

28. STRUCTURI DE DATE ȘI ACCES ÎN BAZE DE DATE

28.1 Stocarea datelor

O bază de date utilizează mai multe dispozitive de stocare a datelor. Aceste dispozitive, numite și memorii, se deosebesc prin capacitatea lor de păstrare, viteza lor, modul de accesare a datelor (secvențial sau direct) și, în sfârșit, prin persistența lor. *Memoriile volatile* își pierd conținutul când sistemul este întrerupt de la sursa de alimentare. *Memoriile nevolatile*, cum sunt discurile sau benzile magnetice își păstrează conținutul chiar și când sunt decuplate de la sursa de curent electric.

Dispozitive de stocare a datelor

În general, cu cât o memorie e mai rapidă cu atât ea este mai scumpă și, prin urmare, cu atât capacitatea ei de stocare este mai redusă. Memoriile utilizate de un sistem de gestiune a bazelor de date (SGBD) constituie o ierarhie, figura 28.1, care pornește de la cea mai mică, dar mai eficace, la memoria mai voluminoasă, dar mai lentă:

- memoria cache este utilizată de procesor pentru stocarea datelor și instrucțiunilor;
- memoria principală constituie spațiul de lucru al mașinii; datele sau programele sunt încărcate în memoria principală, unde este posibilă tratarea lor de către procesor;
- discurile magnetice constituie principalul periferic de tip memorie; ele oferă o capacitate mare de stocare și permit un acces relativ eficient la citire și scriere;
- benzile magnetice sunt dispozitive ieftine, dar viteza joasă de lucru face ca să fie folosite pentru fișiere de salvare.

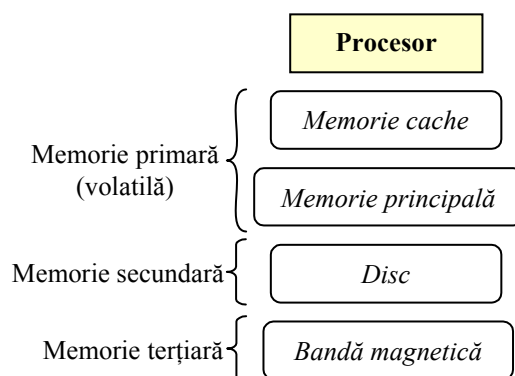


Figura 28.1 Ierarhia de memorii

Memoria primară, care constă din memoria cache și memoria principală, asigură un acces foarte rapid la date. Actualmente, costul unui volum de memorie principală este de 100 ori mai mare decât același volum de memorie pe disc, iar banda magnetică este și mai ieftină. Dispozitivele de stocare lentă, cum sunt discurile, joacă un rol important în bazele de date, deoarece volumul de date este, de obicei, foarte mare.

Există și alt raționament de păstrare a datelor în memoria secundară sau terțiară. În sistemele cu 32 de octeți adresare, numai 2^{32} octeți pot fi referiți direct în memoria principală, or volumul de date este mult mai mare.

O bază de date este aproape întotdeauna păstrată pe disc și din considerente de persistență. Cu toate acestea, datele recuperate sunt plasate în memoria principală pentru a fi prelucrate. Pornind de la această realitate, unde un mic fragment din baza de date poate sta în memoria centrală, SGBD-ul trebuie, deci, să efectueze, în permanență, transferuri de date între memoria principală și memoria secundară. Costul acestor transferuri influențează direct performanța sistemului.

Funcționarea discului magnetic

Un disc este o suprafață circulară, magnetizată, capabilă să înregistreze date. Suprafața magnetizată poate fi situată pe o singură parte (dacă discul are o singură față) sau pe două părți ale discului (dacă discul are față dublă).

Discurile sunt divizate în *sectoare*, un sector constituind cea mai mică suprafață de adresare. Cu alte cuvinte, se pot citi sau scrie zone noi care încep pe un sector și acoperă un număr întreg de sectoare. Mărimea unui sector este, de obicei, de 512 octeți.

Dispozitivul. În mod obișnuit, bazele de date sunt stocate pe disc și datele sunt transferate de pe disc în memoria principală în măsura necesităților. Pentru a limita costul dispozitivului și mări capacitatea de stocare, mai multe discuri sunt montate pe o axă și formează un pachet de discuri. În lucru, axa și discurile sunt antrenate într-o mișcare de rotație cu o viteză mare, figura 28.2.

Cea mai mică unitate de date stocată pe un disc este un bit care poate avea valoarea 0 sau 1. Biții sunt grupați câte 8 pentru a forma octeți, iar o mulțime de octeți formează, pe suprafața unui disc, o cunună circulară numită *pistă*. Mulțimea tuturor pistelor cu același diametru se numește *cilindru*.

Există o mulțime de capete de lectură/scriere situate la sfârșit pe un braț. Capetele se mișcă în grup, astfel că acestea pot fi poziționate asupra tuturor pistelor ce constituie un cilindru. Prin urmare, toate datele unui cilindru pot fi accesate fără a se produce vreo mișcare a brațului, care este o operație mai lentă. Noțiunea de cilindru corespunde, deci, tuturor datelor disponibile fără a avea nevoie de o deplasare a capetelor de lectură.

Capul de lectură nu este antrenat în mișcarea de rotație. El se deplasează pe un plan fixat care îi permite de a se apropia sau de a se îndepărta de axa de rotație a discului și de a accesa o pistă.

Sunt tot atâtea capete de lectură câte discuri sunt (de două ori mai mult, dacă discurile sunt cu fețe duble) și toate capetele sunt poziționate respectiv pe planul lor de deplasare. În orice moment, pistele accesibile sunt cele de pe un cilindru, ceea ce constituie o constrângere de care trebuie să se țină cont când se dorește optimizarea amplasării datelor.

În fine, ultimul element al dispozitivului este *controlorul* care servește drept interfață între pachetul de discuri și sistem. Controlorul primește, de la sistem, cererile de citire sau scriere și le transformă în mișcări corespunzătoare ale brațului cu capete de lectură.

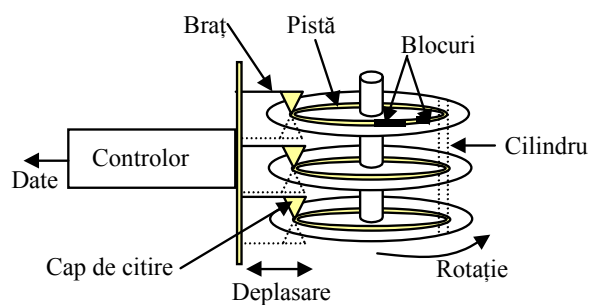


Figura 28.2 Un disc magnetic

Fiecare pistă este divizată în arce numite *sectoare*, mărimea cărora este o caracteristică a discului. Sectoarele sunt diviziuni fizice și, de aceea, au dimensiuni fixe. Ele sunt numerotate astfel încât sectoarele consecutive au numere consecutive. Numărul minim de octeți citiți de capul de lectură este definit de mărimea sectorului, în general, de 512 octeți.

Fiecare pistă este, deci, divizată în *blocuri* (sau *pagini*), care constituie unitatea de schimb dintre disc și memoria principală. Blocul este o diviziune logică și, prin urmare, dimensiunea lui este o mărime variabilă. Dimensiunea unui bloc poate fi setată când discul se inițializează. Atunci se fixează unitatea de intrare/ieșire (adică blocul) mai mare (sau egală) decât a unui sector. Lungimea unui bloc este, de obicei, egală cu un multiplu al lungimii unui sector. Astfel, se obțin *blocuri* a căror lungime (tipică) este de 512 octeți (un sector), 1024 de octeți (2 sectoare) sau 4096 de octeți (8 sectoare).

Toate lecturile sau toate scrierile pe disc se efectuează prin blocuri. Astfel, chiar dacă o lectură nu se referă decât la 4 octeți, întreg blocul este transmis în memoria centrală. Această caracteristică este fundamentală pentru organizarea datelor pe disc. Unul din obiectivele SGBD-ului este de a face totul, încât, atunci când este necesară citirea unui bloc de 4096 de octeți pentru a accesa un număr de 4 octeți (4092 de octeți constituind restul blocului), aceasta să se petreacă într-un timp scurt, adică blocul, deja, trebuie să se găsească încărcat în memoria centrală. Această motivație stă la baza mecanismului de *regrupare*, în special, la baza structurilor de date indexate și dispersate.

Accesul la date. Un disc este o memorie cu acces direct. Spre deosebire de o bandă magnetică, de exemplu, este posibil de accesat o unitate de date situată în orice loc pe disc, fără a avea de parcurs secvențial tot suportul. Accesul direct se bazează pe o adresă dată fiecărui bloc în momentul inițializării discului de către sistemul de exploatare. În general, această adresă este compusă din trei elemente:

- numărul discului în pachet sau numărul suprafeței, dacă discurile sunt cu suprafață dublă;
- numărul pistei;
- numărul blocului pe pistă.

Citirea unui bloc, dată fiind adresa lui, se petrece în trei etape:

- poziționarea capului de lectură pe pista care conține blocul;
- rotirea discului și așteptarea până în momentul când blocul va trece sub capul de lectură – capetele sunt fixate, discul se rotește;

- transferul blocului.

Astfel, există trei factori care afectează direct viteza cu care datele sunt transferate între disc și memoria principală:

- timpul de căutare;
- timpul de rotație sau așteptare;
- timpul de transfer.

Timpul de căutare (sau *poziționare*) este timpul necesar pentru a mișca brațul cu capetele de lectură/scriere din poziția curentă, până în poziția cilindrului adresat. *Timpul de rotație* sau al *stării latente*) (este timpul necesar discului pentru a se învârti până ce capul va fi situat deasupra sectorului de scriere sau citire. *Timpul de transfer* este timpul necesar pentru citirea sau scrierea datelor. Acest timp depinde de numărul de octeți transferați. Timpul de transfer este neglijabil pentru un bloc, dar poate deveni important când trebuie citite mii de blocuri. Mecanismul scrierii este asemănător celui de citire, dar poate să consume puțin mai mult timp, în cazul când controlorul verifică dacă scrierea se efectuează corect.

Optimizarea accesului la date

Acum când se cunoaște cum funcționează un disc, este destul de evident că, pentru același volum de date, timpul de lectură poate varia considerabil, în funcție de factori precum amplasarea datelor pe disc, ordinea instrucțiunilor de lectură/scriere sau prezența datelor într-o memorie *cache*.

Toate tehnicile care permit reducerea timpului consumat pentru accesul la disc sunt intensiv utilizate de SGBD-uri, a căror performanță, în mare parte, este condiționată de eficiența acestor accesări. În această secțiune, vor fi descrise principalele tehnici de optimizare, realizate într-o arhitectură simplă, constituită dintr-un singur disc și un singur procesor.

Regruparea datelor. Din cele relatate mai sus, reiese că timpul de executare a operațiilor asupra datelor în baza de date este afectat semnificativ de faptul cum sunt stocate datele pe disc. Timpul de transfer al blocurilor, din sau spre disc, de obicei, domină timpul consumat de operațiile bazei de date. Pentru a minimiza acest timp, este necesar de a alege o strategie de amplasare a datelor pe disc, ținând cont atât de geometria discului, cât și de mecanica discului.

De exemplu, fie SGBD-ul trebuie să citească 5 secvențe de caractere a câte 1000 octeți fiecare. Dacă un bloc are lungimea egală cu 4096 octeți, două blocuri sunt suficiente pentru a păstra aceste secvențe de caractere. În figura 28.3, sunt prezentate două tipuri de organizare a unui disc. În primul tip, fiecare secvență este plasată într-un bloc aparte și blocurile sunt repartizate în mod aleatoriu pe piste discului. În cel de-al doilea tip de organizare, secvențele sunt stocate în două blocuri consecutive pe aceeași pistă a discului.

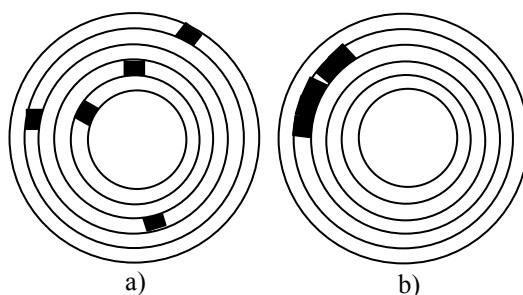


Figura 28.3 a) - organizare rea; b) - organizare bună

Performanțele obținute reprezintă un raport de 1 la 5. Timpul minimal se obține, apelând la gruparea datelor. Gruparea se bazează pe plasarea în același bloc a datelor care au șanse mari de a fi citite împreună. Criteriile folosite la determinarea grupurilor de date constituie bazele structurilor de date în memoria secundară.

În general, câștigul obținut la transferarea a două date este cu atât mai impunător cu cât datele sunt mai „aproprate” pe disc. Această „apropiere” are mai multe valențe.

Evident, apropierea maximal posibilă este obținută când două unități de date se găsesc în același bloc – ele vor fi citite împreună. Două blocuri sunt foarte „aproprate”, dacă sunt stocate consecutiv pe o pistă. Cu viteza discului, două blocuri vor fi citite sau scrise, dacă se găsesc pe aceeași pistă, sub același cap activ de citire. Toate datele de pe o pistă (pe discurile proiectate azi) se citesc și se scriu într-o revoluție a discului, fără a se deplasa capul de citire.

După ce o pistă este citită sau scrisă, alt cap al discului devine activ și altă pistă a aceluiași cilindru este citită sau scrisă. Acest proces continuă până când toate piste de pe cilindrul curent sunt citite sau scrise, și apoi brațul în întregime se deplasează (înainte sau înapoi) spre cilindrul adiacent. Astfel arată noțiunea „apropierea” blocurilor care într-un sens exprimă noțiunile de bloc „următor” și „precedent”.

