singura ce ofera detectie si corectie a erorilor cu acuratete. Celelalte metode pot detecta erori, pot corecta o parte din ele dar in cazuri mai complicate nu se pot descurca fara interventie umana. Principiul este simplu. Datele sunt impartite in Bytes pentru care sunt calculate datele redundante conform codului Hamming pt detectie si corectie a erorilor. Discurile sunt impartite in proportie 4 la 3 (4 discuri date , 3 discuri biti de control).

Dezavantajul este ca discurile trebuie sa se roteasca exact la aceiasi viteza, necesitand un acces sincron asupra lor ceea ce duce la un cost ridicat de productie.

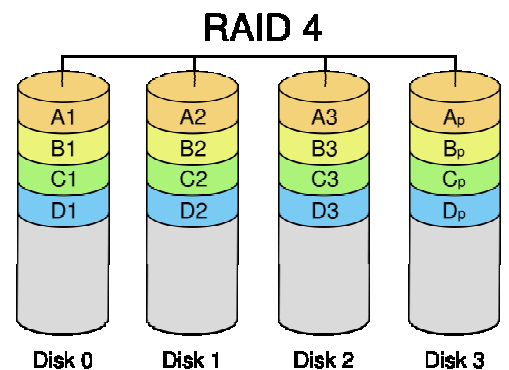
RAID 3

La fel ca si RAID 2 si aceasta metoda este mai putin intalnita. Foloseste imparte blocul de date la nivel de byte si foloseste un disc separat pentru calcularea bitilor de paritate (prin XOR logic). Efectul negativ este ca nu poate face fata la mai multe cereri de date simultan, datele unui pachet fiind imprastiat pe toate discurile.



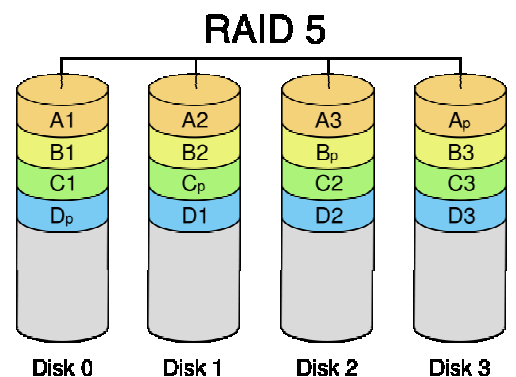
RAID 4

RAID 4 foloseste, spre deosebire de 3, stripe-uri la nivel de bloc de date. Asta permite fiecarui disc sa se comporte independent, accesul fiind mult imbunatatit. Foloseste n discuri pentru date si unul pentru calculul paritatii. De obicei RAID 4 este implementat hardware cu suport pentru calcularea paritatii.



RAID 5

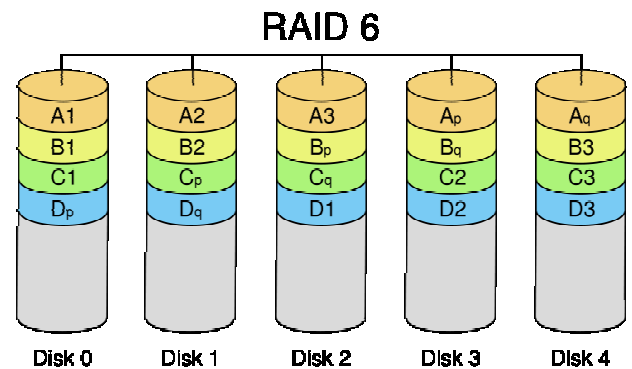
Asemănător cu RAID 4 folosește stripe-uri la nivel de bloc de date, cu condiția ca pachetele de paritate să fie împartite uniform în toate discurile. El păstrează avantajul accesului simultan asupra discurilor. Diferența față de RAID 4 este că, atunci când un disc cedează RAID 4 nu trebuie să calculeze datele de paritate (în exemplul de mai sus, în cazul defectului 1/4 din disc reprezintă date redundante ce nu trebuie să fie calculate). Performanțele sunt asemănătoare cu RAID 0 pentru același număr de discuri. În cazul unui defect în timpul scrierii datelor blocul de paritate poate să devină eronat și va duce la o



reconstituire defecta. Aici intervin controlurile RAID cu memorie cache nevolatila pentru a reduce acest risc.

