

18. FIȘIERE

18.1 Structuri de date externe

Masivele, listele și arborii binari, sunt structuri de date interne, prin faptul că se află în memoria internă a unui sistem de calcul. În momentul în care aceste structuri sunt stocate în memoria externă – pe discuri, benzi magnetice, dischete, compact-discuri acestea sunt numite *structuri de date externe*.

Memoria internă a unui calculator este imaginată ca un șir ordonat de baiți. Fiecare bait se definește prin adresă și prin conținut. Memoria externă este imaginată, de asemenea, ca un șir de baiți, cu deosebirea că pentru calculul adresei sunt luate în considerare elementele specifice modului de structurare a suportului.

Astfel, atunci când suportul este organizat în piste, cilindri și problematica accesului este rezolvată prin poziționarea pe aceste unități de acces, adresarea efectuându-se luând în calcul elemente structurale.

În plus, particularitățile de desfășurare în plan fizic a operațiilor de intrare/ieșire determină analiza distinctă a raportului dintre semnificația instrucțiunilor I/O din programe și operațiile efective de citire/scriere de pe suportul extern.

Presupunând că pentru un sistem de calcul modulele care realizează citiri/scrieri operează cu conținutul unor buffere de lungime L , tot timpul aceste funcții furnizează informații asupra începutului zonei ce este explorată, lungimea acesteia fiind rezultatul logicii de organizare a datelor.

Programatorul are acces direct sau indirect la zona de memorie tampon prin intermediul adresei sale de început sau prin intermediul unei alte zone de memorie definită în programul său, zonă în care este copiat conținutul bufferului.

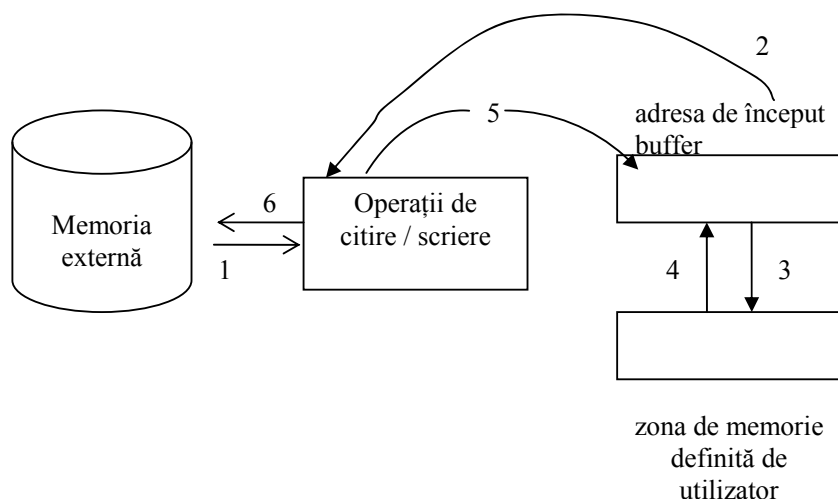


Figura 18.1 Operații de intrare/ieșire

Dinamica operațiilor de intrare/ieșire determină actualizarea variabilei pointer care delimitează partea de început a subzonei din buffer ce este copiată în zona de memorie definită de utilizator.

În cazul funcțiilor de citire/scriere a unui caracter, variabila pointer se modifică cu o unitate, marcând exact schimbul de informații în dialogul om – calculator .

În cazul citirilor/scrierilor cu format delimitatorul acceptat este analizat, iar algoritmul de punere în corespondență este astfel proiectat încât acesta nu este integrat în setul informațiilor.

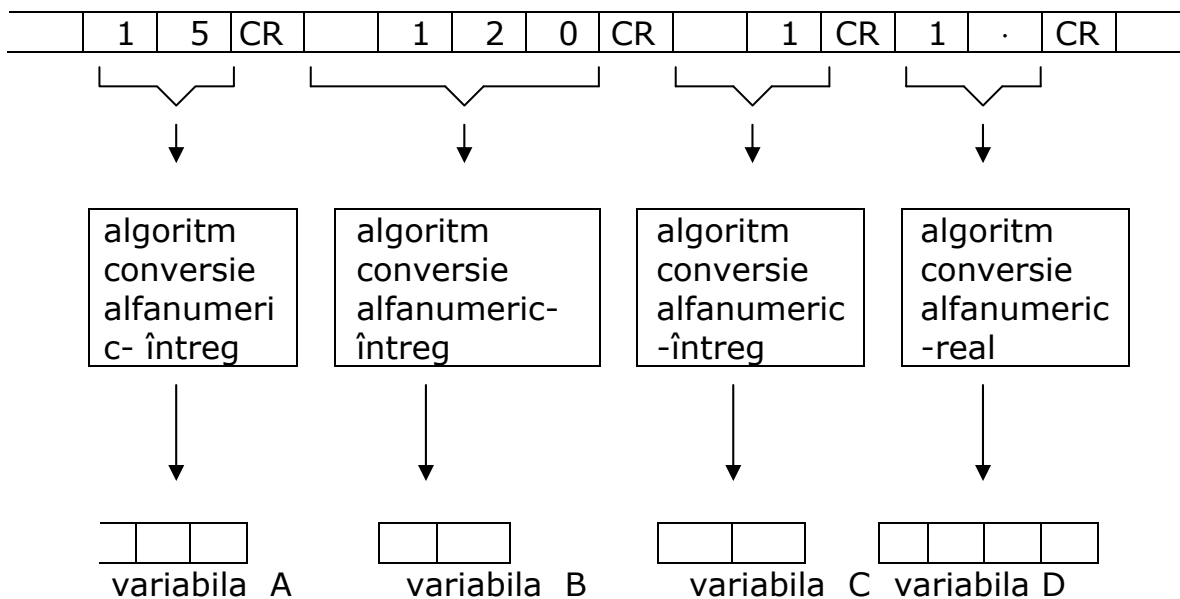


Figura 18.2 Punerea în corespondență a valorilor citite

Șirurile de caractere sunt delimitate prin <CR>, iar algoritmul de parcurgere a bufferului determină un astfel de conținut al variabilei pointer încât este transmis ca parametru funcțiilor de conversie începutul fiecărei subzone de buffer care începe după <CR>.

Modul cum este concretizat <CR> ca delimitator de sfârșit de șir și modul cum este definit descriptorul de format imprimă structurii de parametri ai funcțiilor de conversie anumite particularități.

Se observă că dinamica variabilei pointer este influențată de tipul operației de intrare/ieșire. Acest aspect explică de ce este necesară efectuarea avansului acestuia, când se alternează citiri cu format cu citiri fără format. La operații neomogene există modalități neomogene de definire și de tratare a delimitatorilor de sfârșit de șir ca rezultat al activării tastei <CR>.

Ceea ce pare simplu în cazul structurilor de date interne, nu devine mai complicat în cazul structurilor de date externe, atât timp cât sunt clarificate chestiunile legate de variabilele pointer asociate bufferelor și de faptul că o *citire fizică* efectivă nu înseamnă neapărat o *citire logică*, iar la scriere se întâmplă același lucru. Prin operația logică, înțelegem acțiunea ce corespunde unei apelări de funcție citire/scriere din program.

De exemplu, pentru zona de lungime minimă $L = 256$ octeți, ce este scrisă/citită la o singură operație fizică efectivă pe/de pe suport, la scrierea pe suportul extern a trei variabile de tip articol, A, B și C cu:

$$\begin{aligned}lg(A) &= 120 \text{ octeți} \\lg(B) &= 110 \text{ octeți} \\lg(C) &= 200 \text{ octeți}\end{aligned}$$

variabila pointer permite preluarea datelor din structura *A*, se majorează cu 120, preia datele din structura *B* și realizează o scriere fizică pe suport. Variabila pointer este apoi reinițializată și va prelua datele structurii *C* după care efectuează a doua scriere fizică. În acest caz celor trei *scrieri logice* le-au corespuns două *scrieri fizice*.

Există diferite modalități de realizare a operațiilor de citire/scriere după cum se folosesc sau nu factori de blocare, se definesc sau nu elemente de regăsire a informațiilor.

Structurile de date externe nu diferă mult de structurile de date interne. Apar unele complicații care nu sunt majore de altfel, prin aceea că volumul datelor este foarte mare și elementele repetitive abundă, ceea ce conduce la ideea că structurile de date externe sunt privite ca structuri de date interne dispuse pe suport de memorie externă.

Pentru a realiza regăsirea într-un volum de date de dimensiuni remarcabile, este necesară organizarea, sistematizarea datelor și crearea acelor elemente indispensabile localizării, adresării.

Structurile de date externe sunt contigue, formate din elemente dispuse în continuarea celorlalte și necontigue, distanța dintre elemente fiind o variabilă aleatoare a cărei lege de repartiție este identificabilă, dar care necesită memorarea distanțelor, întrucât nu se construiește un mecanism de generare cu repetare a acestora.

Structurile de date externe se regăsesc în cele mai multe cazuri sub denumirea de *fișiere*. Când ating un nivel de structurare prin adrese suficient de dezvoltat, se formează fișiere interdependente, iar în cazul unor structuri mai complexe se regăsesc sub denumirea de *baze de date*.