Prin urmare, în descreșterea ordinii, două unități de date pe disc trebuie să fie în același bloc, pe aceeași pistă, pe același cilindru sau pe cilindri adiacenți.

Exploatarea noțiunii de „apropiere” și aranjare a datelor astfel ca ele să fie scrise sau citite secvențial este foarte importantă pentru reducerea timpului consumat la accesarea discului. Ea permite *lecturi secvențiale* care, după cum reiese din exemplul anterior, sunt mult mai performante decât lecturile aleatorii, deoarece ele evită deplasarea capului de lectură și minimizează timpul de căutare și stare latentă.

SGBD-urile de azi optimizează distanța dintre date în momentul stocării lor pe disc. De exemplu, o relație este stocată pe aceeași pistă sau, în caz că ea ocupă, pe mai multe piste, pe piste de același cilindru, pentru a putea realiza o parcurgere secvențială eficientă.

Pentru ca SGBD-ul să poată efectua aceste optimizări, administratorul, când definește schema fizică a bazei de date, trebuie să rezerve un spațiu important pe disc, pe care el (SGBD-ul) îl va gestiona. Dacă SGBD-ul se mulțumește de a cere sistemului de exploatare spațiu pe disc atunci când are nevoie, stocarea fizică obținută poate fi extrem de fragmentată.

Reordonarea accesului la date. Bazele de date sunt sisteme ce lucrează în regim de multiutilizator. Astfel, chiar dacă teoretic se admite că un fișier este stocat continuu pe aceeași pistă, citirea secvențială a acestui fișier poate fi segmentată de interogările formulate de alți utilizatori care accesează baza de date în mod concurențial. Și atunci, eficiența organizării optime a datelor pe disc, bineînțeles, scade. Apare, deci, problema ridicării eficienței accesului la date în situația satisfacerii simultane a mai multor utilizatori, cererile cărora trebuie administrate concomitent.

De exemplu, dacă un utilizator U_1 cere lectura fișierului F_1 , în timp ce utilizatorul U_2 cere lectura fișierului F_2 , sistemul, probabil, va alterna lectura blocurilor din aceste fișiere. Chiar și în cazul când ambele fișiere vor fi stocate secvențial, deplasarea capului de citire va minimiza într-o câțva avantajele acestei organizări.

SGBD-ul poate reduce acest dezavantaj prin conservarea temporară a operațiilor de citire/scriere într-o zonă tampon (*cache*) și reorganizarea (*secvențială*) ordinii de acces.

Astfel, fie mai mulți utilizatori formulează instrucțiuni de lectură și scriere asupra fișierelor stocate în baza de date. Evident că multe din aceste cereri se pot suprapune când trec prin controlor. Pentru a evita accesul aleatoriu care survine în urma acestei suprapunerii, cererile de acces sunt stocate temporar într-un *tampon*. Sistemul le triază, atunci, pe piste, apoi pe blocuri în cadrul fiecărei piste și apoi transmite lista ordonată controlorului de disc.

O metodă de sistematizare a acestei strategii este așa-zisul „ascensor”. Adică se presupune că capul de lectură se deplasează de la marginea suprafeței discului spre axa de rotație, apoi revine de la axă spre bord. Deplasarea se efectuează pistă după pistă și, la fiecare pistă, sistemul transmite controlorului cererea de citire/scriere corespunzătoare pistei curente.

Acest algoritm reduce la maximum timpul de deplasare a capului, deoarece căutarea se realizează sistematic pe piste adiacente. El este, în particular, eficient pentru sistemele cu multe cereri, fiecare antrenând mai multe blocuri de date. Dar, bineînțeles că pot fi și unele efecte nedorite în cazul unor cereri mari consumatoare de date. Procesele care cer blocuri de pe pista 1, când capul, deja, trece pe pista 2, trebuie să aștepte un timp considerabil pentru a vedea cererea dată satisfăcută.

Utilizarea memoriei tampon. Utilizarea *memoriei tampon* sau a *buferului* pentru optimizarea accesului la date este pe larg practică în toate SGBD-urile. O memorie tampon este o mulțime de blocuri în memoria principală, care sunt copii ale unor blocuri de pe disc. Când sistemul cere accesul la un bloc, mai întâi se inspectează memoria tampon. Dacă blocul se găsește în bufer, este evitată o citire a dispozitivului secundar de stocare a datelor. Dacă nu, se efectuează lectura și se stochează blocurile în memoria tampon.

Ideea este, deci, de a păstra în memoria principală o copie a unei părți cât se poate de mare a bazei de date, chiar dacă o parte din blocurile plasate în memoria tampon nu este utilă la moment. O latură importantă a administrării unei baze de date o constituie specificarea unei părți de memorie principală în calitate de memorie tampon disponibilă, în permanență, SGBD-ului. În plus, această memorie este utilă și prin faptul că

se poate conserva o parte semnificativă a bazei, câștigând, astfel, în performanță.

Dacă în memoria tampon rămâne loc neutilizat, se poate recurge la *lecturi în avans*. Tehnica lecturii în avans este utilizată frecvent de SGBD-uri pentru a efectua operația de joncțiune a două relații. Această tehnică se folosește pe larg și în lucrul cu fișierele indexate.

28.2 Concepte generale de organizare a fișierelor

O bază de date este concepută ca o mulțime de fișiere stocate pe un suport persistent. Înainte de a vorbi despre conceptele fundamentale ale structurilor de fișiere în baze de date, trebuie menționat, pentru a elimina confuzia, că aceste două concepte (*fișier* și *fișier în baza de date*) sunt două noțiuni distincte. În primul rând, fișierele administrate de un SGBD sunt un pic mai structurate. Dar, de fapt, există trei particularități esențiale care caracterizează fișierele din baze de date. Aceste particularități sunt:

- viziuni diferite ale acelorași date;
- independența date-prelucrare;
- redundanța (gestionabilă) datelor.

De aceea, fără excepție, SGBD-rile au propriile lor module de gestiune a fișierelor și a memoriei cache.

Câmpuri și înregistrări

Când se construiesc structuri de fișiere, datelor li se dă un aspect organizațional și de persistență în același timp, adică o aplicație creează, în memoria centrală, date și le salvează într-un fișier și cel puțin o altă aplicație poate citi și rescrie aceste date în memorie, recuperându-le din fișier. Astfel, obiectivul este organizarea datelor într-o structură comprehensibilă de către ființele umane. Or, pentru un sistem de operare, fișierul este o succesiune de octeți repartizați în unu sau mai multe blocuri.

Prin urmare, se disting nivelul logic și nivelul fizic de concepere a fișierului:

- structura logică este forma fișierului în care este văzută și manipulată de aplicații; cum datele sunt organizate în aplicații (în general, organizate în acord cu obiectele pe care aplicația le manipulează); de exemplu, un fișier poate fi văzut ca o colecție de entități ordonate de o cheie sau o structură ierarhică construită din entități principale și entități subordonate;
- structura fizică este o viziune care reflectă reprezentarea și organizarea datelor în mediul de stocare (sectoare, blocuri,...); este forma în care datele sunt stocate, organizate pe dispozitiv, ținând cont de unitatea de bază pe care dispozitivul o poate manipula (o comandă de citire sau de scriere o manipulează în întregime).

La nivel fizic, fișierele sunt constituite din *înregistrări* (*records*) care reprezintă, din punct de vedere fizic, entitățile de lucru ale SGBD-ului. Conform modelului logic al SGBD-ului, aceste entități pot fi tupluri într-o relație sau obiecte. În cele ce urmează, este tratat primul caz, adică este considerat modelul relațional de date.

Câmpuri cu lungime fixă și lungime variabilă. Un tuplu al unei relații este constituit dintr-o listă de componente (atribute), fiecare având un tip. Acestui tuplu îi corespunde, la nivel fizic, o înregistrare, constituită din *câmpuri (field)*. Fiecare tip de atribut determină lungimea câmpului necesar pentru stocarea unei instanțe a câmpului.

Lungimea unui tuplu este egală cu suma lungimilor câmpurilor care reprezintă attributele sale. În practică, lucrurile sunt ceva mai complicate. Câmpurile (și, deci, înregistrările) pot fi de lungimi variabile. Dacă lungimea unei înregistrări de mărime variabilă crește pe parcursul actualizării, trebuie să se găsească un spațiu liber. De asemenea, apare și problema reprezentării valorii *NULL*.

Tipurile de date pot fi divizate în două categorii: tipuri care pot fi reprezentate printr-un câmp de lungime fixă și tipuri care au lungime variabilă. De exemplu, standardul SQL2 propune, printre altele, două tipuri de date pentru secvențe de caractere: *CHAR* și *VARCHAR*.

Tipul *CHAR* indică o secvență de lungime fixă. Astfel, *CHAR(5)* definește un câmp stocat pe 5 octeți. Atunci apare întrebarea: cum se reprezintă valoarea 'Joc'? Există două soluții:

- ultimele două caractere se completează cu spații;
- ultimele două caractere se completează cu un caracter convențional.

Convenția adoptată influențează compararea, fiindcă, într-un caz, se stochează 'Joc' (cu 2 spații), iar în alt caz - 'Joc' fără caractere de terminare. Dacă se utilizează tipul *CHAR*, este important să se studieze convenția adoptată de SGBD-ul concret.

Tipul *VARCHAR(n)* permite stocarea secvențelor de lungime variabilă. Există (cel puțin) două posibilități:

- câmpul este de lungimea $n+1$, primul octet conține un întreg care indică lungimea exactă a secvenței; dacă se stochează 'Joc' într-un *VARCHAR(10)*, se obține atunci '3Joc', unde primul octet păstrează un 3 în formă binară, urmat de trei octeți cu caracterele 'J', 'o' și 'c', iar următorii 7 octeți rămân neutilizați;
- câmpul are lungimea $l+1$, unde $l < n$, aici nu sunt octeții neutilizați, ceea ce permite economisirea spațiului.

De notat că reprezentarea unui întreg de un octet limitează lungimea maximală a unui *VARCHAR* cu 255. O variantă care poate depăși această limită constă în înlocuirea octetului inițial care indică lungimea cu un caracter de terminare a secvenței (fie un C).

Antetul înregistrării. La fel cum se prefixează un câmp de lungime variabilă cu lungimea sa utilă, în antetul înregistrării se stochează unele date complementare. Aceste date pot fi:

- lungimea înregistrării, dacă lungimea este variabilă;
- un pointer spre schema relației, pentru a ști care este tipul înregistrării;
- data ultimei actualizări etc.

Acest antet, de asemenea, se poate utiliza pentru indicarea valorilor *NULL*. Absența valorii pentru unul dintre atribute este o problemă delicată. Dacă nu se stochează nimic, există riscul să fie perturbată decuparea unui câmp, în timp ce, dacă se stochează o valoare convențională, se pierde spațiu. O soluție posibilă este o mască de biți, câte unul pentru fiecare câmp al înregistrării. Bitul ia valoarea 0, dacă câmpul este *NULL* și 1 dacă nu. Această mască poate fi stocată în *antetul înregistrării* și, deci, nu este

necesar spațiu pentru valoarea *NULL*. Totul rămâne în decodarea corectă a secvenței de octeți.

De exemplu, fie o schemă relațională cu atributele *ID* de tipul *INTEGER*, *Nume_Prenume* de tipul *VARCHAR(50)* și *An_nastere* de tipul *INTEGER* și fie în această relație este înregistrarea (202315, 'Odobescu, Alexandru', *NULL*), adică anul nașterii este necunoscut.

Identificatorul *ID* este stocat pe 4 octeți și numele și prenumele pe 8 octeți, dintre care un octet este rezervat pentru lungimea câmpului. Antetul înregistrării conține un pointer spre schema relației, lungimea sa totală (4+8), și o mască de biți 110 care indică faptul că al treilea câmp are valoarea *NULL*. Figura 28.4 reprezintă această înregistrare. De notat, că citind antetul, se poate calcula adresa înregistrării următoare.

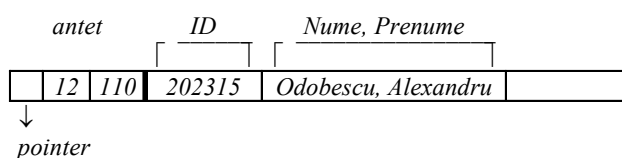


Figura 28.4 O înregistrare cu antet

Blocuri

Unitatea de transfer de date între fișierul memorat pe mediu și memoria internă este *blocul*. Lungimea unui bloc este, de obicei, o putere a lui 2 cuprinsă între 2^9 și 2^{12} octeți. Un bloc în sistemul Oracle, de exemplu, ocupă 4096 sau 8092 octeți. Fiecărui bloc i se asociază o adresă.

Structura blocului. Stocarea înregistrărilor într-un fișier trebuie să țină cont de divizarea în blocuri a acestui fișier. În general, într-un bloc se pot plasa mai multe înregistrări, dar se evită ca o înregistrare să se împartă între două blocuri. Numărul maximal de înregistrări de lungimea L memorate într-un bloc de lungimea B este $\lfloor B/L \rfloor$, unde notația $\lfloor x \rfloor$ desemnează cel mai mare întreg inferior lui x .

De exemplu, fie un fișier memorează o relație a cărei schemă nu conține atribute de lungime variabilă, adică nu utilizează tipurile *VARCHAR* sau *BIT VARYING*. Înregistrările au, deci, o lungime egală cu suma lungimilor tuturor câmpurilor. Fie că această lungime este de aproximativ 84 octeți, iar lungimea blocului este de 4096 octeți. În afară de aceasta, fie că fiecare bloc conține un antet de 100 octeți pentru a stoca datele despre spațiul liber disponibil în bloc, înlănțuirea cu alte blocuri etc. Atunci, se pot memora $\lfloor (4096-100)/84 \rfloor = 47$ înregistrări într-un bloc. De notat că în fiecare bloc rămân $3996 - (47 \cdot 84) = 48$ octeți neutilizați.