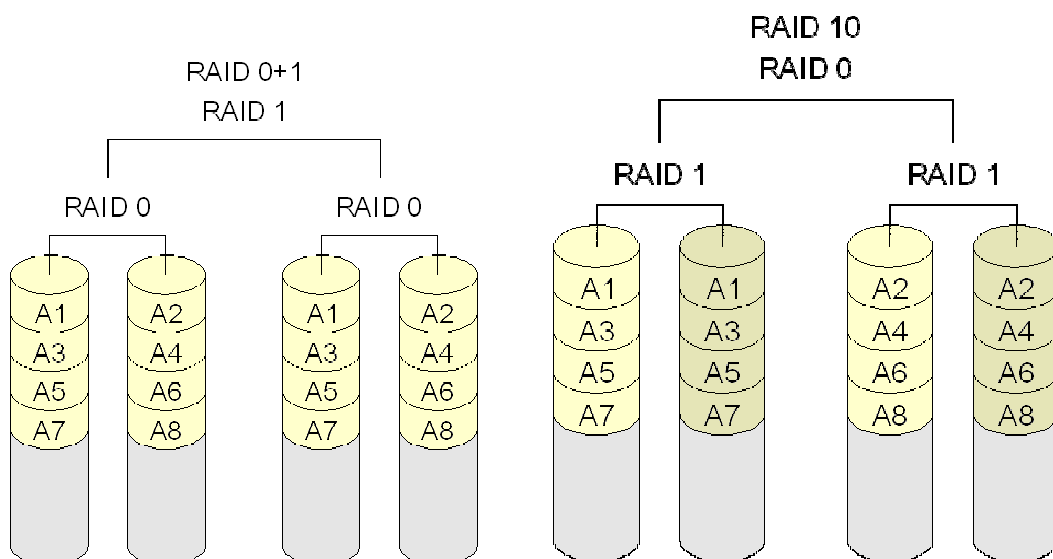
RAID 6

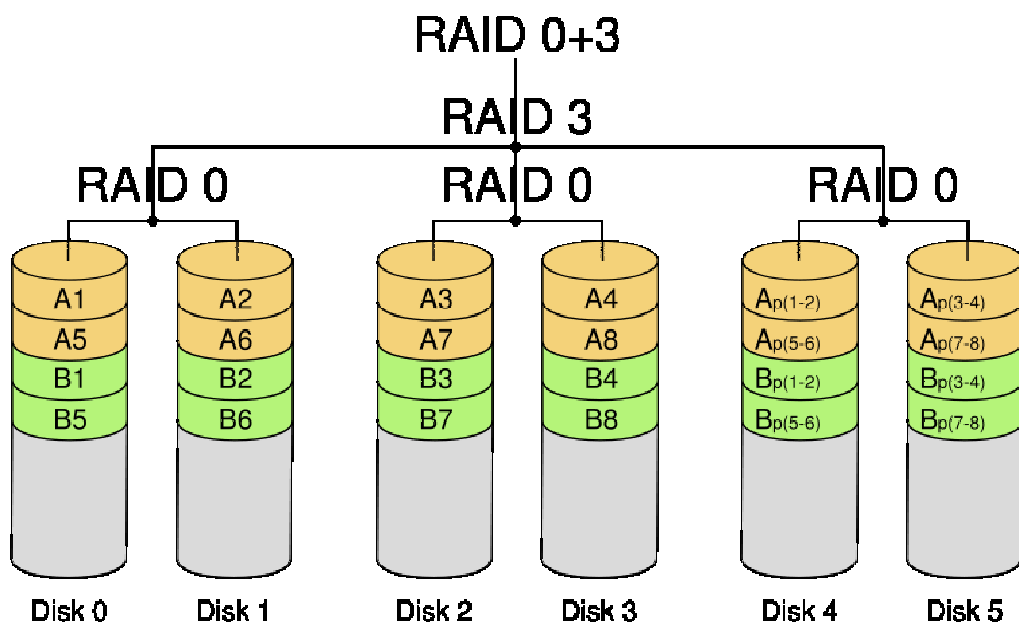
Noutatea adusa fata de RAID 5 este introducerea unui bloc de date de control al paritatii in plus. Aceasta metoda poate fi vazuta ca si un caz mai special al codarii Reed-Solomon. In principiu aceasta tehnica permite (conform definitiei SNIA - Storage Networking Industry Association) continuarea operatiilor de scriere si citire cu doua discuri defecte in sistem.