În cazul în care conținutul de lungime *L* este tratat distinct, se ia în discuție conceptul de *înregistrare fizică*. Se pornește de la faptul că într-un buffer, de regulă sunt stocate datele ce corespund unei structuri de tip *articol*, ce se recunosc sub denumirea de *înregistrare logică* sau *articol logic*. Scopul este de a face deosebirea între modul în care se dispun informațiile pe suportul fizic și modul în care sunt gândite organizările de date în raport cu prelucrările particulare.

Introducerea factorilor de blocare vine să complice lucrurile, dar să îmbunătățească indicele de folosire a zonelor de memorie.

Dacă:

$$L' = lg(\text{structura de date de tip articol}) < L \quad (18.1)$$

și dacă există:

$$k = \left\lceil \frac{L}{L'} \right\rceil \quad (18.2)$$

unde parantezele drepte înseamnă *partea întreagă* a expresiei, raportul *k* reprezintă o expresie mai simplificată a factorului de blocare.

De exemplu, dacă definim o structură de tip articol, ce conține câmpuri ce conduc la o lungime de 80 octeți și $L = 256$ octeți:

$$k = \left\lceil \frac{256}{80} \right\rceil = [3,2] = 3 \quad (18.3)$$

În cazul în care $k' = 1$, pe cei 256 baiți ai bufferului este încărcat un articol, ce este în fișier. Deci fișierul conține în final n articole fizice și tot n articole logice, gradul de încărcare cu informație utilă fiind:

$$g_1 = \frac{n * 80}{n * 256} * 100 = \frac{10}{32} * 100 \approx 30\% \quad (18.4)$$

În cazul în care $k' = 2$:

$$g_2 = \frac{[n/2] + m * 160}{n * 256} * 100 \quad (18.5)$$

unde:

$$m = \begin{cases} 0, & \text{dacă } n \text{ e număr par} \\ 1, & \text{dacă } n \text{ e număr impar} \end{cases}$$

Pentru $k' = 3$:

$$g_3 = \frac{[n/3] + m * 240}{n * 256} * 100 \quad (18.6)$$

Se observă că:

$$g_3 \geq g_2 \geq g_1 \quad (18.7)$$

Se vorbește de factorul de blocare optim care se stabilește pentru fiecare tip de memorie externă și lungime de structură de date de tip articol, ce urmează a fi memorată în fișier.

În continuare, luând în considerare numai aspectele care țin de modul de stocare a informației, strict dependentă de aplicația programatorului, se fac următoarele specificații:

- se consideră fișierul ca structură de date contiguă dacă informațiile utile sunt dispuse unele în continuarea celorlalte, fără baiți care să le separe;
- se consideră fișierul ca structură de date regulat necontiguă dacă între toate articolele sau între grupuri de articole, având număr fix, există baiți nefolosiți, în același număr; există posibilitatea de a construi o formulă de calcul a adresei articolului k , pornind de la adresa altui articol j ;
- se consideră structuri de date necontigue fișierele ale căror articole sunt dispuse unele față de celelalte, la distanțe care sunt variabile aleatoare.

Trecerea de la memoria internă la memoria externă, ia în considerare modul de organizare al fiecărui suport. Dacă la nivelul memoriei interne aceasta este privită ca un vector, în cazul suporturilor externe de informație organizate pe piste baiții sunt priviți ca având dispunerea asemeni elementelor unei matrice. Linia indică pista pe care se află baitul, iar coloana indică poziția baitului pe pistă.

Organizarea pe sectoare determină luarea în considerare a unei matrice tridimensionale. Pentru fiecare suport realizatorii pun la dispoziție formulele de calcul ale adreselor cu luarea în considerare a elementelor de structură a suportului fizic.

Problema fragmentării informațiilor determină stocarea de date necesare localizării părții ce se continuă într-o altă zonă a suportului.

Pentru simplificarea prezentării se consideră un suport extern S , căruia i se asociază un model matriceal de dispunere a baiților. Baitul b_{ij} reprezintă baitul al j -lea, aflat pe pista i a suportului. Suportul are n piste, iar pe o pistă se află m baiți.

Pentru început se presupune că baiții au aceeași destinație, respectiv de a memora informații utile. Nu există baiți care stochează informații privind structura suportului și modul de ocupare a acestuia.

18.2 Criterii de clasificare a fișierelor

Și în cazul clasificării fișierelor, asemenea datelor interne, există o multitudine de criterii, fiecare fișier fiind clasificat cu unul sau mai multe atribute ce corespund criteriilor de clasificare.

- a) Criteriul lungimii articolelor ce alcătuiesc fișierul le împarte în:
 - fișiere cu articole de lungime fixă – sunt formate din elemente având aceeași lungime; fișierele sunt asemănătoare vectorilor ca structuri de date interne; elementele sunt omogene și sunt dispuse unele în continuarea celorlalte;
 - fișierele cu articole de lungimi diferite, dar cunoscute – se consideră m tipuri de articole, având fiecare lungimea $l_1, l_2, \dots, l_m$; fișierul conține aceste elemente dispuse într-o anumită ordine sau în ordine oarecare; în ultima situație este necesară memorarea de informații care să permită identificarea tipurilor de articole și lungimea acestora;
 - fișiere cu articole de lungime diferită, dar necunoscută – ceea ce se cunoaște este legat de faptul că lungimile articolelor se află cuprinse între două limite: lungimile articolelor sunt variabile aleatoare, aparținând unui interval definit; în mod obligatoriu, primul câmp conține lungimea articolului.
- b) Criteriul informațiilor ce definesc regulile referitoare la dispunerea elementelor, împarte fișierele în:
 - fișiere având ca singur mod de dispunere poziția articolelor;
 - fișiere cu elemente sortate după un câmp numit cheie a articolului, în funcție de care se face reperarea în fișier.
- c) Criteriul informațiilor de localizare a articolelor care alcătuiesc fișierul;
 - fișiere care nu conțin informații asupra poziției articolelor – singura modalitate de a selecta un articol este parcurgerea tuturor articolelor care îl preced;

- fișiere care au definite zone ce conțin informații referitoare la adresele unor grupe și subgrupe de articole – pentru a identifica un anumit element se localizează grupul și apoi subgrupul de articole; o dată identificat subgrupul, selectarea elementului căutat este rezultatul parcurgerii articol de articol până la găsirea respectivului element;
 - fișiere în care elementele conțin informații ce permit conturarea de liste înlănțuite sau arbori pe suport de memorie externă – complexitatea legăturilor dintre articolele fișierului determină un volum de informație privind adresele articolelor cu care un element intră într-o anumită relație.
- d) Criteriul operațiilor ce sunt efectuate de fișiere determină gruparea acestora în:
- fișiere destinate scrierii datelor;
 - fișiere destinate citirii datelor;
 - fișiere destinate efectuării operațiilor de actualizare.
- Oricare dintre fișierele unei grupe, în raport cu scopurile prelucrării își modifică atributele. De exemplu, un fișier care se creează este destinat scrierii datelor. La consultare același fișier aparține grupei a doua, iar dacă înaintea consultării se efectuează actualizări, același fișier aparține celei de a treia grupe.
- e) Criteriul gestiunii fișierelor ia în considerare existența unui sistem de funcții de prelucrare care asigură efectuarea operațiilor specifice lucrului cu fișiere, numit sistem de gestiunea fișierelor. Acest criteriu împarte mulțimea fișierelor în:
- fișiere care sunt prelucrate prin apelarea de funcții ale unui sistem de gestiune – rezolvă întreaga problemă legată de organizarea și accesul la date; funcțiilor sistemului de gestiune a fișierelor le corespund instrucțiuni; parametrii instrucțiunilor devin parametri reali ai funcțiilor sistemului de gestiune; programatorul abordează întreaga problemă numai din punct de vedere logic al rezolvării problemei sale;
 - fișiere pentru care sunt definite funcții primitive de realizare a operațiilor elementare cu datele destinate suporturilor externe – programatorului îi revine sarcina să gestioneze totalitatea informațiilor necesare regăsirii elementelor ce alcătuiesc fișierul; pentru programator, fișierul este o masă amorfă de baiți, având conținut și poziții; el efectuează transferuri de baiți din zone de memorie spre suportul extern, indicând adrese acestor zone, lungimea zonei și un mod de reparare a fișierului; există situații în care se efectuează lucru direct în bufferul asociat fișierului cu care se lucrează; în acest caz fișierul apare ca o resursă iar gestiunea acestei resurse este lăsată integral la dispoziția programatorului;
 - fișiere care se construiesc și se utilizează folosind instrucțiunile limbajului de asamblare – programatorul definește totalitatea condițiilor pentru efectuarea operațiilor cu fișiere; sunt definite și încărcate etichetele care sunt scrise; se definesc bufferele și se gestionează; se calculează adresele de pe suportul extern unde vor fi scrise datele; programatorul preia în programele sale tot ceea ce efectuează funcțiile primitive sau sistemul de gestiune a fișierelor; programul în care apar fișierele gestionate

de programator nu apelează alte funcții decât cele scrise de acesta și folosește numai instrucțiuni ale limbajului de asamblare, cel mult seturi de macroinstrucțiuni.

f) Criteriul organizării fișierelor, ia în considerare existența sau inexistența unor informații de regăsire a datelor, după cum urmează:

- fișierele cu organizarea secvențială - succesiune contiguă de articole fără existența unor elemente de informare, altele decât delimitatorii de început și de sfârșit; dacă se ia în considerare similitudinea acestui mod de organizare cu înregistrarea pe o casetă a șlagărelor unei formații rock, avem imaginea clară a tuturor posibilităților de lucru cu fișierele secvențiale; accesul la un șlagăr presupune audierea celor care-l preced; ștergerea unui șlagăr, altul decât cel de la sfârșit presupune existența a două casete; înregistrarea unui nou șlagăr, se face numai după ultimul șlagăr anterior;
- fișierele însoțite de informații de tip index – presupun sortarea articolelor după chei și împărțirea acestora în submulțimi; punerea în corespondență a elementelor din submulțimi cu adresele fizice determină realizarea într-un fișier a unei mulțimi de subfișiere secvențiale; cu ajutorul indecșilor se identifică subfișierul secvențial și articolul căutat este preluat după ce a fost localizat prin parcurgerea articol după articol a subfișierului; operațiile de actualizare, vizează ștergerea de articole, adăugarea de articole la sfârșitul fișierului sau subfișierelor, modificarea de câmpuri, rescrierea de articole și inserarea de articole; subfișierele sunt masive unidimensionale contigue, formate din articole; ștergerea este o operație de dezactivare a elementelor, fără realizarea deplasării celorlalte elemente cu o poziție spre stânga, deci fizic ștergerea unui articol nu se produce, operația fiind numai la nivel logic, în sensul că accesul la date este precedat de un test asupra caracterului activ sau neactiv al articolului; inserarea de articole, pentru a menține criteriul ordonării elementelor care a determinat împărțirea mulțimii în submulțimi de articole, presupune realizarea de liste înlănțuite; articolul inserat este dispus într-o zonă disponibilă numită folcloric, zona de depășire.
- fișiere ale căror articole sunt dispuse aleator pe suport – cunoscându-se dimensiunile fișierului se alocă o zonă pe suportul extern printr-o operație numită preformare; datele sunt puse în corespondență printr-un procedeu oarecare cu adresele unde vor fi scrise; rezultă că procedeul de punere în corespondență este cu atât mai performant cu cât el realizează o dispunere mai uniformă a articolelor în zona rezervată; există posibilitatea ca folosind algoritmi de randomizare să se obțină elemente ce intră în calculul adresei fizice a începutului zonei din suportul extern în care se scrie un articol; pornind de la imposibilitatea să se genereze numere diferite care să îndeplinească condiția de apartenență la subintervale, pe măsură ce se scriu articole în fișiere, are loc fărâmițarea zonei rezervate; se construiesc funcții pentru gestionarea articolelor care au aceeași adresă fizică prin calcul; aceste fișiere cu

organizarea aleatoare, permit efectuarea întregii game de operații, cu condiția ca mecanismul de obținere a adresei fizice să se mențină neschimbat; fișierul organizat aleator apare ca o structură necontiguă în care fiecare element este localizat pe baza unei formule, având ca parametru valoarea unui câmp ce intră în structura articolului, câmp numit cheia articolului; numeroasele lucrări care prezintă limbajul COBOL, redau imagini sugestive ale modului în care se implementează filozofia fiecărui mod de organizare a fișierelor, deci realitatea este alta, dacă se are în vedere organizarea matriceală a suportului și modul static în multe cazuri de alocare a memoriei externe.

g) Criteriul suportului unde este stocat fișierul împarte fișierele în:

- fișiere pe cartele perforate;
- fișiere pe bandă perforată;
- fișiere pe bandă magnetică;
- fișiere pe suport optic;
- fișiere pe disc;
- fișiere pe dischete;
- fișier în memoria internă.

Prelucrările moderne, neîngrădite de restricțiile lipsei de memorie internă, au condus la citirea unui fișier ca un singur articol foarte mare și prelucrarea acestuia în memorie. Logic apare o singură instrucțiune de citire, deci apelarea o singură dată a funcției în realitate au loc:

$$cf = \left\lceil \frac{L_f}{L} \right\rceil + m \quad (18.8)$$

citiri fizice, unde:

- cd - citiri fizice din fișier;
- L_f - lungimea fișierului, dată în baiți;
- L - lungimea unei înregistrări fizice la o citire;
- m - variabila booleană ce este 0, dacă L_f este divizibil prin L și 1 în caz contrar.

h) Criteriul privind modul de efectuare a operațiilor de intrare/ieșire grupează fișierele în:

- fișiere de tip *stream* în care datele la scriere sau la citire sunt câmpuri elementare sau alte tipuri structuri de date, constituite într-un șir, cu indicarea formatului după care se fac mai întâi convențiile și apoi se stabilesc lungimile șirurilor care vor fi scrise pe suport; fișierul apare ca o succesiune de elemente oarecare ale căror poziții sunt deduse, dacă se iau în calcul informațiile pe care le oferă descriptorii de format și locul pe care fiecare element îl ocupă în lista de parametri ai funcției apelate la scriere; este de dorit ca aceleași liste de variabile de la scriere să fie utilizate și la apelul funcțiilor de citire, iar descriptorii de format să se mențină aceiași pentru a oferi succes prelucrărilor;
- fișiere de tip *record* în care datele sunt caracterizate prin lungime și adresă de început; articolele au lungime fixă sau variabilitatea lungimilor este în cadrul unei mulțimi formate din câteva elemente; operațiile de lucru cu fișierele de acest tip nu

sunt în general precedate sau urmate de conversii; operațiile ce preced calculele sunt lăsate la dispoziția programatorului;

De exemplu, fișierul stocurilor de materiale este:

- un fișier de tip record;
- are organizare indexată;
- se află pe disc;
- este gestionat cu funcțiile unui sistem de gestiune a fișierelor;
- are articole de lungime fixă;
- se efectuează toate operațiile de actualizare pe el;
- conține informații asupra poziției anumitor articole;
- articolele sunt identificate după codul materialului care joacă rol de cheie de articol;
- fiecare material are o cheie unică;
- fișierul este sortat crescător.

Astfel, un fișier oarecare este caracterizat înscriind oricare dintre informațiile ce rezultă din criteriile specificate.

18.3 Fișiere secvențiale

Fie mulțimea de elemente $E_1, E_2, \dots, E_n$ oarecare ce corespund structurilor de date $SD_1, SD_2, \dots, SD_n$ definite într-un program P .

Elementele $E_i, i = 1, 2, \dots, n$, sunt șiruri de baiți caracterizate printr-un conținut rezultat din inițializarea structurilor de date $SD_i, i = 1, 2, \dots, n$. Pe un suport extern se alocă elementelor $E_1, E_2, \dots, E_n$ zone de memorie de lungime:

$$\lg(SD_i) + \lg(\alpha) \quad (18.9)$$

fișierul având în final lungimea:

$$L_f = n * \lg(\alpha) + \sum_{i=1}^n \lg(SD_i) \quad (18.10)$$

De obicei, $\lg(\alpha) = 2$. Cei doi baiți atașați fiecărui element vor conține lungimea articolului:

$$\text{cont}(\alpha) = \lg(SD_i) \quad (18.11)$$

Fișierul este caracterizat printr-o adresă a articolelor. Astfel:

$$\text{adr}(E_i) = A + \lg(\alpha) \quad (18.12)$$

unde, A reprezintă adresa fizică a primului bait a primului articol căruia i s-a atașat în față $\lg(\alpha)$ baiți pentru a memora lungimea articolului respectiv:

$$\text{adr}(E_i) = A + \sum_{j=1}^{i-1} (\text{cont}(\beta_j)) + \lg(\beta_i) \quad (18.13)$$

sau:

$$adr(E_i) = A + \sum_{j=1}^{i-1} \lg(SD_j) + i * \lg(\alpha) \quad (18.14)$$

unde, β_k reprezintă grupul de baiți de lungime $\lg(a)$ atașat elementului E_k în fișier. Dacă:

$$\lg(SD_1) = \lg(SD_2) = \dots = \lg(SD_n) \quad (18.15)$$

atunci:

$$cont(\beta_1) = cont(\beta_2) = \dots = cont(\beta_n) \quad (18.16)$$

caz în care, în eticheta de început a fișierului se specifică tipul articolului *lungime fixă*. De asemenea, se memorează și lungimea, iar $\lg(\beta_k)$ devine zero.

Adresa unui element oarecare E_i , este obținută prin relația:

$$adr(E_i) = A + \sum_{j=1}^{i-1} \lg(SD_j) \quad (18.17)$$

Cu fișierele secvențiale se efectuează următoarele proceduri:

- crearea unui fișier secvențial constă în dispunerea elementelor $E_1, E_2, \dots, E_n$ pe suport în această ordine;
- consultarea integrală sau parțială a fișierului baleiază din aproape în aproape elementele fișierului și se prelucrează cele care prezintă interes;
- adăugarea elementului E_m se efectuează la adresa:

$$adr(E_m) = adr(E_n) + \lg(S_n) + \lg(\alpha) \quad (18.18)$$

ceea ce pune în evidență că adăugarea se face numai la sfârșitul fișierului;

- interclasarea a două fișiere;
- sortarea fișierelor, care constă în obținerea unui fișier în care articolele au o astfel de dispunere încât pentru un câmp x există relația:

$$cont(E_1.x) > cont(E_2.x) > \dots > cont(E_n.x) \quad (18.19)$$

dacă sortarea s-a făcut descrescător, sau:

$$cont(E_1.x) < cont(E_2.x) < \dots < cont(E_n.x) \quad (18.20)$$

dacă sortarea s-a făcut crescător.