Transferul, în memorie, al înregistrării 563 a acestui fișier este simplu de efectuat: se determină în ce bloc se găsește ea (fie $\lfloor 563/47 \rfloor + 1 = 12$), se încarcă al 12-lea bloc în memoria centrală și din acest bloc se extrage înregistrarea. Prima înregistrare a blocului este $11 \cdot 47 + 1 = 518$, iar ultima înregistrare este $12 \cdot 47 = 564$. Înregistrarea 563 este, deci, penultima în blocul cu numărul intern 46, figura 28.5.

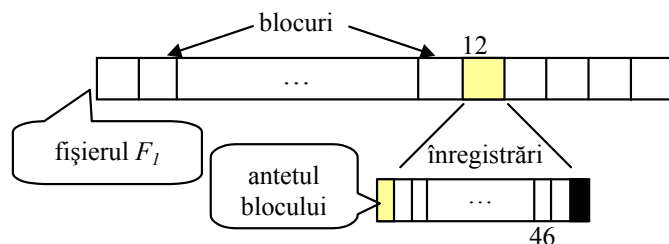


Figura 28.5 Un fișier cu blocuri și înregistrări

Calculul de mai sus arată cum se poate localiza fizic o înregistrare: prin fișierul ei, apoi prin blocul ei, apoi prin poziția ei în bloc. Dacă fișierul e nominalizat cu F_1 , adresa înregistrării poate fi reprezentată prin ' $F1.12.46$ '.

Există multe alte metode de adresări. Dezavantajul utilizării adreselor fizice, de exemplu, este că nu se poate schimba locul unei înregistrări fără a genera adresări invalide ale pointerilor asupra acestei înregistrări (de exemplu, în indecși).

Pentru a permite deplasarea înregistrărilor, se poate forma o *adresă logică*, ce ar identifica o înregistrare independent de locația ei. Un tabel de corespondență permite administrarea asocierii între adresa fizică și adresa logică, figura 28.6. Acest mecanism face ca organizarea și reorganizarea bazei de date să fie o procedură flexibilă. Acum e suficientă referirea unei înregistrări prin adresa sa logică, iar modificarea adresei fizice în tabel să se efectueze când are loc o deplasare. În schimb, această metodă necesită un cost suplimentar pentru că, sistematic, trebuie examinat tabelul de corespondență pentru a accesa datele.

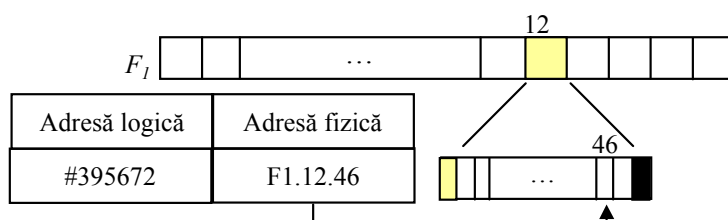


Figura 28.6 O adresare indirectă

O soluție intermediară este îmbinarea adresărilor logică și fizică. Pentru a localiza o înregistrare se indică adresa fizică a blocului, apoi în blocul propriu-zis se administrează un tabel care dă localizarea în bloc sau eventual în alt bloc.

Fie, din nou, înregistrarea $F1.12.46$. Aici $F1.12$ indică blocul 12 al fișierului F_1 . Se presupune că 46 este identificatorul logic al înregistrării administrate în interiorul blocului. Figura 28.7 prezintă această adresare în două niveluri: în blocul $F1.12$, înregistrarea 46 corespunde unei amplasări în același bloc, pe când înregistrarea 57 este plasată în alt bloc.

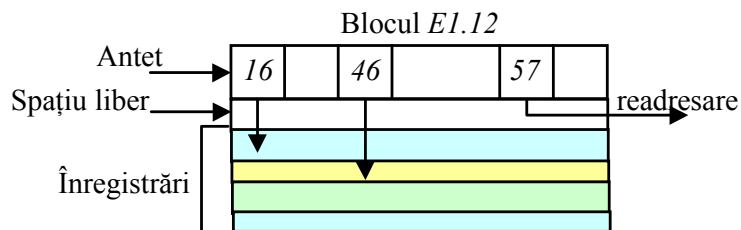


Figura 28.7 O îmbinare a adresărilor logică și fizică

Trebuie menționat că spațiul liber din bloc este situat între antetul blocului și înregistrări. El permite mărirea simultană a acestor două elemente în cazul unei inserări, de exemplu, fără efectuarea reorganizării interne a blocului. Acest mod de identificare oferă multe avantaje și permite reorganizarea suplimentară a spațiului intern al unui bloc.

Blocuri cu înregistrări de lungime variabilă. O relație care e definită pe atribute de tipul *VARCHAR* sau *BIT VARYING* este reprezentată prin înregistrări de lungime variabile. Când o înregistrare este inserată într-un fișier, lungimea se calculează nu după *tipul* de atribute, ci după numărul real de octeți necesari reprezentării *valorilor* atributelor. Această mărime trebuie stocată la începutul înregistrării curente, pentru ca SGBD-ul să poată determina începutul înregistrării următoare.

Se poate întâmpla că înregistrarea este actualizată, adică este modificată valoarea unui atribut sau unui atribut inițial i-a fost dată valoarea *NULL*. În acest caz, se poate întâmpla ca locul rezervat inițial să fie insuficient pentru noile date și, deci, înregistrarea urmează să fie memorată în alt loc al aceluiași fișier. Astfel, e necesară crearea unei legături între înregistrarea anterioară și curentă, memorată în alt bloc.

La locul deplasării înregistrării întregi, unele SGBD-uri aplică tehnica de fragmentare a înregistrării și de memorare în alt bloc a unui fragment, organizând, bineînțeles, o înlănțuire la nivel de înregistrare. Deplasarea (sau fragmentarea) înregistrărilor de lungime variabilă, evident, influențează eficiența de accesare a datelor. În afară de aceasta, înregistrările de lungime variabilă sunt ceva mai complicat de administrat decât cele de lungime fixă. În schimb, un fișier care conține înregistrări de lungime variabilă utilizează, frecvent, mai puțin spațiu decât i se atribuie.

Chei

Dacă fișierul stocat este organizat ca un grup de înregistrări secvențiale, trebuie să fie forme de recuperare a unei înregistrări specifice, executând un număr minimal de accesări. Astfel, înregistrările trebuie să fie ordonate pentru același criteriu de căutare. Pentru aceasta, fiecărei înregistrări i se asociază o *cheie* bazată pe conținutul său pentru a fi utilizată pentru identificarea univocă. În acest context, cheia este un instrument conceptual important pentru menținerea consistenței datelor și pentru asigurarea procesului de restabilire. De exemplu, cheia în fișiere este mai bine de folosit împreună cu unele tehnici de căutare speciale (binară, pe blocuri etc...), în baza valorii ei, pentru a verifica înregistrările în mod secvențial.

O *cheie primară* reprezintă unul sau mai multe atribute, care, univoc, identifică (una și numai una) o înregistrare.

O *cheie secundară* nu identifică univoc o înregistrare și poate fi utilizată pentru căutarea simultană a mai multor înregistrări cu aceeași valoare a cheii. Astfel, definiția cheilor secundare trebuie să folosească unul sau mai multe câmpuri cu o semnificație pentru utilizator și a căror date sunt, de obicei, folosite pentru căutarea datelor în fișier. Deci, este o cheie de ordonare pentru recuperarea unei mulțimi de date, organizate diferit de organizarea cheii primare.

Pe parcursul acestui capitol, se vor utiliza și alte noțiuni legate de cheie. Urmează o explicare a acestor concepte:

O *cheie compusă* este o cheie formată din mai multe atribute sau câmpuri.

O *cheie externă* este o mulțime de atribute care constituie cheia primară a altui fișier. Mulțimea de valori a cheii externe constituie o submulțime a mulțimii de valori ale cheii primare.

O *cheie de acces* este cheia utilizată pentru identificarea (recuperarea) unei înregistrări. Aceasta reprezintă unul sau o mulțime de atribute pentru căutarea înregistrărilor în fișier.

Un *argument de căutare* este valoarea cheii de acces la înregistrare într-o operație de căutare.

O *cheie de ordonare* reprezintă o mulțime de atribute folosite la ordonarea înregistrărilor unui fișier.

O *cheie a unei înregistrări* este valoarea cheii primare pentru această înregistrare.

Accesul la datele stocate pe un suport periferic, spre deosebire de aplicațiile care manipulează datele în memoria centrală, reprezintă una dintre particularitățile esențiale ale SGBD-urilor. Aici apar probleme de îmbunătățire a performanței, deoarece timpul de citire a unei unități de date pe un disc este considerabil mai mare decât timpul de acces în memoria principală. Organizarea datelor pe un disc în fișiere, structurile de indexare și algoritmi de căutare utilizați constituie aspecte foarte importante ale SGBD-urilor. Un sistem performant trebuie să utilizeze eficient tehnicile disponibile în scopul micșorării timpului de acces.

Organizarea fișierelor este o aranjare a unui tip de structură construită astfel încât să furnizeze un mediu eficient pentru stocarea și manipularea datelor în memoria secundară, economisind spațiul utilizat, și mărin d viteza de acces (minimizând timpul total de accesare și transfer al blocurilor de date) la înregistrările din fișier. În organizarea fișierelor, se ține cont de mai mulți factori, cum sunt frecvența cu care se efectuează anumite operații, câmpurile implicate în operațiile de căutare (cheile de căutare), dimensiunile înregistrărilor (fixe sau variabile).

Cele mai utilizate tipuri de organizare a fișierelor în bazele de date sunt: secvențial, indexat secvențial, indexat și cu dispersie. Aceste tipuri de organizare a fișierelor sunt descrise în continuare.

28.3 Fișiere secvențiale

Prin natura lor, fișierele secvențiale sunt asociate dispozitivelor cu acces secvențial, precum sunt, de exemplu, benzile magnetice. Ele se caracterizează prin prelucrarea în lot, mari volume de date, costuri minimale, spațiu redus de păstrare, utilizare în benzile magnetice. Uneori, fișierele secvențiale sunt stocate și pe dispozitive cu acces direct, în

particular, când este necesară prelucrarea datelor, în mod secvențial, cu o viteză mai înaltă.

Într-un mediu multiutilizator, cum este baza de date, trebuie ținut cont că mai mulți utilizatori pot partaja același dispozitiv de stocare și accesul la o nouă înregistrare, de obicei, necesită o poziționare nouă a capului de citire/scriere pe cilindrul care o conține. În afară de aceasta, timpul de mișcare a capului de citire depinde de schema memoriei intermediare utilizate și de faptul dacă sistemul realizează sau nu citirea anticipată.

Utilizarea acestui tip de organizare se recomandă pentru fișierele mici și în cazul când nu sunt frecvent modificate. De exemplu, în fișierele pentru salvare, există puține includeri de date și, în general, ele sunt făcute în lot, periodic. Modificările și suprimările sunt puține sau absente. Consultările, evident, sporadice, pot fi făcute rapid, iar parcurgerea secvențială este cea dezirabilă.

Fișiere secvențiale ordonate

Fișierele secvențiale ordonate reprezintă o structură de fișiere secvențiale, unde înregistrările sunt clasificate de valorile câmpurilor care formează cheia de sortare (de obicei, cheia primară) și stocate astfel încât ordinea logică coincide cu ordinea fizică. Fișierele secvențiale sunt stocate pe disc în blocuri pe poziții fizice continue pe pistele unui și aceluiași cilindru și apoi pe pistele cilindrului adiacent. Deoarece înregistrările în fișierele secvențiale sunt stocate în succesiune continuă, accesarea înregistrării n a fișierului presupune că cele $(n-1)$ înregistrări (începând cu prima), de asemenea, sunt citite.

După cum fișierele sunt unice și cheia de ordonare este unică (fiecare instanță poate fi clasificată numai într-un singur mod). În cazul când fișierele secvențiale nu posedă chei de organizare (ceea ce, în bazele de date, nu are loc), înregistrările sunt aranjate în serie, fiindcă, în general, fiecare înregistrare nouă se plasează la sfârșitul fișierului. În această situație, se adaugă un câmp suplimentar care conține ordinea sau numerele de identificare conform cărora fișierul este ordonat și acest câmp poate fi considerat cheie primară.

Organizarea secvențială este mai perfectă decât organizarea serială a fișierelor (*heap file*), dar ele pierd din flexibilitate, deoarece nu se adaptează ușor la operațiile de modificare. Fișierele secvențiale sunt utile pentru clasificarea și accesarea volumelor mari de date și, din motive economice, în general, sunt stocate pe benzi magnetice.

Fișierele sunt ordonate fizic și, pe parcursul perioadei de păstrare a datelor, această ordine (logică și fizică) trebuie menținută.

Fișierul secvențial stocat pe bandă magnetică poate fi deschis ca un fișier de ieșire sau de intrare. Dar, dacă fișierul este stocat pe disc, poate fi deschis pentru intrare, ieșire și modificare. Dacă se deschide un fișier secvențial, capul de citire întotdeauna este fixat pe prima înregistrare a fișierului și de multe ori (pe bandă, de exemplu) nu are cum să se întoarcă la înregistrarea citită anterior. Pentru aceasta, fișierul trebuie să fie închis pentru a fi rebobinată banda magnetică și apoi deschis.

De obicei, fișierul secvențial este simplu de închis, dar, pentru banda magnetică, operația respectivă presupune rebobinarea ei până la începutul fișierului. Această sarcină este o funcție fizică realizată în mod automat.

Operația de acces la un fișier secvențial poate să se producă sub două forme: cu cheia de acces diferită de cheia de ordonare și cu cheia de acces egală cu cheia de ordonare.