Nivelurile RAID Hibride (Nested - cuibarit)

Principiul de functionare a nivelelor hibride este de a folosi in locul discurilor fizice alte nivele RAID. Exemple:





Capitolul 9

Functii Win32 API

O interfata API este un cod sursa oferit de catre sistemul de operare sau o librarie pentru a permite apeluri la serviciile care pot fi generate din API-uri respective de catre un program.

Un program care ofera functionalitatea descrisa de interfata API este implementarea interfetei API. Interfata API in sine este abstracta, in sensul ca specifica instanta dar nu se implica in detalii de implementare.

Termenul API este folosit in 2 sensuri:

- O interfata coerenta care consta din cateva clase sau cateva seturi de functii sau proceduri interconectate.
- Un singur punct de intrare, cum ar fi o metoda, o functie sau o procedura.

Doua Interfete API foarte cunoscute sunt Single UNIX Specification si Microsoft Windows API.

Interfete API sunt deseori incorporate in Software Development Kit (SDK)¹.

Modelul de design a Interfetelor API

Exista o multime de modele de design a Interfetelor API. Cele prevazute pentru executie rapida deseori consta din functii, proceduri, variabile si structuri de date. Exista si alte modele cum ar fi interpretatori (emulatori)² care evalueaza expresii in ECMAScript (cunoscut sub nume JavaScript) sau alt layer abstract, oferind programatorului posibilitatea de a evita implicarea in relatiile functiilor cu nivelul inferior al abstractiei.

Unele Interfete API, cum sunt cele standard pentru un sistem de operare, sunt implementate ca librarii de cod separate care sunt distribuite impreuna cu sistemul de operare. Altele au integrata functionalitatea interfetei API direct in aplicatie. Microsoft Windows API este distribuita cu sistemul de operare. Interfetele API pentru sisteme embedded, cum sunt console de jocuri video, in general intra in categoria API-urilor integrate direct in aplicatie.

¹ SDK este un set de unelte de dezvoltare care ofera posibilitate de creare a aplicatiilor pentru un anumit pachet software, software framework, platforma hardware, sistem de operare, console de jocuri video, etc.

² Interpretatori compileaza codul sursa atunci cand acesta este intalnit in executia programului, spre deosebire de compilatori care separa faza de compilare de cea de executie

O interfata API care nu necesita drepturi mari de acces sunt numite "open" (OpenGL ar fi un exemplu).

Doua linii generale ale diplomatiei de publicare a Interfetelor API:

- Unele companii protejeaza informatiile despre Interfete API. De exemplu, Sony a facut Interfata API oficiala pentru PlayStation 2 disponibila doar dezvoltatorilor cu drepturi speciale. Asta a dat posibilitatea companiei Sony de a controla cine scrie jocuri pentru PlayStation 2.
- Unele companii fac Interfete API disponibile in public. De exemplu, Microsoft are o mare parte a informatiilor despre Interfete API disponibile freeware, ce ofera posibilitatea de dezvoltare a programelor pentru platforma Windows

Cateva exemple de Interfete API:

- Single UNIX Specification (SUS)
- Microsoft Win32 API
- Java Platform, Enterprise Edition API's
- OpenGL cross-platform API
- DirectX for Microsoft Windows
- Google Maps API
- Wikipedia API
- Simple DirectMedia Layer (SDL)
- svgalib pentru Linux si FreeBSD
- Carbon si Cocoa pentru Macintosh OS

Microsoft Windows API's

Windows API, neoficial WinAPI, este numele dat de catre Microsoft pentru un set de Interfete API disponibile in sisteme de operare Microsoft Windows. Aceste interfete au fost construite pentru a fi folosite de catre programatori C/C++ si sunt cel mai direct mod de a interactiona cu sistemul Windows pentru aplicatii software. Accesul la nivel inferior la sistemul Windows, in general necesar pentru drivere, este oferit de catre Windows Driver Foundation in versiunea curenta a Windows-ului.

In sisteme de operare Windows este disponibil un Software Development Kit (SDK), care ofera documentatia si unelte pentru a permite dezvoltatorilor crearea aplicatiilor folosind Interfete API si tehnologii Windows asociate.