Întrucât se lucrează în cele mai multe cazuri cu fișierul în totalitatea lui, nu este justificată memorarea de informații pentru anumite articole.

Explorarea fişierelor secvenţiale corespunde unei structuri de date contigue, asemeni unui vector de structură sau a unei înşirui de diferite structuri, cu posibilitatea calculului adresei unui element oarecare:

$$adr (suc (E_i)) = adr (E_i) + lg (SD_i) + lg (\alpha) \quad (18.21)$$

$$adr (pred (E_i)) = adr (E_i) - lg (SD_{i-1}) - lg (\alpha) \quad (18.22)$$

Se spune că fişierul este închis dacă:

$$\lim_{m \rightarrow \infty} suc^m(E_i) = E_n \quad (18.23)$$

Se spune că fişierul este deschis dacă:

$$\lim_{m \rightarrow \infty} pred^m(E_i) = E_1 \quad (18.24)$$

O astfel de abordare determină continuarea prelucrării chiar dacă există tentative de a închide un fişier deja închis sau de a deschide un fişier deja deschis.

Apare problema privind parametrii funcţiei de deschidere. Dacă sunt luaţi în considerare parametrii primei deschideri, problema este rezolvată, tentativele de deschidere a unor fişiere deja deschise fiind inefective.

Întrucât există formule de calcul pentru adresele fiecărui element din submulţimea $E_1, E_2, \dots, E_n$ nu se justifică construirea unui vector a adreselor în fişier $a_1, a_2, \dots, a_n$ pentru aceste elemente.

Sistemele de operare evaluate gestionează închiderea fişierelor la închiderea execuţiei programelor.

18.4 Fişiere secvenţial – indexate

Se consideră o mulţime de elemente $E_1, E_2, \dots, E_n$ generate după structura SD şi un câmp x astfel încât:

$$cont (E_1.x) < cont (E_2.x) < \dots < cont (E_n.x) \quad (18.25)$$

deci, şirul elementelor este sortat crescător după câmpul x , care joacă rol de cheie a articolelor în structura SD de generare.

Elementele de sortare vor fi memorate, fiecare ocupând o zonă de lungime:

$$lg(SD) + lg(\gamma) \quad (18.26)$$

unde, γ reprezintă o zonă în care este memorată o adresă fizică de pe suport, ținând seama de structurarea contiguă a suportului.

Fişierul are un fond informaţional de lungime:

$$L_f = n * [lg(\gamma) + lg(SD)] \quad (18.27)$$

Se construiește mulțimea perechilor:

$$(cont(E_i.x), a_i) \quad (18.28)$$

a cheilor și adreselor articolelor ce intră în componența fișierului. Se calculează:

$$\sigma^2 = \frac{\sum_{i=1}^n (cont(E_i.x) - \bar{x})^2}{n} \quad (18.29)$$

și rezultă o împrăștiere suficient de mare care să reducă șansele găsirii unei formule de calcul a adresei fizice a unui articol, folosind strict ca informație cheia acestuia.

Dacă se acceptă ideea utilizării unui arbore binar cu 4 noduri terminale, mulțimea elementelor $E_1, E_2, \dots, E_n$ este împărțită în 4 subșiruri, având un număr aproape identic de elemente, reținând adrese și cheile elementelor ce delimitează fiecare subșir de articole:

$$(cont(E_{kj}.x), a_{kj}), \quad j = 1, 2, 3, 4; k = 1, 2, \dots, n \quad (18.30)$$

unde $k_1 < k_2 < k_3 < k_4$.

Perechile de delimitare ale începutului de subșir, se obțin astfel:

- pentru primul subșir:

$$Q_1 = (cont(E_1.x), a_1) \quad (18.31)$$

- pentru al doilea subșir:

$$Q_2 = (cont(E_{1+m}.x), a_{m+1}) \quad (18.32)$$

- pentru al treilea subșir:

$$Q_3 = (cont(E_{1+2m}.x), a_{2m+1}) \quad (18.33)$$

- pentru al patrulea subșir:

$$Q_4 = (cont(E_{1+3m}.x), a_{3m+1}) \quad (18.34)$$

Perechile pentru delimitarea sfârșitului pentru fiecare din cele 4 subșiruri sunt:

$$\begin{aligned} P_1 &= (cont(E_m.x), a_m) \\ P_2 &= (cont(E_{2m}.x), a_{2m}) \\ P_3 &= (cont(E_{3m}.x), a_{3m}) \\ P_4 &= (cont(E_n.x), a_n) \end{aligned} \quad (18.35)$$

S-a considerat că fiecare subsir are m elemente, cu excepția ultimului șir care are 1, 2 sau 3 elemente mai puține.

Setului de date i se asociază arborele binar:

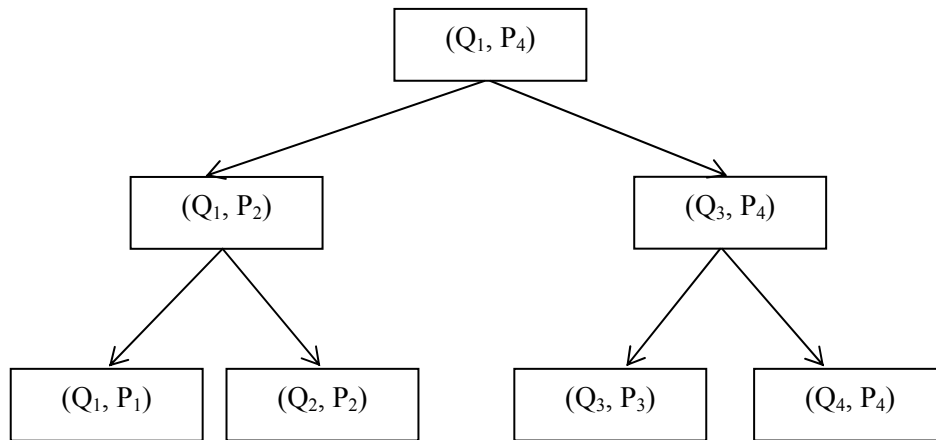


Figura 18.3 Structura setului de date i

Dacă, de exemplu, fișierul sortat este organizat secvențial și se dorește citirea articolului penultim, care are cheia 7233, în mod normal trebuie să fie citite cele $n-2$ articole care îl preced. Dacă însă fișierul este înzestrat cu această informație pe care o conține arborele binar cu 4 noduri terminale, inspectarea nodului rădăcină permite vizualizarea faptului că articolul căutat cu cheia 7233 este în fișier, adică:

$$\text{cont}(E_1.x) \leq 7233 \leq \text{cont}(E_n.x) \quad (18.36)$$

Întrucât:

$$7233 > \frac{\text{cont}(E_n.x) + \text{cont}(E_1.x)}{2} \quad (18.37)$$

rezultă că se parcurge pe nivelul inferior, nodul din dreapta.

Întrucât:

$$7233 > \frac{\text{cont}(E_{2n+1}.x) + \text{cont}(E_n.x)}{2} \quad (18.38)$$

rezultă că se parcurge pe nivelul inferior, nodul din dreapta.

Odată ajungând pe ultimul nivel al arborescenței, se rețin adresele a_{3m+1} și a_n și se baleiază secvențial subfișierul delimitat astfel. Deci, se vor baleia mai puțin $\frac{1}{4}$ din totalul elementelor fișierului.

Pentru construirea arborilor binari asociați există numeroși algoritmi de partiționare a fondului informațional. Important este ca numărul căutărilor secvențiale să fie cât mai redus. Trebuie realizat un echilibru între numărul de nivele ale arborelui binar și numărul subfișierelor, deoarece creșterea numărului de subfișiere conduce la creșterea duratei de acces la acestea din cauza parcurgerii unui număr prea mare de noduri în arborele binar.

În cazul în care se dorește inserarea unui articol E_j într-un astfel de fișier se identifică poziția lui, astfel încât:

$$\text{cont}(E_k.x) < \text{cont}(E_j.x) < \dots < \text{cont}(E_{k+1}.x) \quad (18.39)$$

În acest caz, articolul ca atare este memorat într-o zonă rezervată a fișierului, legătura cu celelalte articole efectuându-se prin:

$$\begin{aligned} \text{cont}(\gamma_k) &= \text{adr}(E_j) \\ \text{cont}(\gamma_j) &= \text{adr}(E_{k+1}) \end{aligned} \quad (18.40)$$

ce corespunde unei inserări de elemente într-o listă.

La un număr mare de inserări, parcurgerea listelor reduce performanța programului de exploatare a fișierului secvențial – indexat. Acestea justifică trecerea la reorganizarea fișierului. Prin reorganizare se înțelege crearea unui nou fișier, în care elementele se dispun într-o aceeași zonă fără a mai exista informații de legătură, ca în cazul în care ar fi dispersate pe suport.

La reorganizare, se construiește un nou binar al indecșilor. Utilizarea arborilor binari are aici numai un caracter exemplificativ. Tipul de arbore depinde în principal de împrăștierea cheilor în intervalul pe care sunt definite. Informațiile privind arborele asociat tablei de indecși se stochează pe suport și face parte din fișier.

Pentru o parcurgere mai rapidă, în cazul în care este posibil, informațiile aferente structurii arborescente a indecșilor se încarcă în memoria internă și se lucrează cu ele folosind funcțiile de parcurgere ale unui arbore.