Manipularea înregistrărilor în ordinea stocării este eficientă, deoarece ele sunt stocate fizic în ordinea în care sunt solicitate – secvența fizică este egală cu cea logică. *Accesul secvențial* (constă în obținerea înregistrării care urmează după ultima accesată, în secvența (de obicei, ascendentă) definită de cheia de ordonare. În acest mod, dacă se folosesc tehnicile de buferizare și *caracteristicile blocului*, în majoritatea accesărilor, înregistrarea dezirabilă deja va fi în memorie, deoarece înregistrarea succesivă va fi în același bloc cu precedenta. Astfel, parcurgerea secvențială după cheia de ordonare este simplă, întrucât fișierul se parcurge de la început la sfârșit. Însă, parcurgerea secvențială după altă cheie va necesita o ordonare prealabilă într-un fișier auxiliar.

Avantajul principal al acestei organizări este facilitatea de realizare a operației de parcurgere secvențială, în afară de simplitatea implementării celorlalte operații. Cel mai mare dezavantaj este timpul de executare a operațiilor de includere și excludere a înregistrărilor, pentru că multe înregistrări sunt deplasate pentru păstrarea ordinii fizice și logice.

Accesul aleatoriu, interogare, este caracterizat de identificarea înregistrării prin specificarea argumentului de căutare. Secvența de acces nu este legată neapărat de ordinea fizică a fișierului, având ca rezultat înregistrări care nu sunt stocate în formă continuă. În general, consultarea aleatorie a unei singure înregistrări după cheia de ordonare se realizează prin folosirea tehnicilor de căutare mai eficiente, precum cea binară. În căutarea binară, numărul maximal de comparații, pentru a atinge înregistrarea căutată, este $(\log_2 n) + 1$, unde n este numărul de înregistrări ale fișierului.

Utilizarea principală a fișierelor secvențiale este legată de prelucrarea secvențială a datelor. Avantajul posibilității accesării rapide a înregistrărilor în formă continuă devine un dezavantaj, dacă fișierul este utilizat pentru accesarea unei înregistrări care nu este una următoare sau dacă atributul de căutare nu este cheie de ordonare. Parcurgerea va fi lentă, creând impresia că fișierul nu este ordonat, necesitând o căutare exhaustivă și citind, în medie, jumătate de fișier pentru a găsi înregistrarea specificată. Astfel, în general, procesarea unui fișier secvențial se face conform modului de organizare a acestuia, adică, predomină procesarea secvențială.

Prin urmare, dacă cheia de acces este diferită de cea de căutare, atunci fișierul este unul serial în care se realizează o căutare secvențială, pornind de la prima înregistrare, până când se localizează cea cu valoarea cheii de acces egală cu argumentul de căutare sau se atinge extremitatea fișierului, adică înregistrarea căutată nu se găsește în fișier. Când cheia de acces este egală cu cheia de ordonare există două alternative:

- pentru dispozitive cu acces secvențial fișierele sunt citite secvențial până când există înregistrări pentru examinare și cheia de căutare este mai mică decât valoarea atributului cheie sau până când înregistrarea căutată este găsită;
- pentru dispozitive cu acces direct sunt utilizate tehnici mai eficiente, precum căutarea binară, prin blocuri sau interpolare.

Consultarea cu cheia de acces diferită de cheia de ordonare. Consultarea cu cheia de acces diferită de cheia de ordonare are loc ca în fișierele neordonate. Căutarea este făcută prin lectura exhaustivă până când

se localizează înregistrarea căutată sau se termină fișierul. Aici $NMC = (n+1)/2$, unde NMC este numărul mediu de comparații, iar n – numărul de înregistrări în fișier. Algoritmul poate fi următorul:

1. *Se merge la poziția inițială a fișierului.*
2. *Până când nu este atins sfârșitul fișierului:*
 - a) *dacă înregistrarea curentă = înregistrarea dorită, terminare cu succes;*
 - b) *avansare cu o înregistrare.*
3. *Terminare cu eșec.*

Consultarea cu cheia de acces egală cu cheia de ordonare. Consultarea cu cheia de acces egală cu cheia de ordonare (sau cu partea ei inițială) a fișierului stocat pe un dispozitiv cu acces secvențial se face cu o căutare secvențială. Singurul avantaj este că, dacă înregistrarea curentă are valoarea cheii mai mare decât a celei căutate, ea nu există și căutarea este întreruptă. Dacă, însă, dispozitivul permite accesul direct, se poate realiza o căutare mai eficientă, cum ar fi căutarea binară și atunci $NMC = (\log_2 n) + 1$. Algoritmul de căutare binară este următorul:

1. *Definirea P_i (poziția primei înregistrări a fișierului) și P_f (poziția ultimei înregistrări a fișierului).*
2. *Până când ($P_i \leq P_f$) se face:*
 - a) *calcularea P_m (poziția medie) = $(P_i + P_f)/2$;*
 - b) *se trece la poziția P_m ;*
 - c) *dacă înregistrarea curentă P_m = înregistrarea căutată, terminare cu succes;*
 - d) *dacă înregistrarea curentă P_m > înregistrarea curentă, $P_i = P_m + 1$;*
 - e) *dacă înregistrarea curentă P_m < înregistrarea curentă, $P_f = P_m - 1$;*
3. *Terminare cu eșec.*

Inserarea înregistrărilor. Inserarea unei înregistrări în timp real are un cost înalt, deoarece trebuie efectuată reordonarea fișierului după cheia de ordonare. Pentru aceasta, se determină poziția adecvată a înregistrării noi (conform cheii primare), se deplasează toate înregistrările care posedă cheia mai mare decât a celei incluse, se inserează înregistrarea nouă. Algoritmul de inserare a înregistrării

1. *Se merge la sfârșitul fișierului.*
2. *Până când argumentul de căutare nu este mai mare decât valoarea respectivă a înregistrării curente și nu s-a ajuns la începutul fișierului:*
 - a) *se deplasează înregistrarea cu o poziție (poziția nouă = poziția anterioară + 1);*
 - b) *se trece la înregistrarea vecină (se micșorează cu o poziție).*
3. *Se inserează înregistrarea nouă în fișier.*

O alternativă pentru inserarea unei înregistrări poate fi următorul algoritm: se creează un fișier auxiliar și se copiază fișierul original, plasând înregistrarea nouă în poziția corectă, se șterge fișierul original și se renumește fișierul auxiliar cu numele fișierului original.

Însă, pentru fișierele mari, stocate pe suport extern, aceste procese au un cost prohibitiv. De obicei, se utilizează un fișier auxiliar care conține înregistrările noi, ordonate în același mod (după aceeași cheie) ca și fișierul principal de date. Dar utilizarea fișierului auxiliar influențează procesul de realizare a tuturor operațiilor care trebuie efectuate asupra ambelor fișiere și nu numai asupra unuia. Toate accesările înregistrărilor, care sunt realizate pentru toate operațiile, trebuie să fie executate în două fișiere.

Suprimarea înregistrărilor. Suprimarea sau eliminarea înregistrării trebuie făcută la nivel fizic, cu reorganizarea fișierului în timpul executării operației. Suprimarea presupune următoarele activități:

1. *Se merge la începutul fișierului.*
2. *Se localizează înregistrarea care trebuie exclusă.*
3. *Se trece, una după alta, fiecare înregistrare de după înregistrarea ce urmează a fi suprimată și se micșorează poziția ei cu 1.*

O alternativă, de asemenea, poate fi utilizarea unui fișier auxiliar. În acest caz, se trece la începutul fișierului, se copiază înregistrările de până la cea care trebuie suprimată, se ignoră înregistrarea care trebuie exclusă, se copiază înregistrările rămase în fișierul auxiliar, se substituie fișierul inițial cu cel auxiliar.

Excluderea fizică, deci, necesită un cost prohibitiv pentru prelucrare în timp real, deoarece trebuie deplasate toate înregistrările care urmează după cea suprimată. Prelucrarea poate fi făcută în lot, dacă operațiile pot fi realizate mai târziu, adică comenzile de excludere pot fi colectate într-un fișier de tranzacții pentru realizarea lor ulterioară. De multe ori, această procedură nu poate fi acceptată. Mai frecventă este procedura care presupune includerea în înregistrare a unui câmp adițional, în care se indică excluderea. Astfel, pentru a elimina o înregistrare se modifică doar valoarea acestui câmp (modificarea înregistrării), prin care se arată că ea a fost eliminată. Cu această procedură, este evitată necesitatea deplasării altor înregistrări pentru completarea spațiilor eliberate în urma eliminării.

Modificarea înregistrărilor. Modificarea unei înregistrări presupune modificarea valorilor unor câmpuri ale înregistrării. Pentru aceasta, înregistrarea este localizată, citită, câmpurile ei sunt modificate și apoi ea este scrisă. În general, sunt actualizate numai atributele care nu fac parte din cheie. Dar, dacă înregistrarea are lungimea variabilă și modificarea mărește lungimea ei, înregistrarea nu poate fi înscrisă în poziția inițială, deoarece lipsește spațiul necesar. Atunci, de obicei, ea este exclusă și apoi inclusă, dar, deja, actualizată.

Dacă se modifică vreun atribut al cheii de ordonare, schimbarea valorii lui implică schimbarea poziției înregistrării. În acest caz, operația este, în general, implementată prin excluderea înregistrării vechi, urmată de includerea în poziția respectivă a variantei modificate. Însă, în majoritatea aplicațiilor, nu este permisă modificarea cheii. În alte aplicații poate fi folosit un fișier de tranzacții. Dacă este utilizată banda magnetică, operația de rescriere poate fi executată doar mai târziu, întrucât nu poate fi rescrisă o înregistrare care este citită.

Lectura exhaustivă a înregistrărilor. Lectura exhaustivă a înregistrărilor este o operație eficientă, deoarece constă în citirea fiecărei înregistrări din fișier. Cea mai adecvată organizare pentru acest tip de operații este cea secvențială.

În general, dacă numărul de operații de actualizare a fișierului este foarte mare, se poate utiliza una din opțiuni: un fișier auxiliar adițional cu actualizări care vor fi făcute mai târziu (procesul *batch*); inserarea la sfârșitul fișierului și reorganizarea (ordonarea) ulterioară a lui.

Fișiere secvențiale ordonate la nivel logic de pointeri

Tehnicile descrise anterior oferă avantaje în cazul prelucrării secvențiale și în cazul căutării înregistrărilor cu metode mai rapide de căutare. Însă dezavantajul lor este timpul cheltuit în operațiile de actualizare (includere, excludere și modificare a atributelor utilizate în ordonarea fișierului). Aceasta se datorează faptului că fișierele trebuie reorganizate în toate operațiile de includere/excludere, deplasând înregistrările pentru a menține ordinea fizică egală cu cea logică. Prin urmare, deoarece au un cost înalt, pentru includerea, modificarea și excluderea înregistrărilor, trebuie propuse tehnici mai eficiente.

Fișierul secvențial cu pointeri a fost elaborat pentru a mări eficiența sistemelor cu multe inserări și suprimări, cum sunt, de exemplu, bazele de date. Structura unui fișier cu pointeri este similară unei liste înlănțuite. În acest caz, nu este necesar a avea o ordine fizică a fișierului egală cu ordinea logică. Această tehnică, deși are dezavantajul că este mai lentă în parcurgerea secvențială a înregistrărilor, poate fi necesară la citirea și recitirea unor sectoare care sunt plasate în diferite locuri ale fișierului (necontinuu).

Acest tip de fișiere posedă unele structuri adiționale, figura 28.8. În primul rând, există un atribut care indică ordinea logică, adică adresa următoarei înregistrări din fișier, în secvența structurii logice conform cheii de clasificare. În al doilea rând, există atributul înregistrărilor șterse, care indică faptul că înregistrarea este activă sau deja a fost eliminată. Altă structură este *înregistrarea zero* sau *antetul*. Înregistrarea zero posedă atributul *Început*, care arată adresa primei înregistrări în fișier, figura 28.8, înregistrarea aflată pe adresa 2. Astfel, structura reprezintă o listă care indică înregistrările în uz și ordinea lor în fișier, și un câmp *Exclusă* pentru înregistrările șterse.

		<i>Început</i>			
0		2		<i>Antet</i>	
	Nume	...	Exclusă	Pointer	
1	Donică	...		3	
2	Amihălăchioaei	...		5	
3	Guțu	...	*	4	
4	Munteanu	...		6	
5	Petrescu	...	*	1	
6	Rotaru	...		*	

Figura 28.8 Un fișier secvențial cu pointeri

Există și implementarea cu două liste, figura 28.9. Deși, prin adăugarea unui pointer în înregistrare structura se complică, această formă de organizare permite utilizarea spațiului înregistrărilor eliminate pentru inserarea înregistrărilor noi. Or, căutarea spațiului disponibil, al unei înregistrări suprimate, într-un fișier cu multe înregistrări este lentă, deoarece necesită multe accesări. În afară de pointerul de păstrare a secvenței logice a înregistrărilor este utilizat un alt pointer pentru păstrarea listei înregistrărilor excluse. Astfel, poate fi utilizat spațiul eliberat, inserând, mult mai rapid, înregistrări noi. Nu se cunoaște doar poziția primei înregistrări și nici prima înregistrare care a fost eliminată. Adică, este necesară adăugarea la fișier a unui antet pentru păstrarea acestor informații. În structura cu două liste, una indică înregistrările în uz și ordinea lor în fișier, iar alta - înregistrările suprimate al căror spațiu poate fi utilizat pentru inserare.