Istoria

Interfetele API Windows au oferit dintotdeauna acces la structura sistemelor Windows. Din acest motiv ele sunt construite in mare parte pentru programatori. Programatorilor li s-a oferit multa flexibilitate si putere in dezvoltarea aplicatiilor. In aceasi timp aplicatiilor Windows li s-a impus mare responsabilitate in manipularea nivelelor inferioare.

Pe parcurs au fost facute multe modificari la sistemul de operare Windows si Interfetele API Windows au fost de asemenea schimbate pentru a tine pasul cu sistemul de operare. Interfetele API pentru Windows 1.0 au avut mai putin de 450 de functii, in comparatie cu API-uri moderne care au mii de functii. Avand asta in vedere, Microsoft a pus mare accent pentru compatibilitatea inversa, adica compatibilitatea API-urilor noi cu API-uri din urma.

Una dintre cele mai mari schimbari facute de Microsoft era schimbarea de la Sisteme de Operare pe 16 biti la Sisteme de Operare pe 32 biti. Pentru a oferi compatibilitate, Microsoft a scris, pentru noile versiuni pe 32 biti, o schema complexa de Interfete API care permiteau codului scris pe 32 biti sa apeleze cod scris pe 16 biti (si invers in unele cazuri limitate).

Aproape fiecare noua versiune a sistemului de operare Windows a introdus schimbari in Interfete API Windows. Numele a ramas consistent intre diferite versiuni. In cele din urma Microsoft a schimbat numele din Win32 API in Windows API.

Desi Microsoft detine drepturile de autor asupra Interfetelor API Windows, in general este acceptat ca alti producatori sa emuleze Windows oferind Interfete API identice. Proiectul Wine este cea mai cunoscuta incercare de a oferi compatibilitate Interfetelor API Windows pe sisteme de operare Unix-like. *ReactOS* a mers cu un pas mai departe si a oferit emulare intregului sistem de operare Windows. *HX DOS-Extender* este alt proiect care emuleaza Interfetele API Windows , pentru a permite rularea programelor Windows mai simple din linia de comanda DOS.

Compilatoare suportate

Pentru a dezvolta software care foloseste Interfetele API Windows, este nevoie de compilator care poate importa si manipula fisierele DLL si obiectele COM caracteristice Microsoft-ului. Compilatorul trebuie sa accepte dialectul limbajelor C sau C++ si sa manipuleze IDL (interface definition language) fisiere si fisiere header care expun denumirile functiilor interioare ale Interfetelor API. Aceste compilatoare, unelte, librarii si fisiere header sunt impreunate in Microsoft Platform SDK (Software Development Kit). Pentru mult timp familia

de compilatoare si unelte Microsoft Visual Studio si compilatoare Borland, au fost singurele care puteau la cerintele sus mentionate. Acuma exista *MinGW* si *Cygwin* care ofera interfete bazate pe *GNU Compiler Collection*. *LCC-Win32* este disponibil pentru utilizare non-comerciala, continand compilator C si intretinut de catre Jacob Navia. *Pelles C* este compilator C gratuit oferit de catre Pelle Orinius.

Componentele Interfetelor API in Windows

Functionalitatea oferita de Interfete API Windows poate fi grupata in sapte categorii:

- Base Services (Servicii de baza)
Ofera acces resurselor fundamentale disponibile in Windows. Sunt incluse sisteme de fisiere, dispozitive, procese, fire de executie, acces la registri Windows si tratarea erorilor. Aceste functii se afla in fisierele *kernel32.dll* si *advapi32.dll* pe sisteme de operare pe 32 biti.
- Graphics Device Interface (Interfata Dispozitivelor Graice)
Ofera functionalitate pentru afisarea continutului grafic pe monitoare, imprimante si alte dispozitive de iesire. Se afla in *gdi32.dll* pe sistemele de operare pe 32 biti.
- User Interface (Interfata utilizator)
Ofera functionalitati pentru a crea si manipula ferestre si majoritatea controlaelor de baza, cum ar fi butoane si scrollbar-uri, pentru a recepta intrare de la mouse si tastatura, impreuna cu alte functionalitati asociate cu GUI (Graphical User Interface). Aceste functii se afla in *user32.dll*. De la Windows XP, controalele de baza se afla in *comctl.dll*, impreuna cu controalele generale (Common Control Library).
- Common Dialog Box Library (Biblioteca Ferestrerol de Dialog Generale)
Ofera aplicatii si ferestre de dialog standard pentru deschiderea si salvarea fisierelor, alegerea culuorii si fontului, etc. Linraria se afla in fisier numit *comdlg32.dll* si mai tarziu de la Windows 95 in fisier *shlwapi.dll*. Este grupat sub categoria *User Interface* a Interfetelor API.
- Common Control Library (Libraria de Control Generala)
Ofera acces aplicatiilor la unele controale avansate ale sistemului de operare. Acestea includ status bars, progress bars, toolbars si tabs. Libraria se afla intr-un fisier DLL numit *comctl32.dll*. Este grupat sub categoria *User Interface* a Interfetelor API.
- Windows Shell

