

**Universitatea Politehnica București**  
**Facultatea de Electronică, Telecomunicații și Tehnologia**  
**Informației**

**Proiect Sisteme de operare avansate**  
**Nucleul sistemului de operare UNIX (KERNEL)**

**Coordonator:**  
***Conf. dr. ing. Ștefan Stăncescu***

**Masterand IISC:**  
**Simona-Ioana Barbu**

**2013**

# **Cuprins**

## **Capitolul 1: Introducere**

- 1.1. Evoluție
- 1.2. Caracteristici
- 1.3. Arhitectura sistemului UNIX

## **Capitolul 2: Structură și funcții**

## **Capitolul 3: Procese și stări**

## **Capitolul 4: Coordonarea și gestiunea proceselor**

- 4.1. Planificarea proceselor
- 4.2. Alocarea memoriei
- 4.3. Divizarea timpului și resurselor

## **Capitolul 5: Sincronizarea proceselor**

- 5.1. Sincronizarea prin intermediul semnalelor
- 5.2. Sincronizarea prin semafoare
- 5.3. Sincronizarea prin directive de sistem

## **Capitolul 6: Comunicația între procese**

- 6.1. Comunicația prin fișiere
- 6.2. Comunicația prin fișiere de tip pipe
- 6.3. Comunicarea prin zone de memorie comune
- 6.4. Comunicarea prin cozi de mesaje

## **Capitolul 7: Gestiunea fișierelor speciale**

- 7.1. Fișierele disc

## **Capitolul 8: Bibliografie**

## Capitolul 1: Introducere

Sistemul de operare Unix este cel mai vechi sistem de operare, acesta rezistând și impunându-se până în zilele noastre.

Unix este un sistem de operare time-sharing universal, a cărui caracteristică principală o constituie portabilitatea, mai exact disponibilitatea să pentru diverse tipuri de calculatoare (microcalculatoare, minicalculatoare, supercalculatoare) și pentru rețele puternice de calculatoare.

Ceea ce a dus la apariția acestui sistem de operare se numără printre urătoarele:

- necesitatea standardizării și unificării sistemelor de operare, în special a interfeței cu utilizatorul;
- păstrarea structurii volumelor și fișierelor în condițiile transferului pe alte sisteme de calcul;
- asigurarea unor niveluri superioare de portabilitate a produselor - program;
- posibilitatea interconectării de sisteme cu arhitecturi, tipuri și puteri diferite, sub același sistem de operare.
- independența software-ului de aplicații, față de evoluția hardware.

### 1.1. Evoluție

În anul 1969 a apărut prima versiune experimentală scrisă în limbaj de asamblare pentru minicalculatoarele PDP-11, al firmei DEC (Digital Equipment Corporation) pentru un singur utilizator (monouser); odată cu primul compilator pentru limbajul C (1972), UNIX a fost rescris în acest limbaj cu scopul asigurării portabilității.

În 1978 s-a realizat prima implementare comercială ÎS/1 (Interactive System One), urmată la scurt timp de versiunea XENIX a firmei Microsoft. În 1980, UNIX s-a impus ca principală soluție de standardizare în domeniul sistemelor de operare, reprezentând o modalitate de realizare a sistemelor deschise pentru toate categoriile de sisteme de calcul.

După 1981, firma AT&T elaborează versiunile UNIX System III și V, iar firma Berkeley Software Distributors (BSD) realizează standardele BSD 1,2,3 urmate de BSD 4.2, 4.3 și din 1993 BSD 4.4. Companiile SUN și AT&T au dezvoltat versiunea System V.4 în 1988, iar IBM, DEC și Hewlett

Packard au format Open Software Foundation (OSF) independent de AT&T (fig. 1.1.1).

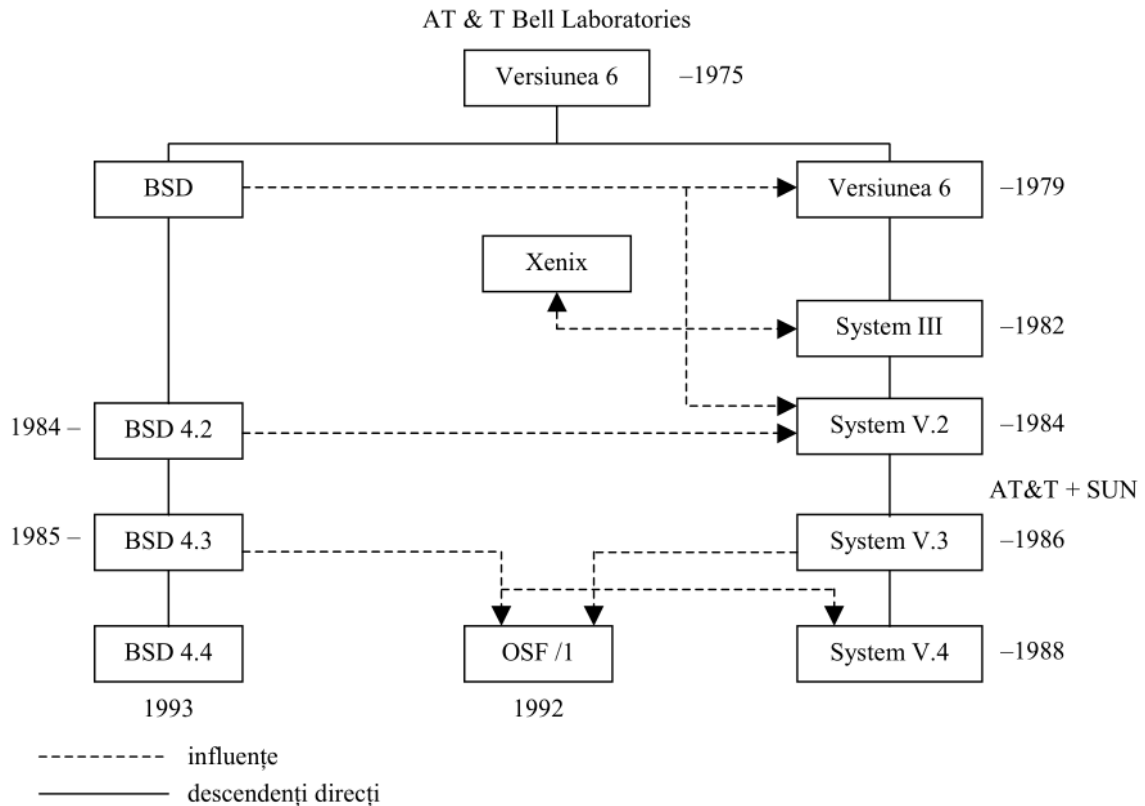


Fig. 1.1.1. Evoluția sistemelor de operare UNIX

## 1.2. Caracteristici

- Sistem de operare time-sharing;
- Protecția fișierelor (prin intermediul parolelor și a drepturilor de acces);
- Intrări/ieșiri generalizate prin care se face integrarea operațiilor de intrare/ieșire în sistemul de fișiere;
- sincronizarea proceselor prin intermediul unui sistem de întreruperi logice;
- transferul de pagini între memoria RAM și cea externă care permite managementul spațiului afectat execuției proceselor și controlul timpului de acces la procesele în așteptare;
- interfețe simple și interactive (SHELL) prin care se asigură dialogul utilizatorului cu sistemul de operare;
- portabilitatea sistemului de operare cât și a software-ului de aplicație;

- gama vastă de aplicații: compilatoare, sisteme de gestiune a bazelor de date, rețele de calculatoare, inteligență artificială, simulare, gestiune, statistică, instruire asistată de calculator, etc.;
- permite execuția aplicațiilor în MS-DOS, în paralel cu execuția de procese sub UNIX (submeniu pentru MS-DOS);
- întreținere și dezvoltare facilă.