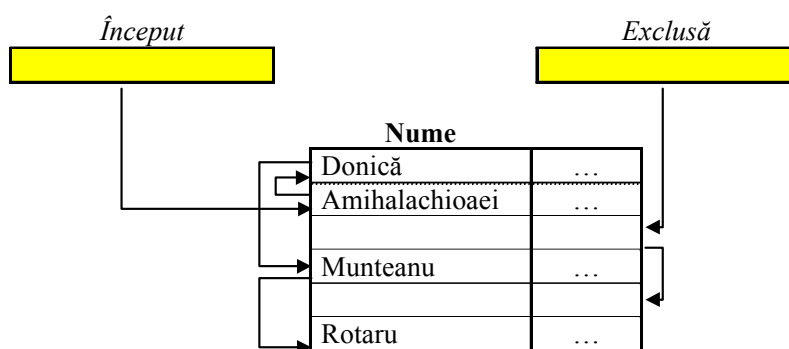


Figura 28.9 Un fișier secvențial cu două structuri de pointeri

În structura cu doi pointeri, în procesul de consultare, se identifică prima înregistrare cu adresa indicată de atributul Început al antetului. Dat fiind faptul că fișierul nu este ordonat fizic, în această structură, nu este posibilă realizarea căutării binare. Apoi, se parcurge fișierul, folosind adresele indicate în atributul Pointer al secvenței de înregistrări până ce câmpul cerut va fi localizat. Algoritmul de consultare:

1. Se trece la poziția primei înregistrări indicată în antet.
2. Până când înregistrarea curentă $\leq$ înregistrarea dezirabilă:
 - a) dacă înregistrarea curentă = înregistrarea dezirabilă, terminare cu succes;
 - b) dacă pointerul secvenței ≥ 0 , se trece la următoarea înregistrare indicată de pointer;
 - c) altfel terminare cu eșec.
3. Terminare cu eșec.

După cum se poate observa, sfârșitul căutării (sfârșitul fișierului) este atins când pointerul secvenței de înregistrări nu indică spre altă înregistrare (nu posedă vreo valoare sau adresă).

Pentru a include o înregistrare, nu mai e nevoie de reorganizarea fișierului, deoarece ordinea fizică, în acest caz, poate fi diferită de ordinea logică (care este determinată de pointeri). Este suficientă inserarea înregistrării în primul loc liber. Primul loc liber poate fi o înregistrare suprimată și pentru a ști că există o astfel de înregistrare, este destul de

verificat atributul *Exclusă* în antetul fișierului. În cazul în care nu există înregistrări excluse, un spațiu nou este creat la sfârșitul fișierului. După selectarea locului de inserare se parcurge fișierul pentru identificarea înregistrărilor precedentă și următoare ale înregistrării curente. Acestea trebuie să aibă valorile pointerilor actualizate, pentru ca înregistrarea nouă să ocupe poziția logică intermediară.

Algoritmul este următorul:

1. *Localizarea unei înregistrări vide (suprimate) sau crearea unui spațiu pentru înregistrarea nouă la sfârșitul fișierului.*
2. *Inserarea înregistrării în locul identificat sau creat.*
3. *Dacă înregistrarea a fost inserată în locul uneia eliminate, se exclude din lista înregistrărilor suprimate.*
4. *Identificarea înregistrărilor precedente și următoare;*
5. *Dacă există înregistrarea precedentă se actualizează pointerul ei pentru ca să arate spre înregistrarea inserată.*
6. *În caz contrar, pointerul antetului este cel ce trebuie actualizat.*
7. *Dacă există înregistrarea următoare, pointerul înregistrării inserate se actualizează pentru a o indica.*

Pentru a exclude o înregistrare, se identifică înregistrarea care precede înregistrarea în cauză și se actualizează pointerul ei. Deoarece înregistrarea precedentă trebuie să dea continuitate logică fișierului, pointerul ei se modifică pentru a indica înregistrarea ce urmează după înregistrarea exclusă. Astfel, valoarea pointerului înregistrării suprimate devine valoare a pointerului înregistrării anterioare.

Algoritmul constă din patru pași:

1. *Localizarea înregistrării ce trebuie suprimată, utilizând lista secvenței de înregistrări (pointerul).*
2. *Identificarea înregistrării precedente și înregistrării următoare;*
3. *Actualizarea pointerului înregistrării precedente, atribuindu-i valoarea pointerului înregistrării ce trebuie suprimată;*
4. *Adăugarea înregistrării excluse în lista înregistrărilor excluse.*

Modificarea unei înregistrări poate fi realizată prin eliminarea înregistrării ce trebuie modificată și inserarea înregistrării cu valorile deja modificate (metodă mai laborioasă, dar mai facilă și practică).

Fișiere secvențiale cu spațiu de tranzacții

Fișierele secvențiale pot fi dotate cu un fișier auxiliar, numit *fișier de tranzacții*. Deplasarea înregistrărilor este o operație costisitoare. Pentru a evita deplasarea înregistrărilor (rearanjarea) în fișierul principal în timpul executării operațiilor, o altă formă de implementare a fișierului secvențial presupune utilizarea a două structuri:

- fișierul principal (master) cu date.
- fișierul de modificări sau fișierul de tranzacții (overflow) este un fișier auxiliar, în care sunt temporar înregistrate modificările curente ale datelor, care vor fi utilizate pentru actualizarea conținutului fișierului principal. În acest mod, nu au loc deplasări de înregistrări în fișierul principal în timpul de executare a

operațiilor. Fișierul de tranzacții poate fi chiar unul virtual, adică, înregistrările pot fi plasate la sfârșitul fișierului principal.

Astfel, pentru a micșora costul înalt de păstrare a ordinii fizice, se utilizează două structuri, plus o listă de pointeri a secvenței de înregistrări pentru menținerea ordinii logice, adică se separă grupul de înregistrări care sunt fizic ordonate de operațiile nou-realizate.

Fișiere secvențiale cu toate operațiile în fișierul de tranzacții. Această variantă de fișiere secvențiale presupune că toate modificările sunt făcute în fișierul de tranzacții. Fișierul principal este actualizat numai în cazul reorganizării, folosind datele din fișierul de tranzacții.

În fiecare înregistrare din fișierul de tranzacții, trebuie să fie un câmp auxiliar, unde se indică tipul operației ce a fost realizată (*I* – inserarea, *E* – eliminarea și *M* – modificarea). Toate operațiile (consultarea, inserarea, eliminarea, modificarea) trebuie să înceapă cu căutarea înregistrării în fișierul de tranzacții și examinarea etichetei acestui câmp. Dacă înregistrarea nu este găsită, atunci este căutată în fișierul principal. Trebuie menționat că fișierul de tranzacții este, de asemenea, un fișier secvențial și, deci, se supune acelorași reguli (înregistrări ordonate de cheia de ordonare).

Inserarea trebuie să fie precedată de o consultare. Dacă cheia este găsită în fișierul de tranzacții cu eticheta *I* sau *M*, sau dacă cheia este întâlnită în fișierul principal, atunci trebuie să fie acționat procesul de tratare a erorii (nu poate fi făcută o inserare cu o cheie deja existentă). Dacă cheia e găsită în fișierul de tranzacții și are eticheta *E*, se inserează o înregistrare nouă în fișierul de tranzacție, care este stocată în ordinea cheii de ordonare. Astfel, pot fi necesare deplasări de înregistrări. Cu toate acestea, deoarece fișierul de tranzacții este mic, executarea operației de deplasare este ușoară. Eticheta *I* trebuie plasată în înregistrarea nouă, indicând că este una inserată.

Utilizarea fișierului de tranzacții schimbă modul de consultare a înregistrărilor. Astfel, o consultare a unei singure înregistrări după cheia de ordonare, începe de la fișierul de tranzacții, deoarece acesta conține ultimele modificări. Trebuie, de asemenea, să fie examinată eticheta operației (*I*, *E* sau *M*). Eticheta *E* indică faptul că cheia nu mai există, iar celelalte etichete indică faptul că consultarea va fi făcută asupra înregistrării în fișierul de tranzacții. Însă, o consultare cu o cheie diferită de cea de organizare are un cost înalt, deoarece fișierul de tranzacții trebuie analizat (cu examinarea etichetelor) în întregime și apoi fișierul principal, la fel, în mod exhaustiv.

Operația de parcurgere secvențială după cheia de ordonare este făcută în paralel în ambele fișiere și urmează aceleași raționamente ca ale operației de reorganizare, cu intercalare. Dacă este necesară parcurgerea secvențială a înregistrărilor în altă ordine, atunci operația trebuie făcută cu algoritmi de ordonare, ținând cont de existența înregistrărilor duplicate în ambele fișiere.

Modificările câmpurilor care nu fac parte din cheia de ordonare trebuie făcute în fișierul de tranzacții. Mai întâi, înregistrarea trebuie căutată (operația de consultare). Dacă ea nu este găsită în nici unul din fișiere, este acționată o procedură de eroare. Dacă ea este găsită numai în fișierul principal, o înregistrare nouă va fi inserată în fișierul de tranzacții, cu aceleași câmpuri ca ale înregistrării inițiale, dar cu modificările necesare și cu eticheta *M*. Dacă ea este găsită în fișierul de tranzacții, se verifică

eticheta ei. Dacă este *E*, ea deja a fost exclusă și operația este întreruptă. Dacă eticheta este *I*, se modifică direct această înregistrare cu eticheta *I* în fișierul de tranzacții, fără a modifica eticheta. Modificările (dacă se permit) în câmpurile care fac parte din cheia de ordonare sunt realizate în două operații: una de eliminare a înregistrării originale și una de inserare a înregistrării modificate.

Operația de suprimare se efectuează după următoarele reguli. Mai întâi, se caută înregistrarea (operația de consultare). Dacă nu s-a găsit, există o greșeală. Dacă există numai în fișierul principal, o înregistrare este adăugată în fișierul de tranzacții (în ordinea cheii) cu eticheta de excludere. Dacă, deja, există în fișierul de tranzacții, eticheta ei va indica următoarele procese. Dacă eticheta este *E*, există o eroare (înregistrarea a fost, deja, eliminată). Dacă este *M*, se schimbă această etichetă cu *E*, iar dacă este *I*, această înregistrare trebuie eliminată fizic din fișierul de tranzacții cu deplasări succesive de poziții (această operație nu va putea fi desfășurată).

Periodic, operațiile înregistrate în fișierul de tranzacții trebuie actualizate în fișierul principal. De aceea, eticheta fiecărei înregistrări din fișierul de tranzacții este analizată și operațiile respective sunt făcute în fișierul principal.

Deoarece această acțiune provoacă deplasări în fișierul principal, o alternativă este crearea unui fișier suplimentar, în care se copiază fișierul principal și în acesta sunt aplicate modificările. În finalul operației de reorganizare, fișierul de tranzacții devine vid, fișierul principal este suprimat și fișierul auxiliar este redenumit cu numele fișierului principal. Deoarece operația de reorganizare are un cost înalt, ea trebuie să fie făcută *off-line*, adică atunci când înregistrările nu sunt actualizate. Periodicitatea depinde de disponibilitatea de timp pentru realizarea acestei operații și de numărul de înregistrări în fișierul de tranzacții (se activează operația de reorganizare, dacă volumul de tranzacții atinge o anumită limită, după care eficiența generală de păstrare a datelor de SGBD scade).

Avantajul utilizării acestui tip de organizare apare în cazul în care sunt efectuate puține modificări *on-line* și când este posibilă întreruperea lor pentru realizarea operației de reorganizare. Alt avantaj îl constituie parcurgerea secvențială, în mod aproape optimal, deoarece înregistrările se găsesc, din punct de vedere fizic, continuu. Cu toate acestea, există tendința de a analiza fișierul principal și cel tranzacțional împreună.

Un mare dezavantaj al acestui tip de organizare îl constituie faptul că consultările pot fi făcute, în mod eficient, numai după cheia de organizare. Dacă sunt necesare consultări sau parcurgeri pe mai multe chei, această organizare nu este recomandată.

Fișiere secvențiale cu unele operații în fișierul principal. Pentru a reduce creșterea fișierului de tranzacții, operațiile care nu cer deplasări de înregistrări sunt realizate direct în fișierul principal. În acest mod, se diminuează viteza de creștere a fișierului de tranzacții și, prin urmare, micșorează perioada de aplicare a operațiilor de reorganizare.

În afară de aceasta, pentru a înviora consultarea când se fac operații asupra fișierului de tranzacții, el, de asemenea, ar trebui să fie menținut ordonat, mărin, astfel, costul operațiilor. De aceea, acest proces nu este acceptat și inserările se realizează la sfârșitul fișierului.

Se poate observa că ordinea fizică și ordinea logică sunt păstrate în fișierul principal. În fișierul de tranzacții, înregistrările nu sunt ordonate, întrucât înregistrările sunt inserate la sfârșit. De aceea, pentru a menține

aceeași ordine în care a fost ordonat, este necesară adăugarea unui pointer pentru fiecare înregistrare cu o adresă de deplasare (*offset*). Această adresă indică următoarea înregistrare care va fi folosită, dacă e necesar, pentru a întocmi lista, menținând, astfel, secvența logică a datelor după cheia de ordonare.

O dată cu adăugarea unui câmp pentru fiecare înregistrare, ocupat de pointer, există dezavantajul ca procesul de parcurgere secvențială se realizează mai lent, deoarece poate fi necesară citirea și recitirea sectoarelor care se găsesc în diferite locuri ale fișierului (necontinuu), fapt ce consumă mai mult timp de căutare și de stare latentă, în afară de supraîncărcarea cu spațiul ocupat de pointeri.

Inserarea unei înregistrări este făcută în fișierul de tranzacții sau, dacă este virtual, la sfârșitul fișierului principal în zona de tranzacții (de parcă nu ar exista cheia de clasificare). Secvența logică este păstrată de lista definită de pointerul care trebuie actualizat la fiecare operație realizată. Mai târziu, fișierul este reorganizat.

În cazul operației de *modificare*, dacă nici cheia primară și nici lungimea înregistrării nu sunt schimbate, înregistrarea este actualizată în aceeași poziție. În caz contrar, ea este eliminată și cea nouă este inserată în fișierul de tranzacții. Această operație trebuie să actualizeze lista secvenței logice a datelor.

În cazul operației de eliminare, fișierul principal poate fi reorganizat (operația necesitând mult timp). Atunci este necesară redefinirea listei secvenței logice, care ar include toate înregistrările care nu au fost eliminate.