În limbajele precum C și C++, fiecare programator implementează algoritmi proprii pentru organizarea secvențial – indexată, iar prin comparația comportamentului lor statistic cu alți algoritmi existenți, sunt dezvoltați sau abandonați.

18.5 Fișiere aleatoare

Sunt acele fișiere care ocupă o anumită zonă din memoria externă și ale căror elemente nu sunt reperate decât prin adresa fizică pe care o au pe suportul extern.

Întrucât abordarea fișierelor *random*, depășește prin amploarea fundamentării acest cadru, se consideră oportună prezentarea unui caz particular de fișiere ale căror adrese ale elementelor se calculează în funcție de poziția pe care acestea o au, tot astfel cum se procedează în cazul masivelor unidimensionale.

Se consideră funcția $\text{poz}(E_m)$ care indică poziția elementului E_m în șirul de articole.

Astfel, dacă mulțimea de elemente $E_1, E_2, \dots, E_n$ are exact n elemente:

$$1 \leq \text{poz}(E_j) \leq n \quad (18.41)$$

$$\text{poz}(\text{succ}(E_j)) = \text{poz}(E_j) + 1 \quad (18.42)$$

$$poz(pred(E_j)) = poz(E_j) - 1 \quad (18.43)$$

Din cele două relații rezultă că:

$$poz(succ(E_j)) - poz(pred(E_j)) = 2 \quad (18.44)$$

sau:

$$poz(E_j) = \frac{poz(succ(E_j)) + poz(pred(E_j))}{2} \quad (18.45)$$

Cu aceste definiții:

$$adr(E_j) = A + (poz(E_j) - 1) * lg(E_j) \quad (18.46)$$

se construiește șirul de perechi:

$$(cont(E_j.x), poz(E_j)) \quad (18.47)$$

care permite, în cazul în care se stabilește o relație de forma:

$$poz(E_j) = \rho(cont(E_j.x)) \quad (18.48)$$

regăsirea elementului după poziție pe care o are în fișier.

În cazul în care nu se identifică funcția $\rho()$, elementele $cont(E_j.x)$, $j = 1, 2, \dots, n$, vor fi memorate într-o structură omogenă unidimensională și poziția elementului în care este memorată această informație, corespunde poziției articolului în fișier.

18.6 Fișiere interdependente

Interdependența fișierelor, este privită sub două aspecte și anume: interdependența la nivelul articolelor unui fișier și, respectiv, interdependența la nivelul a două sau mai multe fișiere.

Se consideră două fișiere:

- fișierul *PRODUSE*, ale cărui elemente sunt generate după structura *PROD*, cu câmpurile *cod produs*, *denumire produs*, *stoc*, *unitate de măsură*, *număr materii prime* care intră în componența sa și materiile prime, gama de operații;
- fișierul *MATERIALE*, ale cărui elemente sunt generate după structura *MAT*, cu câmpurile *cod material*, *denumire material*, *unitate de măsură*, *cantitate existentă în stoc*, *preț unitar*.

Dacă dorim să aflăm costul materiilor prime necesare realizării unui produs, citim din fișierul *PRODUSE* articolul corespunzător respectivului produs și rând pe rând preluăm codurile materialelor ce intră în componența sa. Pentru fiecare material ce intră în componența produsului, citim câte un articol din fișierul *MATERIALE*.

Să ne imaginăm că fișierele sunt secvențiale. Rezolvarea este extrem de greoaie, pentru că accesul la materiale este obținut numai prin baleierea de la început a fișierului *MATERIALE*, dacă materialele nu au fost memorate

în ordine crescătoare după codurile acestora în articolul din fișierul *PRODUSE* și dacă fișierul *MATERIALE* nu este sortat.

Problema devine eficientă dacă, în locul codurilor materialelor, dispunem de adresele articolelor în fișierul *MATERIALE*, dar reorganizarea fișierului de materiale, ar atrage după sine actualizarea adreselor pentru materiale din componența fișierului *PRODUSE*.

Dacă se lucrează cu fișiere indexat – secvențiale, rezultatul este bun, iar dacă se folosesc pozițiile materialelor dintr-o listă, performanța devine și mai bună.

Acest tip de dependență este unidirecțională între două fișiere. Numărul de legături dintre un articol din fișierul *PRODUSE* și fișierul *MATERIALE* este variabil, strict dependent de numărul de materiale ce intră în componența produsului cu care articolul din fișierul *PRODUSE* este pus în corespondență.

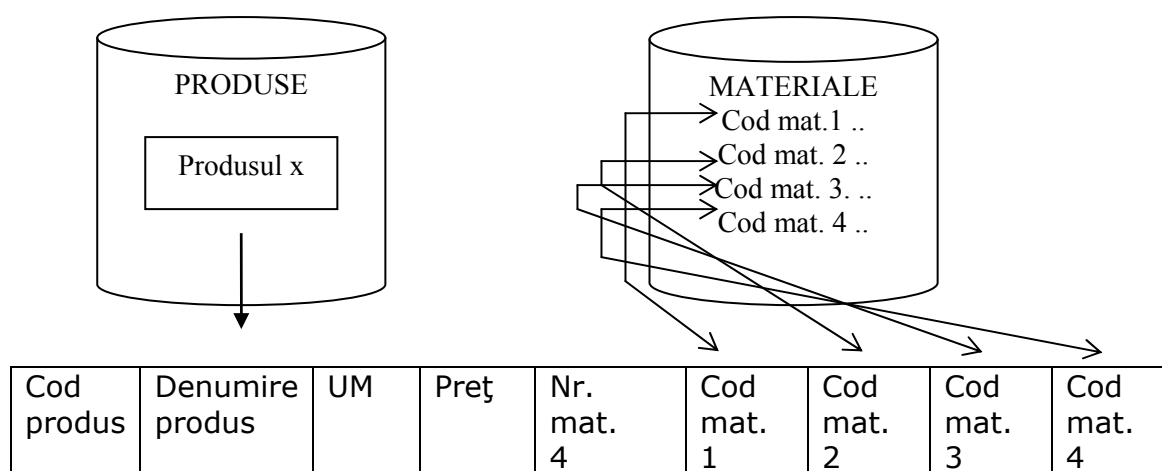


Figura 18.4 Legătura între fișierele *PRODUSE* și *MATERIALE*

Este de dorit ca fișierele *PRODUSE* și *MATERIALE* să se realizeze într-o primă etapă cu codurile materialelor și lângă ele să fie rezervate câmpuri pentru memorarea adreselor fizice ale articolelor ce corespund în fișierul *MATERIALE* respectivelor coduri.

În a doua fază, un sistem de programe identifică poziția fiecărui articol și încarcă în fișierul *PRODUSE* în zonele disponibile de adresă chiar adresa articolului al cărui cod se află în imediata sa vecinătate.

Reorganizarea fișierelor pe fondul definirii codurilor nu implică decât reîncărcarea de adrese, lucru ce se efectuează automat.

Fișierele interdependente privesc legături într-un singur sens, de la fișierul de produse către fișierul de materiale. Se construiesc și legături în sensul opus, pentru fiecare articol al fișierului de materiale, indicându-se în care dintre produse, este folosit respectivul material. Este o informație necesară, dar care lungeste mult articolul *MAT* și de aceea, în acest caz este puțin utilizată o astfel de abordare.

Dacă totuși se dorește și o legătură dinspre fișierul *MATERIALE* către fișierul *PRODUSE*, realizarea se efectuează în același fel, parcurgându-se tot doi pași.

În final se vor obține legăturile dintre articole, ilustrate în figura 18.5.

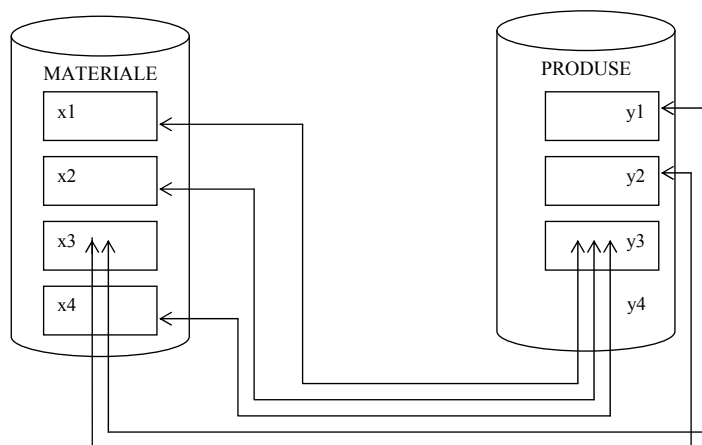


Figura 18.5 Legăturile dintre articolele fișierelor *MATERIALE* și *PRODUSE*

Rezultă că produsul y_3 , utilizează materialele x_1 , x_2 și x_4 , iar materialul x_3 este prezent în compoziția produselor y_1 și y_2 .

Articolele fiecărui fișier rămân în continuare independente unele de celelalte. Ele sunt prelucrate separat, singurele combinații fiind cele legate de posibilitatea de a permite calculul necesarului de materiale, pentru realizarea de k produse de tipul y_3 și de a stabili dacă stocurile materialelor x_1 , x_2 , x_4 sunt suficiente.

La serviciul aprovizionare, prin parcurgerea fișierului *MATERIALE*, se observă care sunt viitoarele cereri din materialul x_3 .

Se observă că o astfel de construcție este limitativă, dar cu o serie de artificii de consultare a celor două fișiere se obțin și alte regrupări de date.