### 1.3. Arhitectura sistemului UNIX

Arhitectura unui sistem de calcul se prezintă din punct de vedere al funcționalității pe trei niveluri (fig. 1.3.1.)

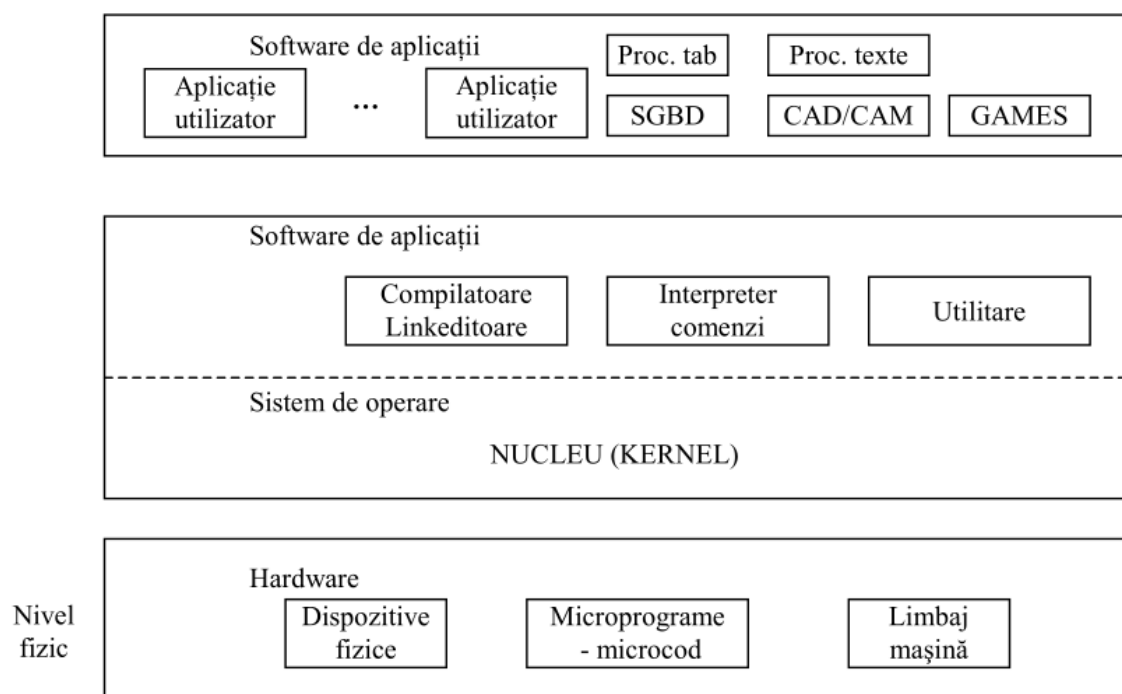


Fig 1.3.1. Arhitectura sistemelor ce lucrează sub UNIX

Nivelul fizic este nivelul inferior. La acest nivel operează hardware-ul, care include dispozitivele fizice, microcodul și limbajul de asamblare.

Următorul nivel - software-ul de bază sau software-ul de sistem are ca principală componentă sistemul de operare având la bază nucleul (Kernel) alături de editoare de texte, compilatoare, interprete de comenzi și utilitare.

Nivelul superior este constituit din software-ul aplicativ și cuprinde practic, o gamă infinită de aplicații.

Se poate observa faptul că orice nivel comunică numai cu nivelul imediat inferior sau superior, furnizând servicii pentru nivelul imediat superior.

Scopul principal al sistemului de operare este acela de a controla resursele hardware ale sistemului, acesta acționând ca o interfață între utilizator (programele de aplicație pe care acesta le lansează în execuție) și componenta hardware (fig. 1.3.2.).

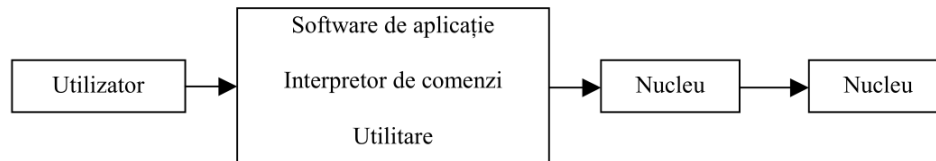


Fig. 1.3.2. Sistemul de operare UNIX- interfața între hardware și utilizator

Sistemul UNIX lucrează în time-sharing, mai exact este constituit dintr-un nucleu (Kernel) și un număr foarte mare de utilitare accesibile prin intermediul interpretorului de comenzi Shell, interpretor ce reprezintă interfața dintre sistemul de operare și utilizator:<sup>1</sup>

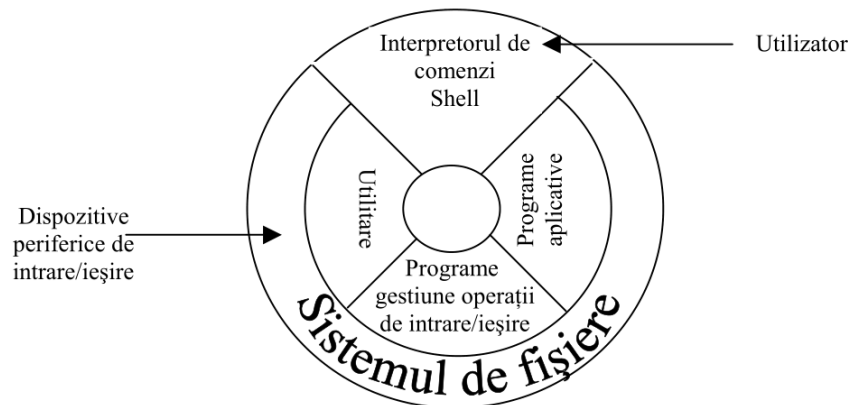


Fig 1.3.3. Componentele sistemului de operare UNIX

- Nucleu (Kernel)
  - partea rezidentă;
  - asigură servicii de sistem pentru programele de aplicații
  - asigură gestiunea proceselor, a memoriei, a intrărilor/ieșirilor și a timpului.