O altă metodă este adăugarea unui câmp auxiliar pentru păstrarea unui indicator al înregistrărilor excluse, folosit la restructurarea ulterioară a fișierului.

Reorganizarea fișierelor. Fișierul de tranzacții, de obicei, este folosit doar pentru stocarea operațiilor de actualizare a fișierului secvențial și în operația de reorganizare care este aplicată numai în cazul când fișierul principal se reorganizează, deoarece acumularea tranzacțiilor, poate diminua eficiența executării lor. Așadar, în general, când fișierul de tranzacții atinge o limită determinată sau se dorește elaborarea unui raport, sau efectuarea unei consultări are loc reorganizarea fișierului secvențial. Or, în prelucrarea normală, este puțin obișnuită utilizarea fișierului de tranzacții ca o prelungire a fișierului secvențial.

Procesul de reorganizare a unui fișier secvențial *S* presupune aplicarea unui algoritm, cum este cel de triere prin fuziune. Pentru aceasta, fișierul de tranzacții *T* trebuie să fie ordonat după același criteriu (cheie) ca și fișierul *S*. Din fișierele (de tranzacții *T* și principal *S*) ordonate după același criteriu se poate construi un fișier secvențial nou *A*. Printr-o tratare secvențială a fișierelor, cu o lectură exhaustivă a lui *S*, se copiază fișierul principal inițial, înregistrare cu înregistrare, în fișierul principal actualizat *A* și, pentru fiecare înregistrare copiată, se intercalează cu înregistrările fișierului de tranzacții *T*, dacă există, făcând comparațiile necesare și producând copia actualizată a fișierului principal.

Algoritmul este următorul:

1. Se citește o înregistrare a fișierului principal *S* (dacă nu este marcată ca fiind suprimată).

2. Se citește o înregistrare a fișierului de tranzacții T (dacă nu este marcată ca fiind suprimată).
3. Se compară aceste două înregistrări. Cea cu valoare mai mică a cheii de ordonare este copiată în fișierul principal actualizat A.
4. În fișierul care conține înregistrarea scrisă în fișierul principal actualizat A, pointerul înregistrării avansează (se localizează următoarea înregistrare).
5. Se citește o înregistrare nouă a acestui fișier (dacă nu este marcată ca fiind suprimată).
6. Se compară aceste două înregistrări, scriind-o pe cea mai mică și repetând procesul până când se atinge sfârșitul unui fișier din cele două.
7. Se înscriu înregistrările nesuprimate ale fișierului, care încă nu a fost parcurs până la sfârșit, în fișierul principal actualizat A.

În acest proces, se observă că:

- înregistrările suprimate nu sunt copiate (sunt eliminate efectiv);
- înregistrările fișierului principal S, care nu figurează în fișierul de tranzacții T, sunt copiate în fișierul actualizat A;
- înregistrările fișierului de tranzacții T sunt prelucrate în lot și plasate în fișierul actualizat A cu păstrarea secvenței ordonate după cheie.

În consecință, fișierul de tranzacții devine vid, iar înregistrările fișierului principal actualizat A sunt ordonate astfel încât secvența logică este egală cu secvența fizică.

Căutarea secvențială cu blocuri de înregistrări

Partea ce mai scumpă (lentă) a operației de accesare a memoriei secundare este căutarea pentru obținerea poziției corecte a unei înregistrări pe disc. Așadar, se cere minimizarea numărului de accesări, pentru că transferarea datelor, odată inițiată, este relativ rapidă, deși este mult mai lentă decât transferarea datelor în memoria principală.

Evident, căutarea (*seek*) și citirea unei înregistrări și apoi căutarea și citirea altei înregistrări au un cost mai mare decât căutarea și citirea concomitentă a două înregistrări. Prin urmare, se poate îmbunătăți căutarea secvențială, dacă, în loc de înregistrări, se recuperează blocuri de înregistrări și aceste blocuri se prelucrează în memoria principală. Astfel, gruparea înregistrărilor în blocuri este o tehnică utilizată pentru îmbunătățirea căutării. În general, dimensiunea blocului este definită în funcție de caracteristicile fizice ale dispozitivului de stocare și ale datelor ce trebuie memorate.

Bineînțeles, fiecare accesare a unui bloc de înregistrări va lua ceva mai mult timp decât o accesare a unei înregistrări, dar beneficiul va fi considerabil datorită reducerii timpului de căutare și timpului de rotație. Astfel, formarea blocurilor de înregistrări:

- mărește eficiența căutării prin diminuarea numărului de accesări;
- profită de diferența dintre costul de accesare în memoria principală și costul de accesare a discului;
- economisește timpul, deoarece reduce timpii de căutare și durata stării latente;
- nu modifică numărul de comparații în memoria principală, dar, probabil, mărește cantitatea de date transferate între disc și

memoria principală (este citit un bloc întreg, în timp ce înregistrarea căutată este una în bloc, iar restul nu sunt necesare).

Astfel, se poate concluziona că deosebirea dintre accesul la memoria principală și disc este ceea ce ghidează proiectarea structurilor de fișiere și, prin urmare, căutarea secvențială este mai eficientă:

- în fișierele cu puține înregistrări;
- în fișierele cu puține căutări (de exemplu, păstrate pe benzi magnetice);
- în căutarea înregistrărilor după chei secundare, ale căror valori au mai multe duplicate.

28.4 Indecși

Indecșii sunt structuri de acces care se utilizează pentru accelerarea accesului la înregistrări și a răspunde anumitor criterii de căutare. Unele tipuri de indecși, numite și căi de acces secundare, nu afectează amplasarea fizică a înregistrărilor pe disc, fapt ce oferă căi alternative de acces pentru găsirea înregistrărilor, în mod eficient, în câmpurile indexate. Există și alte tipuri de indecși care se pot construi numai asupra fișierelor care au o anumită organizare.

Astfel, există tipuri de indecși care se utilizează asupra fișierelor ordonate (indecși cu un singur nivel) și structuri sub formă de arbore (indecși multinivel, B-arbori și B⁺-arbori). În afară de aceasta, se pot construi indecși, folosind funcții de dispersie și alte structuri de date.

Utilizarea tehnicilor de indexare poate fi o alternativă a ordonării, dacă este necesară organizarea unui fișier, pentru a fi căutat cu ajutorul cheilor.

În general, indecșii ameliorează executarea accesului la date. În cazul fișierelor, permit localizarea rapidă a înregistrărilor, cu avantajul că fișierul de date nu trebuie să fie reorganizat, dacă sunt inserate noi înregistrări în acesta. Este de ajuns să fie reorganizați indecșii. Indecșii, de asemenea, permit, în afară de ameliorarea timpului de acces pentru căutarea după o cheie, susținerea mai multor viziuni asupra înregistrărilor dintr-un fișier de date, grație structurilor de indecși secundari. Ba mai mult, cu ajutorul indecșilor, pot fi îmbinate mai multe viziuni particulare.

Indexul poate exista independent de organizarea fișierului de date, ceea ce permite crearea mai multor indecși, dacă se dorește optimizarea accesului la date al mai multor tipuri de interogări. În schimb, crearea fără discernământ a unui număr mare de indecși poate fi costisitoare pentru SGBD-ul care trebuie să-i administreze. Pentru fiecare operație de actualizare a relației, repercusiunile acestei actualizări se extind asupra tuturor indecșilor. O alegere judicioasă a indecșilor, într-un număr optim, este, deci, unul din factorii esențiali ai performanței unui sistem.

Concepte preliminare

Indexul este un fișier. Fișierul index este o structură auxiliară asociată fișierului de date, proiectat pentru a oferi forme mai eficiente de acces și localizare a datelor specificate. Adică, fiind dat un argument de căutare,

scopul indexului este accelerarea procesului de identificare a adresei înregistrării necesare în fișierul de date.

Pot exista unul sau mai multe fișiere index, ale căror înregistrări leagă valorile cheii cu pozițiile lor în fișierul de date. Adică o înregistrare a fișierului index este constituită din două câmpuri, cheia înregistrării din fișierul de date și adresa respectivă pe disc, <cheie, adresă>.

Astfel, un index este întotdeauna specificat de o cheie de acces, iar în calitate de cheie de acces poate fi luată orice submulțime de atribute ale schemei relaționale, dacă această submulțime este definită în calitate de index. Există două tipuri principale de indecși:

- indexul ordonat, care se bazează pe valorile ordonate ale cheilor. Înregistrările fișierului sunt stocate conform unui criteriu de ordonare;
- indexul hash sau index cu dispersie, care se bazează pe distribuirea uniformă a valorilor determinate de o funcție, numită funcție de dispersie (hash function).

Există diverse tehnici atât pentru indecșii ordonați, cât și pentru indecșii cu dispersie. Nici una din ele nu e mai bună, dar fiecare tehnică este mai adecvată pentru anumite aplicații ale bazei de date. Factorii care determină selectarea tipului de index sunt:

- tipul de acces, care este admis în mod eficient; tipurile de acces pot căuta înregistrări după o valoare determinată a unui atribut dat sau înregistrări ale căror atribute iau valori dintr-un interval de valori;
- timpul de acces – timpul consumat pentru găsirea unei înregistrări de date, sau a unei mulțimi de înregistrări, utilizând tehnica respectivă;
- timpul de inserare – timpul consumat pentru includerea unui nou articol de date; acest factor cuprinde timpul cheltuit pentru găsirea locului corect, pentru includerea articolului și, de asemenea, timpul cheltuit pentru actualizarea structurii indexului;
- timpul de suprimare – timpul consumat pentru eliminare, inclusiv timpul afectat găsirii înregistrării și actualizării structurii indexului.
- spațiul de memorie – spațiul adițional ocupat de fișierul index; dacă mărimea spațiului adițional este rezonabilă, în general, spațiul sacrificat se recuperează prin obținerea unei executări mai bune a operațiilor.

Un fișier index este mai ușor de lucrat decât un fișier de date, deoarece înregistrările indexului sunt de lungime fixă (facilitează navigarea cu tehnici mai eficiente de căutare, cum este cea binară) și este mult mai mic decât fișierul de date. Înregistrările de lungime fixă ale fișierului index impun și o limită a lungimii cheii primare, fapt care, uneori, poate crea probleme în practică. Înregistrările indexului pot conține și alte câmpuri (lungimea înregistrării, de exemplu), în afara celor două.

De obicei, pentru unul și același fișier de date pot exista mai multe fișiere index. Aplicând noțiunea de cheie de căutare, dacă sunt mai mulți indecși pentru un fișier de date, există mai multe chei de căutare în acest fișier.

Utilizarea indecșilor are următoarele avantaje:

- fișierul de date poate fi de tip serial (înregistrările sunt inserate la sfârșitul fișierului, în ordinea de intrare);

- adăugarea înregistrărilor este mult mai rapidă, dacă indexul poate fi păstrat în memoria principală;
- în index, se poate localiza rapid o cheie, utilizând căutarea binară;
- pornind de la cheie și adresă, pentru recuperarea unei înregistrări este necesară o singură accesare a fișierului de date.

Astfel, dacă înregistrările unui fișier sunt de lungime variabilă, este dificilă utilizarea metodelor de ordonare și aplicarea căutării binare. Aici o alternativă poate fi construirea unui index. Structura indexului, în acest caz, poate fi foarte simplă: indexul este un fișier cu înregistrări de lungime fixă, în care fiecare are două câmpuri, un câmp pentru cheie și altul pentru poziția inițială a înregistrării în fișierul de date. Fiecărui câmp cheie al fișierului de date îi corespunde un câmp cheie în index. Indexul este ordonat, în timp ce fișierul de date nu este. Fișierul de date poate fi organizat în ordinea intrării înregistrărilor.

Dacă indexul nu încapă în memoria primară, accesarea și menținerea lui trebuie făcută în memoria secundară. În acest caz, utilizarea indecșilor simpli este problematică, deoarece:

- căutarea binară poate necesita multe accesări ale discului;
- deplasarea înregistrărilor, după operațiile de inserare și eliminare, devine inevitabilă pentru menținerea indexului.

În acest context, dacă viteza de acces este caracteristica dezirabilă principală, indexul poate fi organizat, utilizând tehnici cu funcții de dispersie. Dacă nu, B-arborele poate fi o structură acceptabilă atunci când se cer accesări după cheie sau accesări secvențiale eficiente.

Indecși ordonați

Indexul ordonat păstrează ordinea cheilor de căutare și oricărei chei îi asociază adresele înregistrărilor care o conțin. Indexul ordonat poate fi primar sau secundar.

Indecși primari. Indexul ordonat este un *index primar*, dacă fișierul de date și indexul au același criteriu de ordonare, adică, dacă fișierul este ordonat de un câmp cheie, indexul care se definește pe acest câmp este un index primar. Deseori, termenul index primar este utilizat pentru desemnarea unui index al cheii primare (a unei relații din baza de date), dar această formă nu este un standard și poate fi evitată.

Dacă toate fișierele sunt ordonate secvențial după o singură cheie de căutare, aceste fișiere cu un index primar pe cheia de căutare sunt numite *fișiere indexat secvențiale*. Fișierele indexat secvențiale reprezintă una din cele mai vechi scheme de index utilizate în bazele de date. Ele sunt proiectate pentru aplicațiile care necesită atât prelucrarea secvențială a fișierelor, cât și accesul aleatoriu la înregistrări individuale.

Un fișier index primar poate fi:

- un index dens, dacă pentru orice valoare a cheii de căutare din fișierul de date există o înregistrare în fișierul index (sau intrare index) respectivă. Evident, înregistrarea index conține valoarea cheii de căutare și un pointer spre prima înregistrare de date cu această valoare a cheii de căutare.
- un index rar, dacă înregistrările fișierului de date sunt grupate după un criteriu și pentru fiecare grup (organizat în bloc) există o înregistrare index asociată. Astfel, fiecare înregistrare index are o valoare a câmpului cheie egală cu valoarea cheii primei

înregistrări de date a blocului respectiv și un pointer spre acest bloc.