Componenta Interfetelor API Windows care permite aplicatiilor acces la functionalitatile shell ale sistemului de operare, la fel ca si schimbarea lui in vederea imbunatatirii. Acesta componenta se afla in fisierul shell32.dll si mai tarziu incepand cu Windows 95 in shlwapi.dll. Este grupat sub categoria *User Interface* a Interfetelor API.

- Network Services (Servicii de retea)

Ofera acces la posibilitatile sistemului de operare legate de retea. Subcomponentele lui includ NetBIOS, Winsock, NetDDE, RPC si multe altele.

Interfete API asociate Web

Internet Explorer contine multe Interfete API care sunt la randul sau folosite de alte aplicatii. Aceste Interfete API pot fi considerate ca parte a Interfetelor API Windows. Internet Explorer ofera:

- Controlul browser-ului web incorporat, aflat in shdocvw.dll si mshtml.dll.
- Servicii URL, aflate in urlmon.dll. Aplicatiile pot deasemenea oferi serviciile proprii URL.
- Librarii pentru suport multilingvistic si international (mlang.dll)
- DirectX Transforms, un set de filtre pentru imagini
- Suport XML (MSXML componente)
- Acces la Windows Address Book

Interfete API pentru interactionare intre programe

In mare parte Interfetele API Windows se ocupa de interactionare intre Sistemul de Operare si aplicatii care ruleaza pe el. Pentru a face posibila comunicarea intre diferitele aplicatii, Microsoft a dezvoltat o serie de tehnologii incorporate in Interfetele API Windows. Incepand cu Dynamic Data Exchange (DDE), care a fost inlocuit cu Object Linking and Embedding (OLE) si mai tarziu cu Component Object Model (COM)³.

Interfete API asociate multimedia

³ COM este o platforma Microsoft introdusa in 1993. Este folosita pentru a permite comunicare intre procese si crearea dinamica de obiecte in orice limbaj care suporta tehnologia respectiva. Implementarea nu depinde de limbaj. O inlocuire pentru COM este Microsoft .NET framework, iar suport pentru servicii web fiind Windows Communication Foundation (WCF). Windows Vista se bazeaza in special pe aceste elemente.

Microsoft a oferit un set de interfete API numit DirectX ca fiind incorporat in toate kiturile de instalare inca de la Windows 95 OSR2. DirectX ofera un set de servicii pentru jocuri si multimedia.

Microsoft DirectX

Microsoft DirectX este o colectie de Interfete API folosita pentru manipularea taskurilor legate de multimedia, in special programarea jocurilor si video, pe platforme Microsoft.

DirectX este de asemenea folosit si de alti producatori software, in mare parte in sectorul de inginerie, din cauza abilitatii de redare rapida a obiectelor 3D de inalta calitate.

Atat DirectX runtime cat si software development kit (kitul de dezvoltare soft) sunt disponibile gratuit, dar sunt proprietate Microsoft si sunt closed-source (fara posibilitate de schimbare, rescriere, suprascriere). DirectX a fost initial redistribuit de catre dezvoltatori de jocuri impreuna cu kiturile de instalare, dar in ultim timp DirectX a fost inclus in kit de instalare a sistemului de operare (sau in Service Packs). Unii dezvoltatori de jocuri inca mai includ DirectX in kitul de instalare si ofera posibilitate de a-l instala (sau de a face update) dupa instalarea jocului.

Cel mai recente versiuni ale DirectX-ului sunt DirectX 10 si DirectX 9.0L, disponibile exclusiv pentru Windows Vista (motivul fiind, dupa cum sustine Microsoft, faptul ca exista schimbari in arhitectura grafica a Windowsului si din cauza introducerii Windows Display Driver Model).