<sup>1</sup> [Structura UNIX, unibuc.ro](http://Structura UNIX, unibuc.ro)

- Shell
  - interpretor de comenzi
  - mecanismul prin care sistemul de operare realizează interfața dintre utilizator și sistemul de calcul
- Sistemul de fișiere cuprinde programe:
  - utilitare
  - aplicative
  - de gestiune a operațiilor de intrare/ieșire
- Utilitare - servicii accesibile prin intermediul interpretorului de comenzi.

Interfețele disponibile utilizatorului sunt organizate pe trei niveluri:

- nivelul exterior nucleului, care poate fi accesat de către utilizator prin intermediul utilităților;
- nivelul intermediar, accesat prin funcții din biblioteca limbajului C;
- nivelul inferior, a cărui accesare se realizează prin funcțiile de sistem (System Calls).

Prin urmare, există o ierarhizare a straturilor ce se interpun între utilizator și nucleul sistemului de operare redată în fig. 1.3.4.

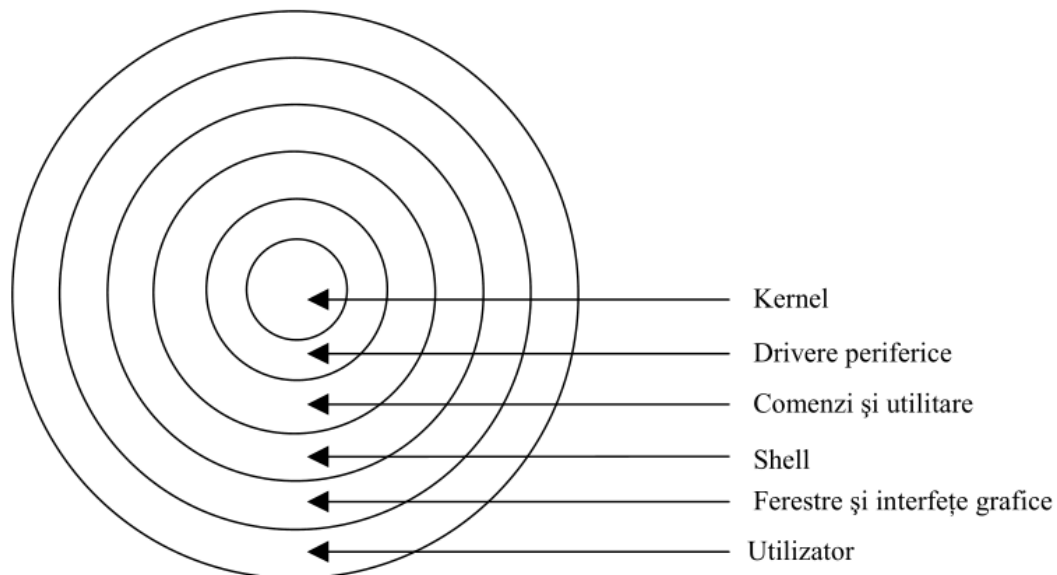


Fig. 1.3.4. Straturile interpușe între utilizator și nucleul sistemului de operare UNIX

- Nucleul (Kernel) UNIX este constituit din două componente principale:
- sistemul de gestiune a fișierelor;
- sistemul de gestiune a proceselor.

## **Capitolul 2: Structură și funcții**

Nucleul sistemului de operare este alcătuit din următoarele componente:

- programul supervisor central;
- rutine de serviciu pentru o serie de activități cum ar fi de exemplu, scrierea în memorie, gestiunea ceasului sistem, etc.

În mod normal, utilizatorul folosește comenzile Shell care dirijează execuția programelor dorite sub supervizorul Kernel, acesta fiind astfel invizibil utilizatorului. Pentru utilizator, sistemul de operare UNIX apare ca fiind alcătuit din:

- un set de programe, fiecare corespunzând unei anumite comenzi;
- Shell-ul care dirijează comenzile și coordonează execuția programelor utilizatorului.

Practic însă, sistemul de operare UNIX este compus din:

- un set de programe de servicii care execută funcțiile legate de sarcinile hardware și software;
- Kernel-ul care coordonează execuția programelor de servicii, specificate sub forma unor comenzi adresate către Shell.

Programele de servicii sunt apelate atunci când sunt necesare (zona de tranziție), în timp ce supervizorul este rezident în RAM, constituind mediul software de bază pentru orice evenimente din sistem.

Principalele funcții îndeplinite de nucleul sistemului de operare:

- planificarea, coordonarea și gestionarea execuției proceselor;
- furnizarea de servicii de sistem cum sunt: tratarea operațiilor de intrare/ieșire și gestiunea fișierelor;
- manipularea operațiilor dependente de hardware, întreruperile și funcțiile de sistem;
- gestiunea memoriei.

Nucleul UNIX este alcătuit din aproximativ 10000 de linii ce constituie codul programului care, în funcție de sistem, se transformă într-un număr mai mare sau mai mic de cuvinte mașină (sau bytes); dintre acestea, 5 - 10% din totalul codului programelor (Shell, utilitare, KERNEL și celelalte) este variabil funcție de sistemul de calcul și de setul de utilitare (fig. 2.1.).

Așa cum rezultă din fig. 2.1, o mare parte din aceste programe este destinată gestiunii memoriei și gestiunii proceselor, această parte evidențiind:

- conținutul stivei;
- conținutul registrelor sistemului;
- detalii de mediu, când procesele sunt introduse în memorie și răspund la întreruperile procesorului.

Această parte conține 7 - 8000 linii codul programelor scrise în limbajul C (deci portabilă pe orice sistem de calcul). Includerea acestei părți în nucleu se justifică prin necesitatea unui răspuns foarte rapid.