Pentru inserarea datelor în poziția corectă în fișierul de date este necesară, în afară de deplasarea înregistrărilor pentru a deschide spațiu pentru cea nouă, modificarea unor înregistrări ancoră (prima înregistrare a blocului). Ca și în indecșii denși, fiecare înregistrare a indexului rar conține o valoare a cheii de căutare și un pointer spre prima înregistrare de date cu această valoare a cheii.

Pentru un fișier index rar, localizarea unei înregistrări după un argument de căutare se face în două etape: în prima este consultat indexul și determinat blocul în care trebuie să se găsească înregistrarea și în a doua - blocul selectat este căutat și localizată în el înregistrarea dorită. Adică, pentru localizarea unei înregistrări, se obține intrarea indexului cu cea mai mare valoare a cheii, mai mică sau egală cu valoarea cheii căutate. Adresa din acea înregistrare permite accesul la blocul care poate conține înregistrarea căutată. Apoi, în blocul respectiv, se face o căutare secvențială sau prin altă metodă.

Deoarece în indexul dens există o înregistrare care indică direct înregistrarea căutată în fișierul de date, se poate crede că recuperarea unei înregistrări cu un index dens este mai eficientă decât a uneia cu un index rar. Însă, indecșii denși sunt prea mari și ocupă mai mult spațiu, fapt ce complică sarcina de actualizare (inserare și suprimare) a lui. În afară de aceasta, indecșii rari explorează bine caracteristicile fizice ale dispozitivelor secundare. Astfel, fiecare bloc definit de indexul rar poate fi încărcat în memoria principală dintr-o singură accesare și manipulat în memoria principală, consumând un timp nesemnificativ. Prin urmare, când trebuie să facă o alegere, proiectantul trebuie să găsească un compromis între timpul de acces și spațiul suplimentar, pentru a lua o decizie mai bună, care este influențată direct de aplicația în cauză.

Indecși secundari. Indexul asociat cheii de ordonare a fișierului de date este numit, în afară de index primar, și index *cluster*. Alți indecși, ale căror chei de căutare specifică o ordine diferită de ordinea secvențială, se numesc *indecși secundari*. În ambele cazuri, intrările indexului sunt ordonate de valoarea de acces, având ca obiectiv creșterea eficienței căutării.

Astfel, dacă un fișier nu este ordonat, orice index care se definește asupra lui este un index secundar. De asemenea, un index secundar este indexul definit pe un câmp diferit de câmpul de ordonare al fișierului.

Indexul secundar poate fi dens sau rar, având aceleași specificații ca și indexul primar. Indexul secundar, ca și indexul primar, este un fișier ordonat cu două câmpuri: câmpul de indexare (orice câmp al fișierului de date diferit de cel de ordonare) și un pointer spre un bloc sau spre o înregistrare. Dacă indexul secundar este definit pe un *câmp cheie* (*cheie secundară*), există câte o înregistrare index pentru fiecare cheie a fișierului, adică indexul secundar este un index dens.

Cu introducerea indecșilor secundari, a dispărut necesitatea de pointeri adiționali în înregistrări și necesitatea de parcurgere secvențială în ordinea cheii secundare. Însă utilizarea acestei structuri îl impune pe administratorul bazei de date să analizeze mulțimea potențială de interogări pentru a selecta câmpurile asupra cărora se vor defini indecșii secundari.

Indecși multinivel. Pentru fișierele cu multe înregistrări, indexul este, totuși, destul de mare și nu poate fi încărcat în memoria principală. De aceea sunt utilizați indecși cu multe niveluri, adică indecși cu niveluri rare.

Un index multinivel este format pentru un fișier index, care se numește primul nivel sau nivel-bază al indexului multinivel. Primul nivel este un fișier ordonat cu o valoare distinctă a câmpului de indexare în fiecare intrare. De aceea, asupra primului nivel se poate crea un index primar. Acest index constituie al doilea nivel al indexului multinivel. Deoarece al doilea nivel este un index primar, în el există câte o intrare în fiecare bloc al primului nivel. Procesul poate continua și asupra acestui nivel, dacă este necesar. Atunci al treilea nivel va fi un index primar pentru nivelul doi.

Această schemă multinivel se poate utiliza asupra oricărui tip de index, fie primar, fie secundar, cu condiția ca, întotdeauna, primul nivel să ia valori distincte în câmpul de indexare și intrările să fie de lungime fixă. Astfel, dacă fiecare bloc se presupune că are r intrări, adică fiecare bloc reduce numărul de intrări de la nivelul anterior de r ori, atunci un index multinivel cu i intrări, la primul nivel, va avea aproximativ n niveluri, unde $n = \lceil \log_r i \rceil$.

Indecși cu dispersie

Se pot crea structuri de acces similare indecșilor, dar bazate pe dispersii. Astfel, intrările indexului (K, Pr) se pot organiza ca un fișier dispersat, care își va schimba dimensiunea, folosind dispersia dinamică, extensibilă sau liniară. Algoritmul aplică funcția de dispersie asupra cheii K . Îndată ce este prezentată o intrare (o cheie), pointerul Pr se utilizează pentru localizarea înregistrării în fișierul de date.

Funcția de dispersie asociată poate fi utilizată atât pentru organizarea fișierelor, cât și pentru crearea structurii indexului.

28.5 Fișiere indexat secvențiale

Fișierele indexat secvențiale sunt un tip de organizare a fișierelor în care înregistrările sunt identificate de un index, numit cheie de acces și fiecare cheie reprezintă valoarea cu care sunt identificate alte date ale înregistrării în fișier. Această cheie de acces trebuie să fie unică pentru fiecare înregistrare, adică nu sunt admise două înregistrări cu aceeași valoare a cheii. În afară de aceasta, toate fișierele (inclusiv indexul) sunt ordonate secvențial după cheia de acces, care, de fapt, este cheia primară a fișierului cu date.

Structuri ale fișierelor indexat secvențiale

Întrucât volumul datelor dintr-un fișier secvențial poate fi foarte mare și numărul de accesări (cu o executare joasă) înalt, se utilizează alte tehnici pentru manipularea eficientă a înregistrărilor. Astfel, se folosește o structură de acces, asociată la fișier, care oferă o mai mare eficiență în localizarea înregistrării prezentate de argumentul de căutare, decât metodele aplicate în fișierele secvențiale.

Un fișier indexat secvențial este constituit dintr-un fișier secvențial și un index (structura de acces). Structura lui constituie trei spații de lucru: spațiul pentru date (fișier principal), spațiul rezervat înregistrărilor excedentare (fișier de tranzacții) și spațiul pentru index (fișier index). Fișierul index poate fi organizat în niveluri, de la indexul principal (rădăcina), până la nivelul de acces la înregistrare.

Astfel, pentru accelerarea căutării înregistrărilor, fișierul indexat secvențial utilizează fișierul de date și fișierul index. Fișierul principal (de date) este de tip secvențial, cu înregistrări grupate în blocuri și ordonate de o cheie de ordonare. Fișierul index este format din perechi <cheie, adresă>. Adresa indică înregistrări sau blocuri în fișierul principal sau indică blocuri de înregistrări în propriul fișier index (formând un arbore). Adică, pot fi mai multe niveluri de indecși, ei fiind stocați în același fișier fizic sau în fișiere diferite. Când se prelucrează blocurile, perechea <cheie, adresă> (fie bloc de date sau bloc de index) care arată spre următorul nivel conține cheia ce corespunde înregistrării cu valoarea cea mai mare (variantă mai comună) în bloc.

Avantajul utilizării structurilor în blocuri constă în faptul că blocurile sunt citite și încărcate în memoria centrală în întregime. Dimensiunea blocului trebuie determinată, ținându-se cont de caracteristicile dispozitivelor fizice de lectură și scriere și limitele sistemului utilizat (responsabil de operațiile de lectură și scriere pe dispozitivele fizice).

O strategie bună este rezervarea în fiecare bloc al fișierului a unui spațiu liber pentru înregistrările noi ce vor fi incluse. Astfel, la prima încărcare a fișierului de date și la fiecare reorganizare a fișierului (rearanjare), este ocupată numai o parte din fiecare bloc. Pentru a evita umplerea rapidă a blocurilor, spațiul liber este rezervat, ținându-se cont de statisticile de includeri de înregistrări, iar reorganizarea este o procedură ce trebuie făcută periodic.

O alternativă mai potrivită este utilizarea unui singur spațiu pentru înregistrările excedentare, deoarece nici o dată nu se poate prezice când vor fi completate unele blocuri. Și dacă această problemă apare, ceva trebuie întreprins pentru a include o înregistrare nouă, care nu poate fi plasată imediat în poziția adecvată în fișier (această operație necesită reorganizarea fișierului care se execută în momente speciale).

Pentru acumularea înregistrărilor excedentare, una din cele mai bune opțiuni este crearea unui fișier de tranzacții (*overflow*). Acest fișier are aceeași structură ca și fișierul de date, cu înregistrări suplimentate cu un câmp care permite înlănțuirea înregistrărilor. Înregistrările excedentare sunt incluse în fișierul de tranzacții (prima poziție vidă) și formează o listă înlănțuită (crescător după cheia de ordonare).

Fișierul de tranzacții poate fi utilizat în diverse moduri. O variantă este folosirea unei liste de înregistrări excedentare pentru fiecare bloc al fișierului de date, figura 28.10. De aceea, în fișierul principal există câte un pointer asociat fiecărui bloc. El indică o înregistrare în fișierul de tranzacții, care este prima în lista de înregistrări excedentare ale blocului. Această listă poate fi formată din cheile mai mari decât cea mai mare din bloc (și mai mici decât cea mai mică a blocului următor), sau, din chei mai mici decât cea mai mică din bloc (și mai mari decât cea mai mare a blocului anterior). Aici trebuie ținut cont că cea mai mare cheie a blocului nu poate fi exclusă fizic, deoarece este utilizată în nivelul intern al indecșilor.

Această variantă necesită reorganizarea înregistrărilor între blocul din fișierul de date și lista de înregistrări excedentare, păstrând, după inserări și suprimări, ordinea cheii în bloc și în listă (deplasând înregistrările respective din bloc spre listă sau din listă spre bloc). Avantajul acestei structuri este o mai bună utilizare a spațiului în blocuri, deoarece spațiul din fișier ocupat de înregistrările excluse (logic) poate fi ocupat de alte înregistrări numai după reorganizare.

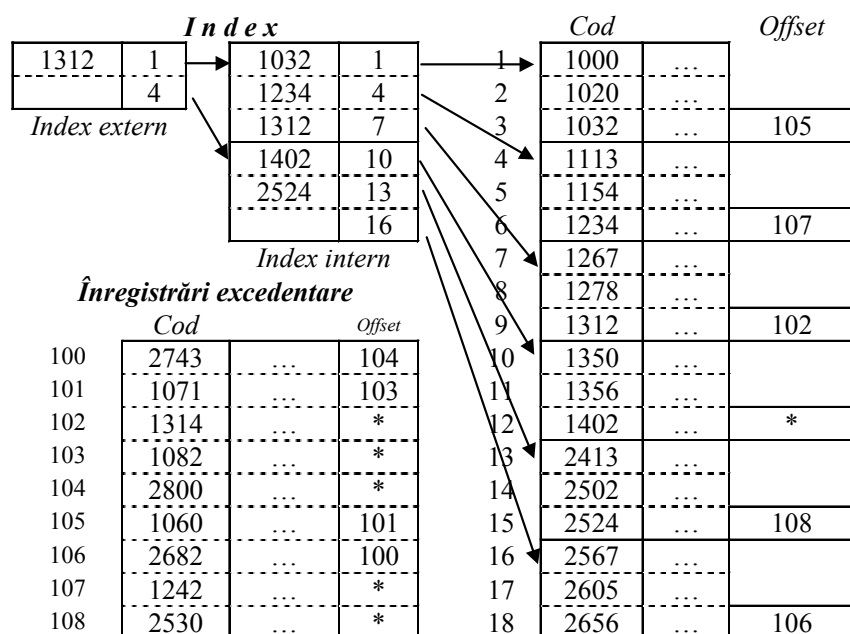


Figura 28.10 Înregistrări excedentare asociate cu blocuri

Există și alternativa când listele cu înregistrări excedentare se asociază cu înregistrările din fișierul principal și nu cu blocurile, figura 28.11. Astfel, fiecare înregistrare are câte un pointer nul sau care indică spre lista respectivă de înregistrări excedentare. Această listă este constituită din înregistrări cu chei mai mici decât a înregistrării asociate și mai mari decât a înregistrării anterioare. În această structură, eliminările de înregistrări în blocuri vor fi făcute numai la nivel logic (cu includerea unei etichete).