O privire de ansamblu pune în evidență că păstrarea articolelor din fiecare fișier ca entități independente reduce diversitatea combinațiilor de date, care pentru un manager competent, reprezintă o sursă sigură de fundamentare a deciziilor.

18.7 Fișiere înlănțuite

Apar întrebările: structurile dinamice cu multe elemente sunt memorate pe suport extern? Se implementează mecanisme de înlănțuire în memoria externă? Cum se generează legăturile exprimate prin adrese dintre elementele deja alocate dacă se specifică numai cheile de identificare a acestora?

Răspunsurile nu sunt simple, dar au fost date cu mulți ani în urmă, obținându-se funcții de creare și prelucrare a fișierelor înlănțuite.

Se consideră că pentru realizarea unui produs, este necesară parcurgerea unor etape, efectuarea unor operații de prelucrare într-o anumită succesiune, numită gama de operații.

Descrierea unei operații, conține informații precum:

- codul operației;
- denumirea operației;
- durata de execuție;
- calificarea cerută;
- tipul de meseriaș care o execută;
- utilajul necesar;
- scule necesare;

- plata pentru efectuarea operației;
- operația precedentă;
- operația următoare;

La execuția programului de încărcare a fișierului gamei de operații, se introduc rând pe rând, datele ce urmăresc șablonul descris mai sus. Pentru primul articol al unei game, operația precedentă este *NULL*, iar pentru ultimul articol din gamă, operația următoare este de asemenea *NULL*.

Sistemul de creare a fișierului cu acest tip de înlănțuire adaugă două câmpuri ce vor fi inițializate cu adresele ce corespund poziției articolelor pentru operația precedentă și pentru operația următoare din gama de operații.

Câte game de operații sunt definite într-un proces de producție, atâtea liste dublu înlănțuite vor fi realizate.

Prin parcurgerea unui fișier cu articole înlănțuite, se determină care este necesarul de specializări și care sunt salariile ce trebuie plătite, dacă sunt realizate k produse de tipul x , pentru care este necesară efectuarea operațiilor din gama de operații cu un cod specificat.

Fișierul gamei de operații, va avea articole ce corespund gamelor $G_1, G_2, \dots, G_n$, care la rândul lor conțin operațiile $O_{11}, O_{12}, \dots, O_{k1}, \dots, O_{21}, O_{22}, \dots, O_{k2}, \dots, O_{n1}, O_{n2}, \dots, O_{kn}$.

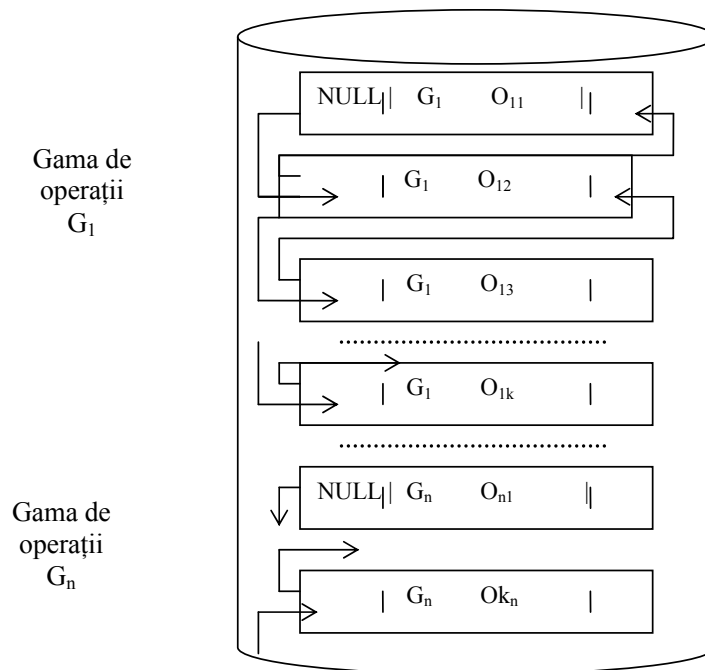


Figura 18.6 Structura fișierului gamei de operații

O altă situație, este aceea corespunzătoare memorării pe suport extern a structurilor arborescente.

Se consideră că la o întreprindere se realizează N produse $P_1, P_2, \dots, P_N$. Fiecărui produs i se asociază o structură arborescentă, deci vor exista N noduri rădăcină și N noduri terminale.

$$M = \sum_{i=1}^N K_i \quad (18.49)$$

unde K_i reprezintă numărul de materii prime, repere sau subansamble nerealizate în întreprindere și care intră în componența produsului P_i .

Se pune problema memorării celor N arborescente, cu condiția realizării cel puțin a unui nivel de redundanță controlat, dacă minimizarea este dificil de realizat sau chiar nu este dorită.

Nodurile reprezintă produse, subansamble, repere sau materii prime, despre care trebuie cunoscut:

- cod de recunoaștere, respectiv cheia;
- denumire;
- unitate de măsură;
- cantitate necesară, respectiv greutate specifică;
- preț unitar;
- gama de operații;
- cod rețetă de fabricație;
- caracteristici de performanță standard;

De asemenea, se impune stabilirea poziției în arborescență, indicând nodul părinte, nodul descendent și nodul vecin din dreapta.

Această multitudine de date, se grupează în două categorii:

- date pentru descrierea componentei ce corespunde unui nod;
- date pentru descrierea poziției nodului în arborescență.

Celor două categorii de date le vor corespunde fișiere descriptive și fișiere de legături.

Fișierele descriptive, au articole formate din două tipuri de date:

- date ce se completează de către utilizatori cu acele elemente ce caracterizează un produs, un ansamblu, subansamblu, reper sau materie primă;
- date ce se completează de un sistem de gestiune a fișierelor înlanțuite și care conțin adrese de articole sau identificatori de marcare a finalului înlanțuirii.

De exemplu, pentru produsul A care este format prin asamblarea reperelor B , C , D și E în ordinea menționată în figura 18.7 fișierul descriptiv, conține 3 articole ce corespund elementelor A , B , C , D și E , iar fișierul de structură descrie legăturile dintre elementele A , B , C , D și E , conform arcelor ce definesc gradul de tip arbore. Deci, acest fișier are articolele: (A ; B), (A ; C), (C ; D), (C ; E).

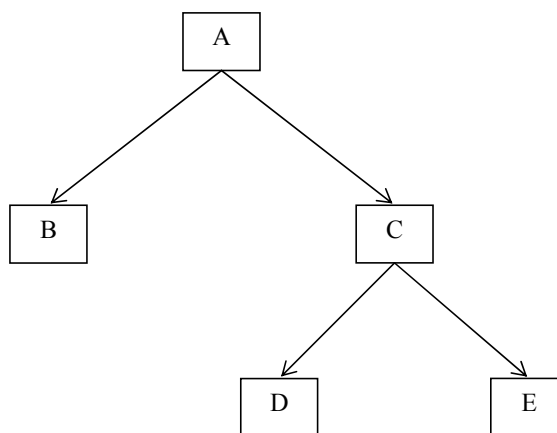


Figura 18.7 Structura componentelor produsului A

Indicarea acestor articole nu necesită o anumită ordine. Analiza perechilor este aceea care permite identificarea tuturor elementelor necesare completării câmpurilor de adresă din fișierul descriptiv.

Astfel, într-o pereche $(x; y)$ rezultă că pentru nodul x , nodul y este descendent, iar pentru nodul y nodul x este părinte.

Nodul A este părinte pentru nodurile B și C , dar el nu apare în nici una dintre perechi ca descendent, deci A este nod rădăcină. În articolul din fișierul descriptiv ce corespunde produsului A , un câmp este completat cu *NULL*.

Nodul B nu apare ca părinte în nici o altă pereche, deci nu are descendenți. Se completează și un câmp din articolul fișierului descriptiv ce corespunde reperului B .

Pentru regăsirea rapidă a informațiilor, înainte de a completa câmpurile din fișierul descriptiv, vor fi completate în fișierul de legătură, câmpurile ce corespund adreselor fizice pe care le ocupă articolele fișierului descriptiv.

Astfel, legăturii corespunzătoare perechii $(A; B)$ i se asociază în fișierul de legătură, articolul:

cod A	cod B	adr (articol A)	adr (articol B)
-------	-------	-----------------	-----------------

Figura 18.8 Structura articolului în fișierul de legătură

Dacă este analizat un produs, problema nu diferă prea mult cu modul de reprezentare a structurilor dinamice în memoria internă. Dacă însă fișierul descriptiv conține totalitatea nodurilor ce descriu cele N produse care se realizează, iar fișierul de legături are atâtea elemente câte arce au cele N arborescențe după care se descompun produsele $P_1, P_2, \dots, P_n$, avem o imagine mai exactă a ceea ce înseamnă reprezentarea în memoria externă a datelor ce formează articole interdependente și care în final, dau fișierele înlanțuite.

Obținerea controlului asupra redundanței revine la a deplasa informații privind adresele din articolele fișierului descriptiv în fișierul de legătură. Dacă se minimizează redundanța, în sensul că fișierul de produse conține repere și materii prime descrise o singură dată, reconstituirea arborescenței se obține prin construirea tuturor adreselor de regăsire nu în articolele din fișierul descriptiv, ci în articolele de legătură.

Problema traversării unui arbore se menține din punct de vedere al semnificației, dar în cazul fișierelor înlanțuite revine la a utiliza, alternativ, informații din fișierul de legături și din fișierul descriptiv.