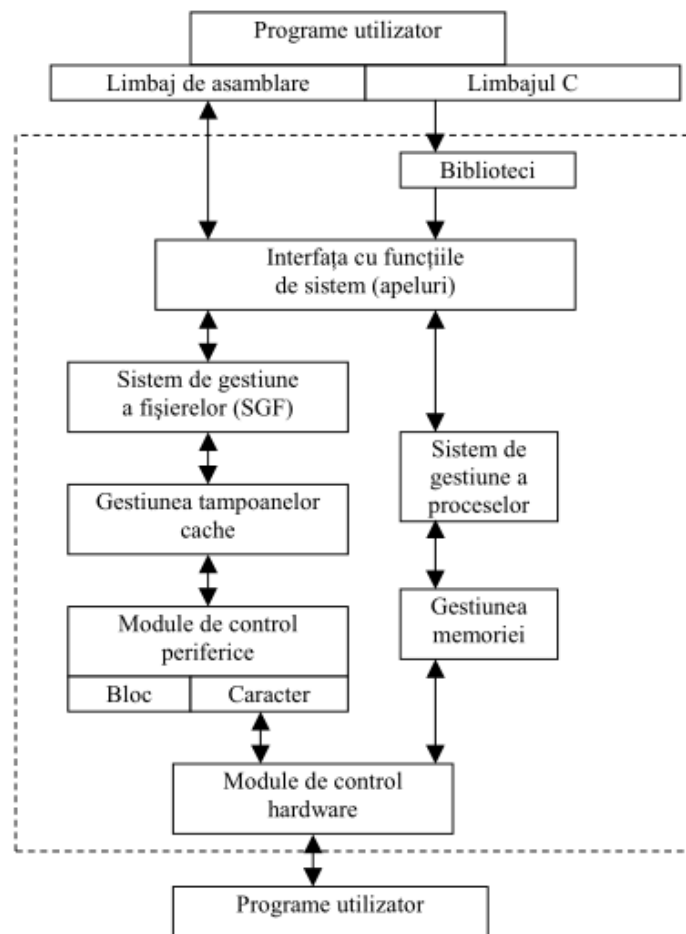


Fig. 2.1.1. Structura Kernel

- Altă parte a Kernel-ului conține driverele dispozitivelor periferice de intrare/ieșire constituite din programe ce realizează:
- controlul adreselor de citire/scriere, adresarea registrelor de date ale perifericelor de
- intrare/ieșire;
- manipularea întreruperilor generate de aceste dispozitive;
- efectuarea recuperării erorilor.

Conține aproximativ 1000 linii codul programelor scrise tot în limbajul C, dar acest număr este variabil în funcție de numărul perifericelor de intrare/ieșire.

Primitivele de sistem sunt specifice fiecărui sistem de calcul (scrise în limbaj de asamblare, aproximativ 1000 linii codul programului), ele conținând:

- operații de intrare/ieșire de bază;
- comutarea execuției între procese;
- permiterea sau inhibarea întreruperilor hardware;
- resetarea priorităților întreruperilor;
- alte operații.

Accesarea primitivelor de sistem se realizează prin apeluri (directive) de sistem (system calls) din programe în C sau în limbaj de asamblare.

### **Capitolul 3: Procese și stări**

Un sistem de calcul poate să lucreze la un moment dat în două moduri:

- utilizator, când execută un program sau proces;
- sistem (Kernel), când execută un cod sistem.

Comutarea între modul utilizator și KERNEL se realizează prin următoarele mecanisme:

- ceasul - care întrerupe orice alt program cu frecvența de 60 Hz, rutina de ceas permițând reevaluarea priorităților proceselor și implicit schimbarea procesului; în absența altor întreruperi, ceasul realizează divizarea timpului, ceea ce permite ca sistemul să fie împărțit între mai mulți utilizatori;

- apeluri de sistem prin care utilizatorul solicită diverse servicii oferite de sistemul de operare; cele care realizează operații de intrare/ieșire conducând la suspendarea procesului apelator pe durata transmiterii datelor;
- cereri de serviciu ale perifericelor de intrare/ieșire.

Procesul este conceput fundamental de organizare a sistemului de operare UNIX. În esență, un proces reprezintă un program în execuție. Pentru un program activ, pot exista mai multe procese active - numite instanțe:

- din punct de vedere al procesului, operațiile nucleului sunt prioritare;
- din punct de vedere al nucleului, procesele sunt structuri de date catalogate.

Informațiile necesare procesului sunt memorate în:

- a) tabela proceselor constituită în memorie, ce conține o intrare pentru fiecare proces detaliind starea acestuia:
  - localizarea procesului;
  - adresa de memorie și adresa de evacuare;
  - mărimea procesului;
  - numărul de identificare;
  - identificatorul utilizat;
- b) tabela utilizator, care este atașată fiecărui proces activ la rutinele nucleului lui (fig. 3.1.).

Crearea unui proces implică astfel inițializarea unei intrări în tabela procesului, care inițializează o tabelă a utilizatorului creând totodată textul real și datele pentru proces.

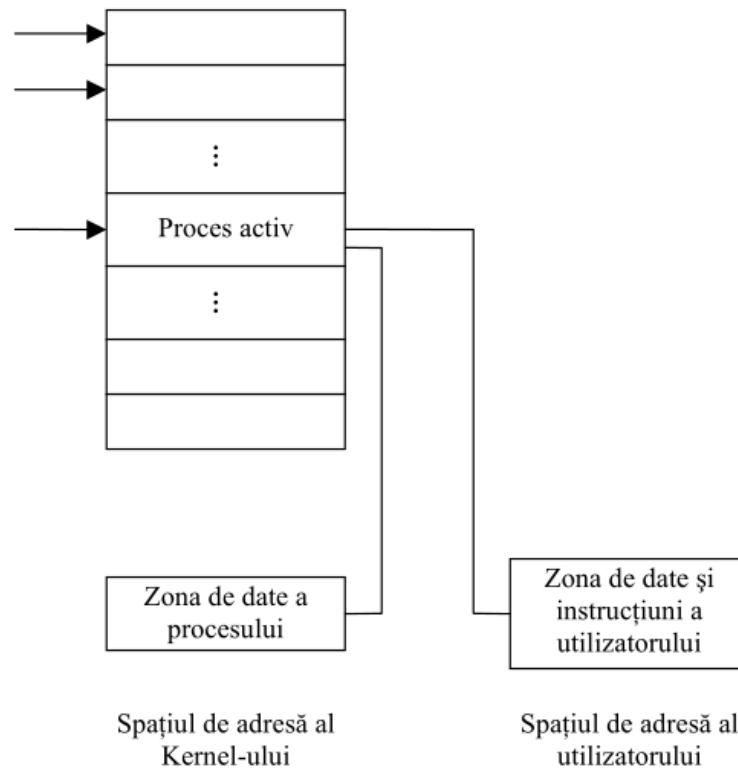


Fig. 3.1. Inițializarea tabelii utilizator

Schimbarea stării unui proces (se execută, așteaptă, este evacuat, este încărcat, primește un semnal) este o activitate ce se focalizează pe tabela procesului (fig. 3.2.):