Interfetele API din DirectX

Majoritatea Interfetelor API din DirectX sunt in forma de obiecte COM. Componentele care le contine DirectX sunt:

- DirectX Graphics, alcatuit din doua Interfete API (DirectX 8.0 sau mai mult):
 - DirectDraw: pentru generarea de obiecte grafice 2D (acum dezaprobat, desi este inca folosit de multi programatori)
 - Direct3D (D3D): pentru generarea obiectelor grafice in 3D
- DirectInput: pentru tastatura, mouse, joystick sau alte controlere pentru jocuri (nu mai este folosit decat pentru Xinput, in controlere la Xbox 360)

- DirectPlay: pentru comunicare jocurilor in retea (Impreuna cu DirectInput a fost folosit ultima data in DirectX 8. Acuma este dezaprobat)
- DirectSound: pentru redarea sunetului si inregistrarea sunetului in forma wave
 - DirectSound3D (DS3D): pentru redarea sunetelor 3D
- DirectMusic: pentru redarea sunetelor autorizate de catre DirectMusic Producer
- DirectX Media: contine DirectAnimation pentru animatii web 2D, DirectShow pentru redare multimedia, DirectX Transform pentru interactionare web si Direct3D Retained Mode pentru obiecte grafice 3D de nivel inalt. DirectShow mai contine DirectX Plugins pentru procesarea semnalului audio si DirectX Video Acceseration pentru a accelera redarea video.
- DirectX Media Objects: suport pentru pentru obiecte cu streamuri cum sunt encodere, decodere si efecte
- DirectSetup: pentru instalarea componentelor DirectX (care de fapt nu prea este API)

Istoria

Spre sfarsitul anilor 1994 Microsoft era pe punct de a lansa urmatorul sistem de operare - Windows 95. Factorul care era sa determine succesul sistemului de operare era ce, de fapt, va putea fi rulat pe acest sistem de operare. Trei angajati la Microsoft (Craig Eisler, Alex John si Eric Engstorm) erau ingrijorati de faptul ca programatorii considerau sistemul de operare DOS (lansat tot de Microsoft) o platforma mai buna pentru dezvoltarea jocurilor. Asta insemna ca putine jocuri vor fi dezvoltate pentru Windows 95 si asta va influenta succesul acestui sistem de operare.

DOS oferea acces direct la placa video, tastatura, mouse-ul, dispozitivele de sunet si alte componente ale sistemului, in timp ce Windows 95, avand memoria protejata, restrictiona acces la toate acestea. Microsoft avea nevoie de o modalitate de a oferi programatorilor ceea ce doresc. Eisler, John si Engstorm au dezvoltat o solutie care in cele din urma era sa fie numita DirectX.

Prima versiune a DirectX-ului a fost lansata in Septembrie 1995 avand nume Windows Games SDK. Ea era, de fapt, inlocuire pentru Interfetele API din sistemele pe 16 biti, care aveau un design foarte slab. In mare, prima versiune de DirectX ofera posibilitate de a incorpora elemente multimedia de inalta performanta.

Înainte de existența DirectX-ului, Microsoft avea deja inclus setul de API-uri OpenGL pe platformele Windows NT. În acel timp OpenGL necesita hardware de înaltă performanță și era limitat pe inginerie și utilizatori CAD⁴. Direct3D (introdus de persoane care au dezvoltat DirectX ca o alternativă pentru OpenGL) urma să fie o concurență pentru (atunci) mai costisitor (din punct de vedere hardware) OpenGL pentru dezvoltarea jocurilor. Cu timpul, după ce puterea plăcilor video au crescut, OpenGL a devenit standard și conducătorul pieței. În acel punct a început "batalia" între adepți ai cross-platform OpenGL și Windows-only Direct3D, care după cum mulți susțineau era încă un exemplu de adoptare, extindere și eliminare (embrace, extend and extinguish)⁵ de la Microsoft. Alte Interfețe API din DirectX sunt deseori combinate cu OpenGL în crearea jocurilor video datorită faptului că OpenGL nu include toate funcționalitățile care le are DirectX (cum ar fi suport pentru sunet sau joystick). Totuși, combinația între OpenGL și OpenAL (Open Audio Library) pentru acest scop a devenit populară.