Operațiile care sunt specifice structurilor arborescente construite în memoria internă sunt definite și în cazul arborescențelor de pe mediile externe. Modificările vizează pointerii, al căror conținut reflectă adrese ale suportului extern.

Funcțiile specifice pentru lucru cu fișiere înlanțuite au ca parametri toate elementele necesare ștergerii unui nod, modificării unor legături sau a unor câmpuri. Adresele care apar la descrierea legăturilor dintre noduri permit parcurgerea arborelui de la rădăcină spre bază sau de la orice nod către bază sau de la orice nod către rădăcină.

Când se proiectează o anumită structură, nu numai listă sau arbore, trebuie ca programatorul să se familiarizeze cu ideea că oricărui arc îi corespund două adrese. El trebuie să găsească algoritmi pentru încărcarea

în câmpuri, pe care are obligația să le definească, a acestor două adrese. Programatorul este obligat să cunoască funcțiile care întorc adresa fizică a începutului zonal de memorie pe care o ocupă un anumit articol.

Dacă sistemele de gestiune a fișierelor înlănțuite realizează aceste operații pentru structuri arborescente sau pentru liste, cazuri particulare de arborescențe, în cazul structurilor oarecare programatorul trebuie să-și construiască funcții proprii. El are grijă să rezerve zone pentru adrese, își definește structura de date adecvată descrierii structurii produsului. De fiecare dată, va avea grijă ca ceea ce realizează să nu genereze ambiguități sau să conducă la reprezentări ce nu concordă cu structura pe care dorește să o implementeze în fișiere.

18.8 Fișierele bazei de date

Din punctul de vedere al modului de structurare a datelor, bazele de date reprezintă o formă mai generală de reprezentare a datelor, care reflectă caracteristici ale elementelor omogene ce definesc mai multe mulțimi.

Fie mulțimile $M_1, M_2, \dots, M_k$, formate fiecare din $n_1, n_2, \dots, n_k$ elemente, așa fel alese încât caracterizarea completă a unui element $x_j \in (M_1)$ este efectivă dacă sunt prezentate datele $d_{j1}, d_{j2}, \dots, d_{jk}$, unde $d_{jk} \in (M_k)$.

În plus, fiecare mulțime are elementele structurare după o anumită regulă. Punerea în corespondență a elementelor celor k mulțimi, conduc în numeroase situații, ca unui element din mulțimea M_i să-i corespundă mai multe elemente din mulțimea M_j sau mai multor elemente din mulțimea M_i să le corespundă un singur element din mulțimea M_h .

Se observă că necesitatea de a separa informațiile în fișiere distincte este dată în principal din dorința de a diminua redundanța dintr-un fișier, pe de o parte, și pentru a permite noi facilități de exploatare a structurilor de date, pe de altă parte.

În continuare, se consideră un exemplu clasic, foarte frecvent utilizat în descrierea principiilor și fundamentelor de realizare a bazelor de date.

Fie mulțimea M_1 a persoanelor dintr-o localitate, care se află într-una din relațiile:

x este vecin cu y
x este fiul lui y
x este tatăl lui y
x este soția lui y
x este soțul lui y

Un individ al colectivității, este descris prin:

- nume;
- adresă;
- raport cu alți indivizi, respectiv vecini, rude;
- tip automobil/marca;
- culoare automobil,
- an cumpărare;
- loc de muncă.

Fie mulțimea M_2 mulțimea automobilelor, în care elementele se află în relațiile:

marca x este produsă de y
marca x are capacitate z
capacitatea z are consumul specific w

Un autoturism este descris prin:

- nume producător;
- model;
- culoare;
- an fabricație;
- capacitate;
- tip combustibil;
- caracteristici motor;
- consum la 100 km.

Fie M_3 mulțimea locurilor de muncă, formată din elemente caracterizate prin:

- nume instituție;
- capital social;
- tip instituție;
- număr salariați;
- nume compartiment;
- nume meserii acceptate la compartimentul respectiv;
- nume lucrători din compartiment.

Fie M_4 mulțimea impozitelor care se aplică mijloacelor de transport, formată din elementele:

- limita inferioară a capacității cilindrice;
- limita superioară a capacității cilindrice;
- impozit;
- taxa CASCO pentru primii 3 ani de funcționare;
- taxa CASCO pentru mașinile cu vechime cuprinsă între 4 – 10 ani;
- taxa CASCO pentru mașinile cu vechime mai mare de 10 ani.

Deși se discută foarte mult, cea mai costisitoare etapă din activitatea de manipulare a bazelor de date o reprezintă încărcarea acestora.

În cazul de față, datele nu suferă o uzură morală rapidă pentru că vizează indivizii dintr-un cartier sau bloc. În cazul în care fenomenul are o dinamică accelerată, se observă că la terminarea încărcării 60 – 80% din date sunt uzate moral și baza de date practic este inoperantă.

În cazul considerat, cele patru mulțimi conduc la date ce alimentează baza de date și se trage concluzia că în continuare ea este complet încărcată.

Pentru a vedea puterea de prelucrare pe care software-ul bazei de date o are, exemplificăm câteva cereri:

- listarea locurilor de muncă ale membrilor familiei lui x;
- listarea proprietarilor de autoturisme cu culoarea z;
- listarea tuturor celor care lucrează la locul de muncă w și au autoturisme pentru care plătesc taxa CASCO cuprinsă în intervalul [a, b];
- există o relație între capitalul social al firmelor și uzura morală a autoturismelor pe care le au salariații lor?

- listarea proprietarilor care plătesc un impozit mai mare decât C lei și taxa CASCO mai mare decât D lei;
- există meseriași de tipul z care au mașini de tipul u absolut noi și care au vecini în aceeași situație?

Dacă toate datele despre un individ, sunt înregistrate într-o structură cuprinzătoare, care conține elementele de descriere ale celor patru mulțimi, se observă o multitudine de repetări, ceea ce conduce de fapt la regruparea informațiilor. Cele 4 mulțimi nu sunt un dat, ci sunt rezultatul unei analize a modului în care sunt sistematizate și concentrate datele.

Pentru elementele din mulțimea M_1 se asociază arborescența:

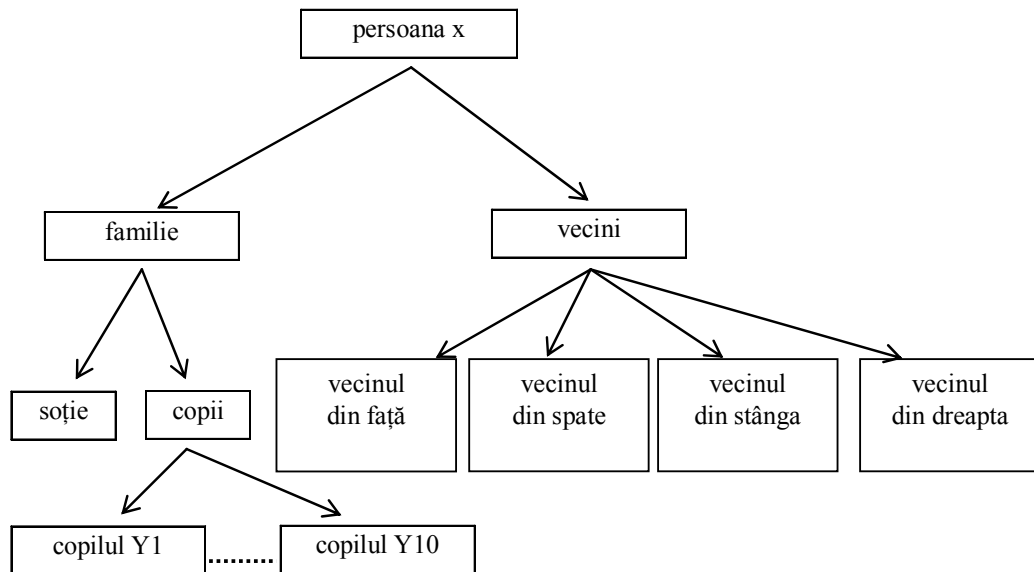


Figura 18.9 Arborescența asociată elementelor mulțimii M_1

Pentru elementele mulțimii M_2 se asociază arborescența din figura 18.10.

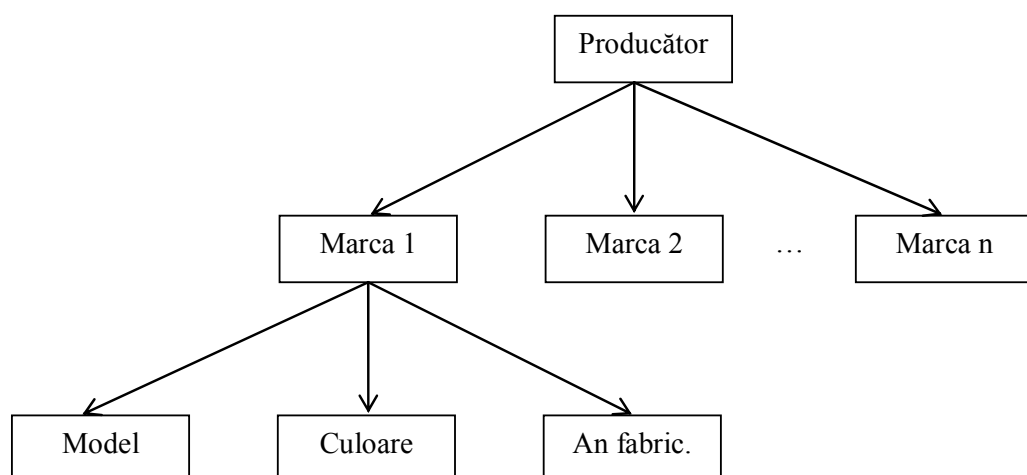


Figura 18.10 Arborescența asociată elementelor mulțimii M_2

Și pentru elementele mulțimii M_3 și M_4 se asociază, de asemenea, arborescențe.