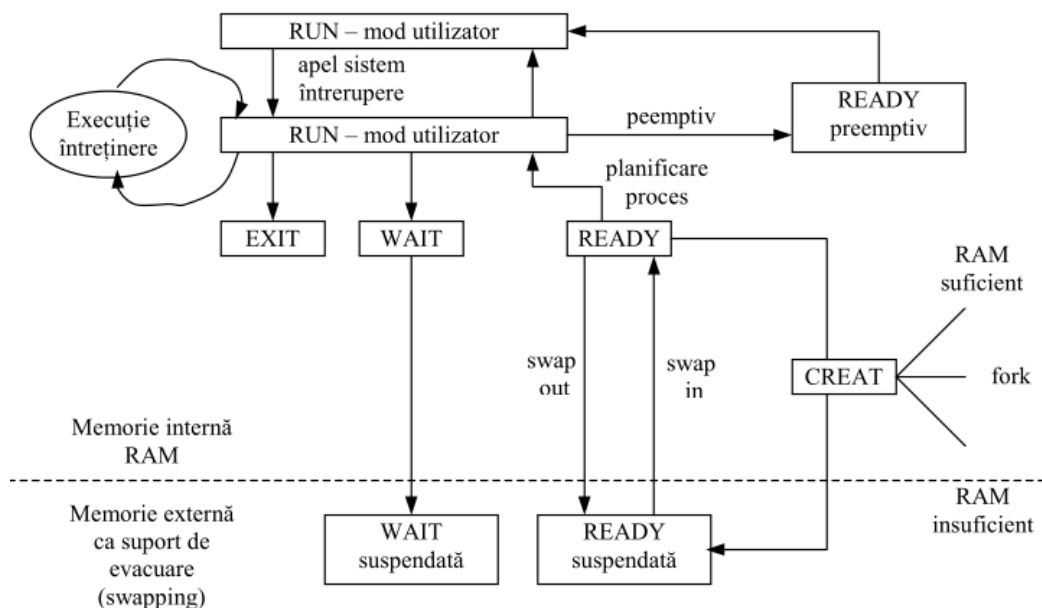


Fig. 3.2. Tranzitia starilor unui proces

- terminarea unui proces implică eliberarea intrării sale în tabela proceselor, putând fi folosită pentru alte procese;
- când procesul este activ, au loc alte evenimente, de exemplu așteptarea terminării unei operații de intrare/ieșire;
- pe durata suspendării unui proces, tabela utilizatorilor nu este accesată sau modificată;
- dacă procesul este evacuat pe disc, este evacuată și tabela utilizator în cadrul imaginii procesului;

Tabela utilizator (structură a utilizatorului) conține în acel moment:

- numele de identificare al utilizatorului și al grupului din care face parte, pentru stabilirea drepturilor de acces;
- pointerii din tabela de fișiere existente în sistem, pentru toate fișierele deschise;
- un pointer către inod-ul directorului curent în tabela de inod-uri;
- o listă a răspunsurilor pentru diverse semnale.

Informația curentă despre un proces se schimbă prin execuția directivei *chdir* când este schimbată valoarea pointerului către inod-ul directorului curent.

## **Capitolul 4: Coordonarea și gestiunea proceselor**

Într-un sistem de operare multiutilizator, se pot executa mai multe programe simultan. Existând o singură CPU, în realitate un singur program se execută efectiv la un moment dat, iar celelalte se vor executa atunci când CPU le va acorda cuante de timp. În timpul execuției unui program, celelalte programe sunt rezidente în RAM, cu condiția ca nici unui să nu solicite toată memoria. De aici decurg 2 sarcini ale nucleului:

- planificarea proceselor - împărțirea de cuante de timp între procese;
- gestiunea memoriei - atribuirea de zone de memorie liberă la procese și eventuala evacuare au încărcare a unui proces în/din disc.

## 4.1. Planificarea proceselor

Inițierea unui proces poate fi efectuată numai de către un alt proces activ. Când este conectat primul utilizator, Kernel-ul asigură o copie a Shell-ului ce se rulează doar pentru el, ivindu-se existența unei alte structuri ierarhice de procese create printr-un mecanism numit bifurcație (fork) prin care Kernel înlocuiește un proces existent, prin două procese:

- procesul inițial (părinte);
- procesul inițiat de procesul inițial (fiu), care împarte (partajează) toate fișierele cu procesul părinte.

După bifurcație, ambele procese se execută independent; excepție: când se solicită explicit că procesul părinte să aștepte terminarea procesului fiu prin directiva WAIT; în continuare, procesul fiu poate genera la rândul lui, o nouă bifurcație (fig.4.1.1.).

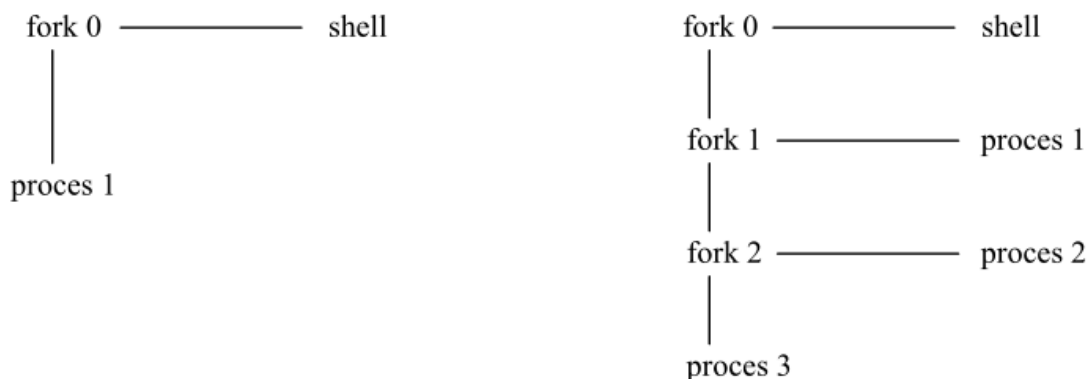


Fig. 4.1.1. Mecanismul *fork*

- De remarcat că *proces3* are acces la toate fișierele deschise de procesele anterioare, în timp ce invers nu se permite accesul. Prin urmare, fișierele sunt accesibile de pe nivelurile situate la periferia ierarhiei.

Când Shell inițiază un proces, bifurcația este aranjată astfel încât Shell să aștepte terminarea procesului (procesul este executat în foreground). Dacă nu așteaptă, înseamnă că Shell a dat naștere altor procese mai puțin prioritare (processe executate în background) adăugând după comanda corespunzătoare semnul **&**.

## 4.2 Alocarea memoriei

Fiind un sistem de operare multitasking, UNIX ține evidența taskurilor concurente și menține controlul diverselor programe care sunt rezidente în RAM la un moment dat, furnizând și informații privind suficiența spațiului de memorie:

- în mod normal, fiecare program este încărcat în diverse zone de memorie RAM;
- în modul de operare cu divizarea timpului (time-sharing), poate opera fără evacuarea conținutului memoriei pe disc; fiecare program se execută în cuanta de timp alocată rămânând în memorie numai programele a căror execuție nu s-a încheiat (practic, numai registrele sistemului sunt divizate între programe, astfel că evacuarea programelor implică de fapt, numai evacuarea conținutului registrelor);
- pe durata execuției proceselor, UNIX alocă porțiuni distincte de RAM pentru:
  - segmentul de cod (instrucțiunile programului) protejat la scriere;
  - segmentul de date ce conține toate datele definite de utilizator: valori, variabile, etc.;
  - segmentul de stivă care conține toate informațiile de sistem necesare menținerii intacte a unui proces, atunci când este evacuat din RAM pe disc sau când este încărcat în RAM de pe disc.