DirectX este folosit ca și bază pentru Interfețele API de la console Xbox și Xbox 360, oferite de către Microsoft. Interfețele API au fost dezvoltate de către Microsoft și NVIDIA (care a dezvoltat hardware folosit de consola). Interfețele API din Xbox sunt similare cu cele din DirectX 8.1, dar fără posibilitate de update. Xbox a fost numit DirectXbox, dar numele a fost prescurtat pentru motive comerciale.

În 2002 Microsoft a lansat DirectX 9 cu suport pentru programe *shader*⁶ mult mai lungi cu *pixel* și *vertex shader*, versiunea 2.0. Microsoft a continuat cu update, și în August 2004 a introdus *shader* model 3.0 în DirectX 9.0c.

Compatibilitate

Producători hardware sunt nevoiți să scrie drivere și să testeze fiecare bucată aparte, pentru a le face compatibile cu DirectX. Unele dispozitive hardware au doar drivere compatibile cu DirectX (ceea ce înseamnă că utilizatorul trebuie să instaleze DirectX pentru a folosi acele dispozitive). Versiunile mai vechi de DirectX includeau și o bibliotecă care conținea toate driverele cunoscute și disponibile la momentul lansării versiunii respective. La această practică s-a

⁴ Computer - aided design (CAD) este un set de unelte folosite pe calculator pentru a oferi asistență pentru ingineri, arhitecți și alți designeri profesioniști pentru dezvoltările lor.

⁵ Este o frază pentru a descrie strategia de la Microsoft care presupune acceptarea standardelor folosite, extinderea lor cu capacități proprii și apoi folosirea diferentelor obținute pentru a distruge concurența. Prima dată apărut în New York Times în 1996 în articolul "Microsoft Trying to Dominate the Internet".

⁶ Shader, în grafica calculatoarelor, este un set de instrucțiuni software, care este folosit de resurse grafice pentru efecte de redare.

renunatat in favoarea sistemului de reactualizare, bazat pe web, Windows Update care permite utilizatorilor sa descarce doar drivere relevante pentru dispozitivul lor.

Inainte de DirectX 10, DirectX era considerat ca fiind compatibil cu versiuni precedente (versiunile noi erau compatibile cu cele vechi). Asta este consecinta pozitiva a modelului COM folosit pentru API-ul DirectX.

Cu Windows Vista si Direct3D 10 cu schimbari radicale, compatibilitate nu mai este posibila. Din acest motiv DirectX 10 ofera Interfata API Direct3D 9, astfel incat sa fie posibila rulara jocurilor si aplicatiilor mai vechi.

DirectX 10 si 9.0L

Windows Vista este lansata cu DirectX 10 (si 9.0L pentru compatibilitate cu versiuni mai vechi) si este singurul sistem de operare din familia Windows care il contine. Schimbari radicale au fost facute: *DirectInput* este defavorizat si inlocuit cu Xinput, din Xbox. De asemenea, *DirectSound* este defavorizat si inlocuit cu XACT⁷. DirectX 10 a renuntat la suport pentru accelerarea dispozitivelor audio, iar sunetul este redat in softearul pe CPU. *DirectPlay* este defavorizat si inlocuit cu *Games for Windows - LIVE*⁸, in timp ce *DirectShow* este defavorizat si inlocuit cu Media Foundation, un set diferit de Interfete API lansate impreuna cu Windows Vista, pentru manipularea de secvente audio si video. *DirectMusic* va fi probabil singurul ramas intact.

Direct3D

O proprietate noua majora in DirectX 10 este Direct3D (initial numit Windows Graphics Foundation). Folosind noul Windows Display Driver Model, Shader Model 4.0 si noile, mai stricte, cerinte pentru producatori de GPU (Graphical Processing Unit) care trebuie satisfacute pentru a sustine compatibilitate cu Direct3D, versiunea 10 de Direct3D reprezinta indepartarea de la practicile versiunilor mai vechi. Pentru a oferi compatibilitate cu versiunile precedente de Direct3D, DirectX 10 contine de fapt trei versiuni de Direct3D:

- Direct3D 9: aceasta Interfata API emuleaza comportamentele ale Direct3D 9 pe Windows XP pentru a oferi compatibilitate cu aplicatii

⁷ XACT este biblioteca audio lansata de catre Microsoft ca fiind o parte din DirectX SDK. Initial a fost dezvoltata pentru Xbox, iar mai tarziu a fost modificata sa fie compatibila si cu Microsoft Windows.