Deci, dacă fiecare mulțime este concretizată în câte un fișier, articolele în cadrul fiecărui fișier sunt legate între ele, funcție de structura arborescenței asociate.

Baza de date presupune atât legături între articolele fiecărui fișier, cât și legături între fișiere.

Se ridică întrebarea: pentru a răspunde la toate solicitările, este necesară constituirea unor structuri de adrese; toate structurile de adrese sunt create de la început sau acestea se creează pe măsură ce necesitățile de prelucrare impun acest lucru?

În fișierul persoanelor, câmpul corespunzător autoturismului conține și adresa articolului ce corespunde mărcii și culorii autoturismului.

Dacă interesează listarea proprietarilor de mașini de culoare z se, procedează astfel:

- se construiește șirul:

$$S = (s_1, s_2 \dots s_n) \quad (18.50)$$

unde s_i este adresa articolului pentru descrierea autoturismului al cărui proprietar este persoana a_i din fișierul de persoane.

- din fișierul autoturismelor se extrage subșirul:

$$P = (p_1, p_2, \dots p_n) \quad (18.51)$$

al adreselor articolelor ce corespund autoturismelor de culoare z :

$$P \cap S = (s_{k1}, s_{k2}, \dots s_{km}) \quad (18.52)$$

ceea ce înseamnă că persoanele ale căror înregistrări ocupă pozițiile $k_1, k_2, \dots k_m$ cu mașini de culoare z și afișarea:

$$\text{cont}(\text{ref}(k_i) . \text{nume}), i = 1, 2, \dots, m \quad (18.53)$$

conduce la tabelarea numelor de persoane care posedă mașini de culoare z .

Interogarea unei baze de date revine la constituirea mai multor șiruri, din fiecare mulțime câte unul. Numărul de șiruri maxim este de regulă egal cu numărul mulțimilor de articole definite.

Astfel, dacă se dorește listarea tuturor celor care au locul de muncă w și au autoturism pentru care plătesc o taxă CASCO cuprinsă în intervalul $[a, b]$, se procedează astfel:

- se construiește șirul adreselor persoanelor care au locul de muncă w :

$$(\gamma_j = \text{const}(w.\text{adresa_persoana}_j)), j = 1, 2, \dots n \quad (18.54)$$

- se construiește șirul capacităților cilindrice:

$$(C_1, C_2, \dots C_h) \quad (18.55)$$

al autoturismelor, însoțite de adresele articolelor din fișierul asociat mulțimii M_2 , ce corespund autoturismelor pentru capacitățile identificate:

$$m_1 m_2 \dots m_e \quad (18.56)$$

De regulă, $e \geq h$, întrucât există mai multe mărci de mașini cu aceeași capacitate cilindrică.

Prin analiza conținutului fișierului M_4 , rezultă că autoturismele pentru care se plătește o taxă CASCO cuprinsă în intervalul $[a, b]$, sunt cele care au capacitățile cuprinse între:

$$[\alpha_1, \beta_1] \quad (18.57)$$

$$[\alpha_2, \beta_2] \quad (18.58)$$

Aceste limite conduc la filtrarea elementelor mulțimii $C_1, C_2, \dots, C_h$, astfel încât rezultă submulțimea adreselor:

$$m' = \{m_1', m_2', \dots, m_p'\} \quad (18.59)$$

de regulă $p < e$.

Se construiește șirul de adrese:

$$\delta_u = (\gamma_i / 0 \leq 1 < k; \text{cont}(\gamma_i.\text{adresa_autoturism}) \in m') \quad (18.60)$$

Scrierea elementelor șirului:

$$(\text{cont}(\gamma_u \rightarrow \text{nume_persoana})), \quad u = 1, 2, \dots, w \quad (18.61)$$

rezolvă cererea formulată inițial.

În exemplul dat, s-a procedat la utilizarea de variabile pointer. Pentru a face deosebire între pointerii folosiți la referirea elementelor din memoria internă și pentru a evita confuziile ce apar folosind conceptul de pointer spre fișier, întrucât este deja concentrat tipul de dată *FILE*. În continuare se utilizează conceptul de pointer extern.

Fiind dat un suport extern al căror baiți au adresa cuprinsă între valorile $[A, B] \cap N$, $A < B$ și $A, B \in N$, variabila v se spune că este pointer extern dacă și numai dacă:

$$\text{cont}(v) \in [A, B] \cap N \quad (18.62)$$

Poziționarea citirii este de fapt atribuirea sau modificarea conținutului unei variabile de tip pointer extern, așa fel încât ea să conțină adresa baitului de unde începe citirea/scrierea.

În enunțul problemei, M_1, M_2, M_3 și M_4 reprezintă fișiere, adică variabile pointer care sunt inițializate cu adresele baiților de unde începe alocarea memoriei externe pentru fiecare din cele patru fișiere. Dacă se are în vedere că fișierul are și o etichetă de început de lungime, $\text{cont}(M_i) + \alpha$ reprezintă adresa primului articol din fișierul M_i .

Presupunând că fișierului M_i i se alocă o zonă contiguă, cuprinsă între adresele $[A_i, B_i]$, rezultă că:

$$\text{cont}(M_i) = A_i \quad (18.63)$$

$$\text{cont}(\delta_i) = B_i - \lg(SD_i) \quad (18.64)$$

unde δ_i reprezintă adresa ultimului element de structură SD_i al fișierului M_i care are o etichetă de sfârșit de fișier de lungime β .

Revenind la fișierele interdependente, fișierele înlănțuite, se observă că structurile de date ce se asociază fișierelor, pe lângă informațiile utile date de programator, se definesc încă multe câmpuri de tip pointer extern care asigură traversările prin structura externă pe timpul execuției.

Folosind definirea pointerilor extern, locul de muncă w este identificat prin pointerul extern pw ce corespunde adresei articolului respectivului loc de muncă.

Mulțimea $m_1, m_2, \dots, m_e$, reprezintă valori ale variabilei pointer extern r , asociat fișierului M_2 :

$$m = \{ r_j / \text{cont}(r_j \rightarrow \text{capacitate}) \in \{C_1, C_2, \dots, C_h\}, i < j < nrr(M_2) \} \quad (18.65)$$

unde $nrr()$ este o funcție care definește numărul de articole existent într-un fișier:

$$nrr : \{F_1, F_2, \dots, F_k\} \rightarrow N \quad (18.66)$$

unde F este mulțimea fișierelor, date prin valoarea inițială a pointerilor lor externi:

$$m' = \{ m_i / \text{cont}(m_i \rightarrow \text{capacitate}) \in [\alpha_1, \beta_1] \cup [\alpha_2, \beta_2] \} \quad (18.67)$$

Selectarea elementelor reprezintă construirea de șiruri de adrese și apoi prin operații de reuniune, intersecție, se obțin șirurile de elemente care prezintă rezolvarea problemei.

Utilizarea bazelor de date reprezintă un domeniu al informației aplicate. Construirea de sisteme de gestiune a bazelor de date este o preocupare de mare importanță pentru toți realizatorii de software, iar elementele prezentate definesc doar o serie de trăsături generale ale filozofiei bazelor de date.

Fiecărui mecanism particular îi corespund reguli precise de definire, inițializare și modificare a pointerilor externi ce se asociază fiecărui articol.

De menționat că în spatele oricărei cereri de informație furnizată de utilizatorul bazei de date, se află pointeri externi, care în unele cazuri sunt deja inițializați și se folosesc ca atare, iar în alte cazuri, sunt numai definiți și își încarcă conținutul în funcție de cerere.

Există situații când pentru a rezolva o problemă, se definesc noi pointeri externi și se activează proceduri de inițializare, în concordanță cu problema de rezolvat. În acest caz se creează noi legături între elemente, cu noi posibilități de selectare a informațiilor din baza de date.

Statistic, diversitatea de fișiere face ca utilizarea unora să fie în anumite cazuri mai eficiente decât a altora. Experiența fiecărui programator este cea care hotărăște tipul de fișier cu care se lucrează pentru fiecare aplicație.

La alegerea tipului de fișier cu care se lucrează, concură o serie de factori din care se enumără:

- numărul de elemente care alcătuiesc fișierul;
- tipul prelucrărilor: integrale, după o cheie, prin selecție;
- frecvența și tipul operațiilor de actualizare;
- durata impusă de prelucrare și necesitatea sortării datelor;

- timpul disponibil pentru elaborarea programelor;
- costuri de elaborare și utilizare;
- sistemul de calcul disponibil;
- sistemele de gestiune a fișierelor cunoscute și disponibile.

Alegerea tipului de fișier și comportamentul algoritmilor implementați pentru regăsirea informațiilor, sunt rezultatul unei analize statistice a datelor, înregistrate prin urmărirea în timp a comportamentului la prelucrare.