Această divizare este realizată din două motive:

- evacuarea segmentului de cod implică reducerea cantității de informații necesară a fi evacuate; fiind de tip read-only, acesta este identic cu imaginea sa din disc, deci nu trebuie rescris când se evacuează;
- Shell este încărcat pentru fiecare utilizator, deci există atâtea copii Shell, câți utilizatori sunt conectați; dacă mai mulți utilizatori solicită un program (segment de cod), nu este necesară copierea acestuia deoarece el poate fi folosit simultan; este suficient să se creeze și când este necesar, să se evacueze un segment de date și un segment de stivă pentru fiecare utilizator.

### 4.3. Divizarea timpului și resurselor

În sisteme de operare multitasking, procesele individuale nu sunt executate până la terminare, având alocate cuante de timp de către procesor într-o planificare de tip Round- Robin;

Cuantele de timp nu sunt egale, lungimea depinzând de prioritatea taskului, de accesibilitatea datelor de intrare solicitate și a perifericelor de ieșire.

Alocarea se face având în vedere maximizarea utilizării resurselor hardware, cu respectarea priorității taskurilor critice.

Prioritățile sunt periodic reînnoite astfel:

- taskurile cu timp de execuție mai mare, vor avea prioritate mai scăzută;
- taskurile de tip întrebare-răspuns vor avea răspuns instantaneu;
- taskurile ce constă în schimbarea de caractere au cea mai scăzută prioritate.

După expirarea timpului alocat, taskul este evacuat într-un fișier de evacuare pe disc unde vă rămâne până i se dă din nou controlul; în acel moment, imaginea salvată anterior se reîncarcă în RAM; ea va cuprinde:

- conținutul părții de memorie inscriptibilă;
- conținutul registrelor;
- numele directorului curent;
- lista fișierelor deschise de procesul considerat.

În acest sens, sunt utile prezentarea apelurilor de sistem:

***I=fork ()***

Care permite crearea unui nou proces fiu, având active procesele părinte și fiu. Practic, se introduce un nou proces în tabela de procese a sistemului UNIX, dar nu se lansează în execuție nici un program. Se crează astfel o bifurcație în care procesul fiu este realizat că o copie a procesului părinte, nefiind necesară evacuarea programului din memorie ci doar copierea zonelor de date.

***J=execl (nume, arg1, arg2,..., argn, o)***

Solicită execuția unui alt program; nume, arg1, arg2,..., argn sunt pointeri către șiruri de caractere care specifică numele programului și numele

argumentelor sale; *o* - pointer către o zonă de pointeri, fiecare din ei punctând câte un șir din mediu.

Prin acest apel, nucleul va determina ce program original să fie înlocuit cu alt program.

Secvența *fork-execl* crează o copie a procesului părinte, apoi înlocuiește procesul fiu cu programul care dorește să-l execute (procesul fiu nu este alterat de *execl*).

### ***K=wait (status)***

Determină procesul părinte să aștepte terminarea execuției procesului fiu; **Status** este pointer la o valoare întreagă, iar **k** - valoarea returnată ca număr de identificare al procesului fiu terminat.

O aplicație frecventă a lui *execl* se întâlnește la înlănțuirea programelor; de exemplu, o compilare în mai mulți pași implică încărcarea programului și datelor pentru execuție la primul pas; în continuare, se procedează asemănător și la al doilea pas; o dată ce primul pas a fost executat complet, nu există nici un motiv de a se întoarce la el, deci a se suprascrie memoria ocupată la prima trecere cu programul și zona de date.

Înlănțuirea ne este totdeauna practică pentru secvențe în care părți din program solicitate sunt dependente de date; astfel, procesul părinte crează un proces fiu (prin *fork*) cu cerere de așteptare (*wait*); procesul fiu va exista ca proces evacuabil și va avea acces la același spațiu de memorie ca și procesul părinte; prin *execl*, procesul fiu poate executa următorul program, putând el însuși să fie rescris; deci, când procesul fiu modificat și suprascris va fi încheiat, controlul va reveni procesului părinte.

Deoarece numărul de procese concurente nu este limitat, această schemă permite execuția unui mare număr de programe cu condiția să fie posibil a se divide în programe și date astfel încât fiecare execuție să încapă în RAM.

## Capitolul 5: Sincronizarea proceselor

Sincronizarea proceselor poate fi sub controlul sistemului de operare prin intermediul conductelor de comunicație (*pipe*) sau al utilizatorului, prin intermediul evenimentelor care reprezintă modul de determinare al momentului în care un proces devine gata pentru execuție (*READY*).

La un moment dat un proces se află în execuție; blocarea procesului se poate realiza prin funcția *sleep*, moment în care nucleul trece toate procesele care îndeplinesc condiții stabilite, în starea gata de execuție; dintre acestea, numai un singur proces se va lansa în execuție, celelalte fiind trecute din nou în starea de blocare.

Mecanismele de sincronizare prin intermediul evenimentelor privesc sincronizarea prin intermediul semnalelor sau al semafoarelor, respectiv al directivelor de sistem.

### 5.1. Sincronizarea prin intermediul semnalelor.

Așa cum s-a menționat anterior, la îndeplinirea condițiilor prestabilite pentru lansarea în execuție, nucleul trece procesele în starea gata de execuție prin intermediul funcției *wakeup*.

O altă modalitate de sesizare a unor evenimente apărute în sistem o constituie plasarea unor biți de atenție (*lock-bit*) în tabela de procese, nucleul efectuând verificarea activării bitului corespunzător care se asimilează cu emiterea unui semnal la trecerea din modul sistem în modul utilizator înainte, respectiv după blocarea fiecărui proces.

La apariția unui semnal se pot ivi următoarele acțiuni ce se execută în continuare:

- terminarea procesului, când se afișează de regulă, și imaginea procesului în memorie în fișierul *core*;
- acțiuni specifice care sunt desemnate prin funcția:

*Signal (nr\_semnal, nr\_rutină\_tratare\_semnal)*

De remarcat că după tratarea semnalului, se revine la valorile standard ale acțiunilor ce se vor executa ulterior.

Un caz particular îl constituie utilizarea funcției de sistem wait, care permite blocarea unui proces până la terminarea unuia dintre fiii acestuia sau până la recepționarea unui semnal:

***P=wait (&stare)***

Unde:

- *p* reprezintă valoarea returnată a funcției (identificatorul de proces al fiului):
  - - 1 dacă s-a recepționat un semnal, o eroare sau s-a terminat deja execuția fiilor procesului;
  - - numărul de identificare al fiului a cărui execuție s-a terminat;
- *stare* adresa unde se specifică modul de terminare al execuției inițiate de fiu.