⁸ Games for Windows - LIVE este serviciul online de jocuri pentru Windows, care permite calculatoarelor, ruland Windows Vista, sa se conecteze la LIVE service (acest serviciu a fost initial dezvoltat pentru Xbox)

mai vechi. Toate detaliile si avantajele ale Windows Display Driver Model (WDDM) din Vista, sunt ascunse de aplicatie (in caz ca WDDM drivere sunt instalate). Aceasta este singura Interfata API disponibila in caz ca exista doar drivere grafice XP.

- Direct3D 9Ex (initial numit 9.0L): permite acces total la noile capabilitati ale WDDM, in timp ce pastreaza compatibilitate pentru aplicatii Direct3D existente folosind o interfata API separata. Efectul de transparenta ("glass effect") din *Windows Aero* se bazeaza pe codul din Direct3D 9Ex. Cand 9Ex era inca sub numele de cod 9.0L, existau zvonuri ca acesta va fi Direct3D 10 pentru Windows XP. Pana la urma s-a demonstrat ca nu este asa, in mare parte din cauza incompatibilitati al WDDM pe Windows XP.
- Direct3D 10: Temporar, singurul GPU compatibil cu Direct3D este NVIDIA GeForce 8 Series, care de fapt profita foarte putin de capabilitatile ale Direct3D 10.

Alternative

Exista alternative pentru aceasta arhitectura, unele mai complete decat altele. Nu exista o solutie unica care sa ofere toate functionalitatile care le ofera DirectX, dar prin combinare de librarii - OpenGL, OpenAL, SDL, OpenML, FMOD, etc - utilizatorul pote sa implementeze o solutie comparabila cu DirectX si in aceasi timp gratuita si cross-platform (poate fi folosite pe diferite platforme).

Bibliografie

- [1] „Microsoft Windows Internals 4th Edition” - by Mark Russinovich
- [2] „Windows NT/2000 Native API Reference” - by Gary Nebbett
- [3] „The Windows NT Device Driver development” - Peter G. Viscarola, W. Anthony Mason
- [4] „Windows NT File System Internals: A developer's Guide” - R. Nagar - capitolul despre driverele filtru
- [5] „Microsoft Extensible Firmware Initiative FAT32 File System Specification” - FAT32 Whitepaper al firmei Microsoft.
- [6] „Microsoft Developer Network” - pe scurt MSDN - <http://msdn.microsoft.com/>
- [7] „Platform SDK Documentation”
- [8] „Windows 2000/XP Device Driver Kit Documentation”
- [9] NIST, "Advanced Encryption Standard (AES)", Federal Information Processing Standards Publication 197, November 26, 2001, - <http://csrc.nist.gov/publications/fips/fips197/fips-197.pdf>.
- [10] C. Adams, „The CAST-128 Encryption Algorithm”, Request for Comments 2144, May 1997, - <http://www.rfc-editor.org/rfc/rfc2144.txt>.
- [11] RSA Laboratories, PKCS #5 v2.0: „Password-Based Cryptography Standard”, RSA Data Security, Inc. Public-Key Cryptography Standards (PKCS), March 25, 1999 - <ftp://ftp.rsasecurity.com/pub/pkcs/pkcs-5v2/pkcs5v2-0.pdf>
- [12] H. Krawczyk, IBM, M. Bellare, UCSD, R. Canetti, IBM, HMAC: Keyed-Hashing for Message Authentication, Request for Comments 2104, February 1997, - <http://www.rfc-editor.org/rfc/rfc2104.txt>.
- [13] Peter Gutmann, „Software Generation of Practically Strong Random Numbers”, 1998 Usenix Security Symposium, - <http://www.cs.auckland.ac.nz/~pgut001/pubs/usenix98.pdf>.
- [14] Victor Valeriu Patriciu, Monica Ene Pietroseanu, Ion Bica, Costel Cristea - "SECURITATEA INFORMATICA ÎN UNIX si INTERNET" - Editura Tehnica, 1998
- [15] Dr. Brian Gladman – „Rutine AES, Serpent, și Twofish”.
- [16] Eric Young – Librăriile libdes, libcast și blowfish.
- [17] Multe alte resurse electronice de pe INTERNET.
- [18] www.Wikipedia.org