Pentru sincronizarea procesului cu un anumit fiu se poate folosi:

***While (wait (&stare)! =pid)***

Care reia ciclic testarea terminării execuției acțiunii fiului cu identificatorul pid, ceea ce evident implică blocarea procesului apelant (tatăl).

## **5.2. Sincronizarea prin semafoare**

Particularitatea sincronizării prin intermediul semafoarelor constă în aceea că se pot executa operații pe o mulțime de semafoare printr-un singur apel de funcție. Structurile de date asociate semafoarelor sunt constituite din:

- identificatorul mulțimii de semafoare;
- descrierea mulțimii de semafoare care include:
  - utilizatorii care au acces la mulțimea de semafoare;
  - numărul de semafoare al mulțimii;
  - momentul ultimei accesări și al ultimei actualizări al mulțimii semafoarelor;
- descrierea semaforului care cuprinde:
  - valoarea semaforului;
  - identificatorul ultimului proces care a accesat semaforul;

- numărul proceselor care așteaptă ca valoarea semaforului să devină mai mare decât valoarea sa curentă;
- numărul proceselor care așteaptă ca valoarea semaforului să devină zero.

Algoritmul de execuție al unei operații pe semafoare se desfășoară astfel:

a) dacă semaforul operației este mai mare ca zero, atunci acesta se va aduna la valoarea semaforului;

b) dacă semaforul operației este zero atunci:

- dacă valoarea semaforului este zero, funcția se termină imediat;
- dacă valoarea semaforului este zero și indicatorii de condiție (*flags*) arată *NOWAIT*, funcția se termină imediat; în caz contrar:
- se incrementează numărul de procese care așteaptă ca valoarea semaforului să fie zero;
- se suspendă procesul curent până când valoarea semaforului devine zero (prin decrementarea numărului de procese ce așteaptă ca valoarea semaforului să fie zero) sau se primește un semnal prin care se efectuează aceeași decrementare și se reia procesul.

C) dacă semaforul operației este mai mic decât zero și:

- dacă valoarea semaforului este mai mare sau egală cu modulul semaforului operației, atunci din valoarea semaforului se va scădea valoarea absolută a semaforului;
- dacă valoarea semaforului este mai mică decât modulul semaforului operației și indicatorii de condiție specifică *NOWAIT*, funcția se va termina imediat; în caz contrar:
- se incrementează numărul proceselor care așteaptă ca valoarea semaforului să devină mai mare ca valoarea sa curentă, suspendându-se procesul curent până când valoarea semaforului devine mai mare sau egală cu valoarea absolută a semaforului operației (când numărul de procese ce așteaptă ca valoarea să devină mai mare ca valoarea curentă se decrementează, iar din valoarea semaforului se va scădea valoarea absolută a semaforului operației sau se primește un semnal (când numărul de procese care așteaptă ca valoarea semaforului să devină mai mare decât valoarea curentă se decrementează, iar procesul suspendat se reia).

### 5.3 Sincronizarea prin directive de sistem

Sincronizarea se poate realiza și printr-o serie de funcții disponibile utilizatorului cum sunt:

***Kill (identificator\_proces, număr\_semnal);***

Care realizează trimiterea unui semnal pentru un proces;

***Pause ();***

Realizează blocarea procesului curent până la primirea unui semnal;

***Allarm (număr\_secunde);***

Trimite un anumit semnal (întrerupere de ceas) după intervalul de timp specificat în număr de secunde.

## Capitolul 6: Comunicația între procese

La sistemele de operare UNIX este necesară realizarea unor modalități practice de comunicație între procese; mecanismele actuale permit alternativele prezentate în continuare.

### 6.1. Comunicația prin fișiere

Comunicația prin intermediul fișierelor oferă avantajul simplității de realizare și o modalitate directă prin care se poate efectua comunicarea. Dezavantajul constă în aceea că este sarcina programatorului să controleze accesul concurent, prin testarea ciclică a momentului eliberării zonei critice.

Exemplu

```
While ((f=creat ("fișier", 0111)) <0) sleep (2);  
    < zona critică >  
Unlink ("fișier");
```

Prin care primul proces care reușește să creeze fișierul, va trece la execuția instrucțiunilor din regiunea critică; până ce nu va executa funcția *unlink*, orice alt apel al funcției de creare a unui fișier (*creat*), va bloca procesul prin apelul funcției *sleep*.

## 6.2. Comunicația prin fișiere de tip pipe

Fișierele de tip *pipe* sunt fișiere obișnuite la care însă operațiile de citire/scriere se realizează într-o ordine prestabilită *FIFO* (*First Input First Output*); aceste fișiere sunt tranzitorii, datele citite într-o manieră strictă a ordinii în care au fost scrise respectând regula de sincronizare producător/consumator (o dată citită dintr-un fișier nu mai poate fi reluată, iar memorarea se face ca la orice fișier utilizând însă numai blocuri adresate direct).

La UNIX există două tipuri de fișiere pipe, fiecare generând anumite particularități ale modului în care se realizează comunicația:

1) **fișiere fără nume** care se crează cu directivă

*Pipe (p)*

Unde *p* reprezintă descriptorul de fișier astfel:

*P [0]* pentru citire

*P [1]* pentru scriere.

De remarcat că în blocurile alocate fișierelor (de regulă numai primele 10 blocuri de disc) alocarea se realizează conform algoritmului Round-Robin.

Scrierea mărește dimensiunea fișierului iar citirea o micșorează; dacă un proces încearcă să scrie un articol ce depășește capacitatea fișierului, nucleul execută transferul până la umplerea acestuia, apoi blochează procesul; când se citește din acest fișier se poate scrie din nou, dar s-ar putea să existe mai multe procese care așteptau scrierea în fișier ducând la apariția competiției.

Comunicația prin fișier pipe fără nume se poate realiza numai între un proces și fiii săi, deoarece fișierul pipe nu se identifică decât prin descriptorul său și intrările în tabela cu fișiere deschise în zona de procese ale utilizatorului, aceasta fiind moștenită de către fiii săi; astfel, zonele de date ale proceselor nu pot fi partajate în mod normal.

2) **Fișierele pipe cu nume** oferă câteva facilități comparativ cu fișierele pipe fără nume:

- se crează într-un director ca orice fișier obișnuit prin directiva `mknod` cu tipul fișierului `010000`;
- la deschiderea sa prin directiva `open`, se poate specifica blocarea procesului ce încearcă să deschidă pentru citire un fișier pipe nedeschis anterior pentru scriere;
- se va șterge explicit prin directiva `unlink`;
- datorită unui nume de identificare este permisă accesarea lui de către orice procese independente.

### 6.3. Comunicarea prin zone de memorie comune

*Zonele de memorie comune* permit proceselor să utilizeze în comun, o zonă de memorie prin transferul unei zone din spațiul virtual al proceselor; accesul este controlat integral de programator.

O *zonă de memorie comună* este definită prin:

- identificator;
- segment de memorie;
- o structură de date ce include:
  - numărul de identificare al utilizatorului care a creat zona, al utilizatorului curent și al grupului;
  - dimensiunea segmentului;
  - identificatorul procesului care a accesat zona de memorie;
  - numărul de procese care partajează această zonă;
  - alte informații.

### 6.4. Comunicarea prin cozi de mesaje

Principiul de comunicație prin intermediul cozilor de mesaje constă în comunicarea prin intermediul unei cozi circulare de mesaje gestionată de nucleul sistemului de operare.

Transmiterea se realizează cu sincronizare, deci procesele ce solicită primiri de mesaje se vor bloca până ce recepționează mesajul solicitat, iar cele ce doresc să emită, se vor bloca dacă nu există nici un proces care să recepționeze iar coada de mesaje este plină.

Mesajele ce se transmit sunt caracterizate printr-un tip (reprezentat printr-un număr la alegere), iar recepția se poate face prin specificarea tipului.

## Capitolul 7: Gestiunea fișierelor speciale

Toate operațiile de intrare/ieșire sunt privite ca operații cu fișiere, Shell nerecunoscând existența unităților periferice, deci Kernel ascunde unitățile fizice reale sub forma unor fișiere aparente.

Fiecare unitate fizică este înzestrată cu un program special (*driver*) care translatează acțiunile solicitate de unitățile virtuale astfel încât utilizatorul poate comunica cu orice unitate atașată sistemului fără nici o deosebire, dacă aceasta dispune de driverul corespunzător.

Deoarece aceste fișiere sunt deosebite de fișierele utilizator, ele sunt referite ca **fișiere speciale** de tip **caracter** sau **bloc**, fiind recunoscute două tipuri de unități periferice (disc și terminal) a căror inod-uri conțin:

- numărul major - indică tipul dispozitivului;
- numărul minor - indică numărul unității.

### 7.1. Fișierele disc

Fișierele disc sunt organizate fizic în blocuri de 512 bytes, fiecare fișier având alocat un număr întreg de blocuri. Blocurile nu trebuie să fie neapărat contigue pe disc, spre deosebire de alte sisteme ce impun acest lucru. Atribuirile standard pentru fișiere implicit sunt:

- 0 fișier standard de intrare - tastatură;
- 1 fișier standard de ieșire - ecranul terminalului;
- 2 ecranul de intrare al terminalului, ce servește gestionării mesajelor de eroare și a informațiilor de diagnostic.

Identificarea fișierelor se realizează prin:

- intrările din directori care identifică fișierele printr-un număr de index asociat fiecărui fișier;
- numărul de index pentru orice volum fizic, este punctator într-o altă tabelă rezidentă pe volum (*index-listă*);
- intrările în lista de indecși ale fiecărui volum, conțin un număr de intrări numite *inod-uri*, motiv pentru care lista de indecși se mai numește tabelă de *inod-uri*;
- un *inod* conține o serie de informații cu privire la fiecare fișier:
  - numărul de identificare al ușor care a creat fișierul;

- starea de protecție a fișierului (*ReadOnly*, *Open*, etc);
- 13 cuvinte ce arată blocurile unității ocupate de fișier;
- dimensiunea fișierului;
- data la care a avut loc ultima modificare;
- de câte ori a fost referit în directori;
- biții de identificare a directorilor, fișierelor speciale și a fișierelor mari.

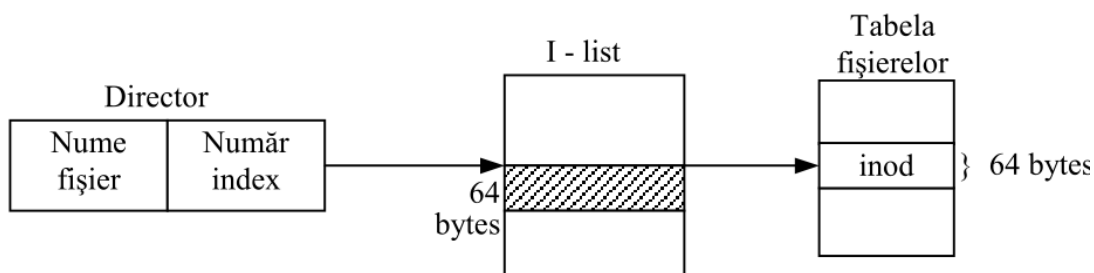


Fig 7.1.1. Identificarea unui fișier pe disc

- Corespondența nume fișier-index este disponibilă prin comanda **ls-i** și se numește legătură (*link*), aceasta servind drept mijloc de identificare și gestiune a fișierului; alte comenzi ce afectează legăturile:
  - *mv* - mută legătura;
  - *ln* - crează o legătură;
  - *rm* - șterge o legătură;

Când se lucrează cu un fișier în program, în RAM există o altă structură ce aparține nucleului (*File Table*) unde se copiază un inod din *i-LIST*, dar zona alocată în memorie va avea 64 B ai inod-ului plus o zonă ce conține informații necesare în timpul lucrului (pointeri de citire/scriere, contor pentru procesele ce folosesc fișierul, etc.).

Ca și în MS-DOS, fișierele se identifică în program printr-un *FD* (*File Descriptor*). Fiecare proces are o zonă de informații în care se păstrează toate *FD* - urile pentru fișierele deschise în timpul lucrului și o zonă de cod a procesului neapartenând nucleului ci în *UOF* (*User's Open File*).

Accesul la sistemul de fișiere se face la trei nivele:

- 1) prin utilitare - nivelul 1;
- 2) prin funcții de bibliotecă <stdio. H> - nivelul 2 care apelează 3);

3) prin apeluri de sistem - nivelul 3: *open, close, read, write* prin care se gestionează fișierele la cel mai scăzut nivel.

Pentru reducerea numărului operațiilor de intrare/ieșire la periferice ce lucrează la nivel de bloc, nucleul organizează în memoria internă o zonă de tampon cache (structuri de date software) diferite de memoria cache.

*Tampoanele cache* sunt organizate în două liste dublu înălțuite (unele putând face parte din ambele); lista tampoanelor în uz este organizată conform unei funcții de dispersie (*hashing*) ce ține seama de numărul logic al dispozitivului periferic și numărul blocului al cărui conținut se încarcă în zona de date atașată tamponului. Funcțiile sunt apelate de nivelul inferior al sistemului de gestiune al fișierelor (SGF).

## Bibliografie

1. Gheorghe D, Oancea B, Vasilescu A, ed. Editura Economica, 2003
2. McMullen J, Advanced UNIX user's handbook, Prentice Hall PTR, 2000
3. Tanenbaum A, Computer Networks, 3<sup>rd</sup> ed, PrenticeHall, 1996
4. [Sistemul de operare UNIX](#) , unibuc.ro
5. [Teaching Unix](#), www.ee.surrey.ac.uk
6. [Kernel \(computing\)](#) , wikipedia.org