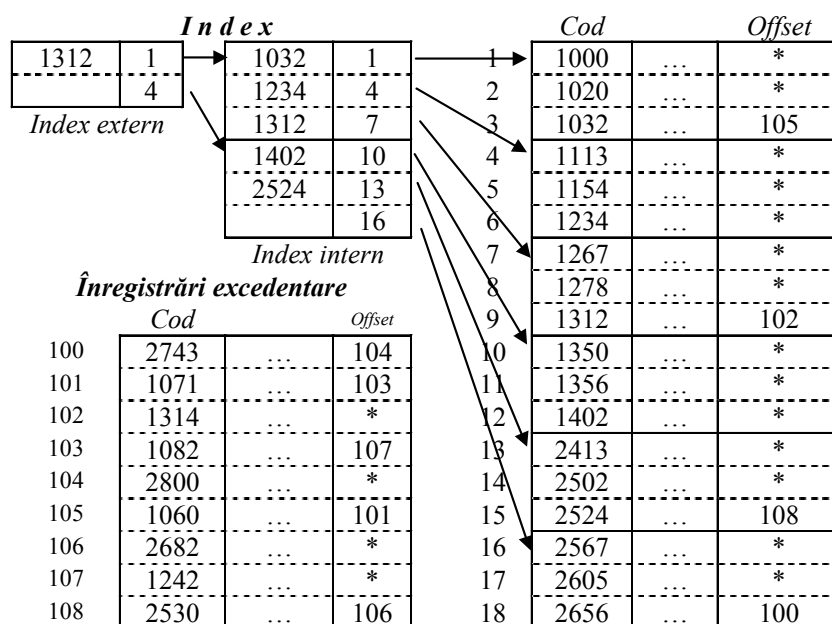


Figura 28.11 Înregistrări excedentare asociate cu înregistrări

Avantajele organizării indexat secvențiale sunt cele ale organizării directe și organizării secvențiale. Adică, un avantaj al acestei organizări este posibilitatea de a îmbina consultarea rapidă a unei înregistrări (prin index) cu parcurgerea facilă secvențială a tuturor înregistrărilor (după cheia de ordonare).

Cu toate acestea, structura nu este adecvată pentru consultarea după mai multe chei (în afară de cea de ordonare) și pentru sistemele în care operațiile de actualizare sunt făcute 24 ore pe zi, în toate zilele. În sistemele, unde sunt multe inserări (cu multe liste de înregistrări excedentare), are loc degradarea generală a executării operațiilor din cauza deplasărilor de înregistrări între blocuri și liste în momentul inserărilor sau eliminărilor (prima alternativă), sau din cauza consultărilor exhaustive în listele cu înregistrări excedentare (alternativa a doua).

Astfel, indexul trebuie actualizat încontinuu, pe măsură ce înregistrările sunt adăugate sau eliminate. În afară de aceasta, în raport cu organizarea directă, sunt necesare etape suplimentare pentru localizarea înregistrărilor - indexul trebuie să fie utilizat pentru orice accesare.

Fișierele indexat secvențiale și dispozitivele de stocare

Metoda indexat secvențială presupune că înregistrările sunt stocate pe dispozitive cu acces direct, iar pentru operațiile de acces direct există indecși care permit căutarea adresei unei înregistrări individuale. În cazul în care se solicită accesul aleatoriu, folosind diverse chei de acces, trebuie construit un index pentru fiecare din ele. În general, indecșii au o structură ierarhică (pe niveluri), pentru o localizare rapidă a pistei care conține înregistrarea căutată.

Pentru a profita de caracteristicile dispozitivelor secundare de stocare a datelor și a face mai eficient procesul de căutare, este indicat să se construiască indecși structurați, în mai multe niveluri, în acord cu nivelurile de organizare fizică a acestor dispozitive. Numărul de niveluri este

proporțional cu numărul de intrări (numărul de indecși). Acești indecși pot fi implementați, de exemplu, în formă de B-arbori, datorită flexibilității mari a acestei structuri în ceea ce privește operațiile de inserare și eliminare în fișierul index și eficienței de căutare pe care o oferă.

Astfel, fișierul index se construiește structurat în acord cu nivelurile fizice ale modului de stocare pe discurile magnetice. Nivelurile fizice sunt discuri, cilindri, pista și sectorul. Indexul de cilindru conține pentru fiecare cilindru ocupat de fișier, o intrare ce indică cea mai mare valoare stocată a cheii de ordonare. Cu fiecare cilindru este legat un index de pistă, care indică, pentru fiecare pistă a cilindrului, cea mai mare valoare a cheii care o conține. Indexul de pistă al unui cilindru este stocat pe cilindru. Indexul de nivel mai înalt, asupra indexului de cilindru, este numit index principal și conține, pentru fiecare disc (sau grup de cilindri) ocupat de fișier, valoarea cea mai mare a cheii stocate.

Indexul principal și de cilindri, când este posibil, se păstrează în memoria principală, astfel ca, în procesul de căutare, cilindrul, unde se găsește înregistrarea căutată, să fie identificat, fără accesarea discului. După obținerea cilindrului, acesta este accesat pentru lectura indexului lui de piste, care permite localizarea pistei, unde se găsește înregistrarea. Pista selectată este, atunci, citită și în ea este localizată (de o metodă de căutare) înregistrarea solicitată. În acest proces, numai un cilindru este accesat. E important de menționat că indexul de pistă nu precizează (nu are nevoie) conținutul pentru fiecare înregistrare sau adresă a pistei, întrucât el este determinat de secvența în care apare la intrarea în index. Același lucru se întâmplă și cu indexul cilindrilor.

Chiar dacă nu se manipulează direct structura *disc/cilindri/piste* este importantă structurarea indexului în niveluri, în primul rând, pentru că nivelurile externe (care nu sunt direct legate cu fișierul de date) antrenează puține intrări, ceea ce permite să fie stocate în memoria principală, astfel devenind eficientă căutarea în indexul primar (numai un bloc), fapt ce, la rândul său, eficientizează căutarea fișierului de date (numai un bloc).

Operații cu fișiere indexat secvențiale

Accesul aleatoriu la înregistrări este realizat, folosind indexul. Argumentul de căutare definește calea în index, care duce spre adresa înregistrării solicitate. Adresa obținută de index poate fi adresa înregistrării sau adresa blocului care o conține. În ultimul caz, pentru localizarea înregistrării, este necesară efectuarea unei căutări în blocul respectiv care, la rândul său, poate cere accesarea spațiului rezervat înregistrărilor excedentare.

Fișierele indexat secvențiale trebuie, de exemplu, să fie reorganizate când înregistrările excluse logic nu sunt refolosite de algoritmul de includere a înregistrărilor. În acest caz, spațiul eliberat de înregistrările excluse trebuie să fie ocupat după reorganizarea fișierului, care presupune și o reorganizare completă a fișierului index, deoarece pozițiile fizice ale unor înregistrări vor fi modificate.

Deschiderea și închiderea fișierelor. Deschiderea unui fișier indexat secvențial implică identificarea fișierului de date, cheilor de acces pentru care există indecși și care sunt fișierele index, fără a cunoaște modul cum sunt structurate acestea, câte niveluri are fișierul index asociat. De multe ori, indecșii pot fi creați numai când se utilizează fișierul (adică indexul nu e

stocat pe disc). Prin urmare, întotdeauna când un fișier este deschis, trebuie să se creeze în memorie structurile respective de acces.

Închiderea fișierului este, de obicei, o sarcină simplă și depinde de faptul cum se lucrează cu el. Dacă indecșii sunt manipulați numai în memoria principală, fișierul de date se poate închide fără nici o operație adițională. Cu toate că manipularea indecșilor are loc în memoria primară, deseori (cazul când această structură este modificată), este necesară salvarea fișierului index pe disc, or generarea fișierelor index poate fi o sarcină foarte retardă. În acest caz, este important să se cunoască care indecși sunt temporari și care indecși sunt permanenți. De asemenea, trebuie prelucrat spațiul de înregistrări excedentare, pentru că se recomandă ca fișierul închis să aibă spațiul de tranzacții liber. Prin urmare, mai întâi se face reorganizarea fișierului, apoi se închide.

Consultarea înregistrărilor. Căutarea unei înregistrări după cheia de ordonare este făcută mai întâi în fișierul index (pornind de la nivelul extern). Se caută cheia imediat mai mare sau egală cu cea solicitată. Cu adresa obținută se caută în nivelul următor al indexului, citind un bloc de înregistrări din index. În acest bloc, se caută cea mai mare cheie. Se repetă căutarea până nu mai există niveluri de indecși și adresa obținută corespunde unui bloc de înregistrări în spațiul de date al fișierului principal. Apoi se citește acest bloc, căutând cheia cerută. Dacă este găsită, trebuie să se analizeze eticheta de eliminare (dacă eticheta există, atunci cheia nu mai există). Dacă cheia nu se întâlnește în acest bloc, se verifică dacă se află în lista de înregistrări excedentare, fie în bloc sau înregistrare. Dacă este, se ia adresa indicată (corespunzătoare unei poziții în fișierul de tranzacții) și se parcurge lista în căutarea cheii cerute. În general, în fișierul de tranzacții nu se lucrează cu etichete de eliminare, deoarece aici eliminările sunt, de obicei, fizice.

Consultarea unei înregistrări după cheia care nu e de ordonare necesită o căutare exhaustivă a tuturor înregistrărilor atât în blocuri, cât și în listele de tranzacții.

Inserarea înregistrărilor. Pentru inserarea unei înregistrări, inițial se caută în fișierul index și se determină în ce bloc de date înregistrarea trebuie să fie inclusă. Îndată după determinarea poziției, înregistrarea este inclusă în spațiul de înregistrări excedentare și se actualizează legăturile, dacă spațiul de tranzacții există pentru fiecare înregistrare.

Dacă spațiile de înregistrări excedentare sunt rezervate blocurilor, înregistrarea este inserată în blocul selectat în spațiul rezervat pentru date. În cazul când acesta este plin sau nu există spațiu pentru inserarea înregistrării respective, astfel ca înregistrările acestui fișier să păstreze ordinea secvențială, este inserată în spațiul de tranzacții al blocului, dacă cheia sa este mai mare decât cheia ultimei înregistrări a blocului principal și se actualizează legăturile.

În caz contrar, ultima înregistrare a blocului principal este trecută în spațiul de tranzacții și înregistrarea nouă este inserată în poziția sa corectă din blocul principal (pot fi deplasări înăuntru blocului). Poate exista un câmp similar, în acest spațiu, care ar indica următoarea adresă în spațiul de date sau de tranzacție.

În unele implementări inserările pot fi realizate în blocuri, în spațiile eliberate de înregistrările care au fost excluse sau în spațiile rezervate pentru aceasta la sfârșit. În acest caz, spațiile pentru înregistrările excedentare pot să nu fie folosite. Aici, sistemul localizează în index blocul

unde va fi stocată înregistrarea. După această operație, se găsește, în fișierul de date, blocul dezirabil și se include înregistrarea în poziția corectă, adică în poziție ordonată.

Înregistrarea este scrisă pe prima adresă liberă din fișier și sunt actualizate pozițiile relative ale acestei adrese (cheia înregistrării, adresa înregistrării), apoi e scrisă în index (indecși). Înregistrarea nouă poate fi stocată pe orice adresă vidă din spațiul alocat pentru fișier. În continuare, trebuie inclusă o intrare în fișierul index primar (ce se referă la cheia primară) și în fiecare din celelalte fișiere index care există. Inserările în fișierele index atrag după sine restructurarea lor pentru ca înregistrarea să fie scrisă în poziția adecvată în fișierul index.

Excluderea înregistrărilor. Ca și în alte organizări analizate anterior, în fișierele indexat secvențiale se pot realiza atât excluderi fizice, cât și logice. Excluderea fizică presupune eliberarea spațiului în fișierul de date (poate dura mult timp), precum și eliminarea înregistrării din fișierul de index.

O soluție mai bună, în majoritatea cazurilor, este excluderea logică. În acest caz, se poate opta pentru mai multe variante. Este posibilă includerea unui câmp în fișierul de date care ar indica dacă înregistrarea a fost exclusă. Astfel, nu e necesară modificarea indexului, însă devine obligatorie verificarea acestui câmp la fiecare accesare a înregistrării.

Altă opțiune este utilizarea în index a unui indicator de excludere asociat cheii primare. În acest mod, la accesarea fișierului index deja se va putea verifica dacă înregistrarea a fost exclusă. În afară de aceasta, înregistrarea ar putea fi detectată de algoritmul de inserare, care ar localiza ușor spațiile eliberate pentru inserări.

Dacă fișierul posedă mai multe structuri index (asupra diferitelor chei de acces), sarcina de excludere a înregistrărilor se poate complica. În orice caz, pentru orice structură index, înregistrarea exclusă nu trebuie luată în seamă.

Modificarea înregistrărilor. Pentru modificarea înregistrărilor, mai întâi, se localizează înregistrarea ca într-o consultare (prin structura index), folosind un argument de căutare în calitate de cheie de acces. În continuare, se modifică câmpurile respective și se înscrie înregistrarea.

Pot fi întâlnite două situații. Dacă nu se modifică câmpul cheii de ordonare și lungimea înregistrării este mai mică, înregistrarea poate fi scrisă în locul în care se găsește. Dacă modificarea datelor înregistrării cere modificarea poziției ei în fișier, adică are loc modificarea unui câmp al cheii de ordonare sau a lungimii ei, atunci înregistrarea se exclude din fișierul inițial și se inserează, în calitate de înregistrare nouă, în fișierul de tranzacții.

Astfel, sistemul de operare localizează în index adresa blocului unde poate fi stocată cheia căutată. Dacă blocul este găsit în fișierul de date, înregistrarea este căutată secvențial. Modificarea unei înregistrări este făcută cu o nouă înscriere a înregistrării actualizate, în cazul când nu se modifică dimensiunea ei, pe aceeași adresă. Iar în cazul când lungimea este modificată, înregistrarea veche este exclusă, iar cea nouă este inserată pe altă adresă, unde înregistrarea nouă poate fi scrisă.

Modificarea unei înregistrări poate provoca modificarea indecșilor, deoarece aceasta poate modifica valoarea câmpurilor pentru care există indecși. Acest fapt implică eliminarea intrărilor anterioare în fișierul de index și în includerea unei intrări noi, ce va implica o reorganizare a indexului. De

obicei, nu se admite modificarea cheii specificate pe câmpurile cheii primare.

Lectura exhaustivă a înregistrărilor. Lectura exhaustivă este realizată fără utilizarea indexului. Se citește fiecare înregistrare din fișierul principal și continuă această lectură, folosind lista de pointeri ce leagă fișierul principal de cel de tranzacții, în mod analogic, ca în fișierul secvențial.

Însă, din cauza accesării spațiului de înregistrări excedentare, această operație este mai puțin eficientă, în comparație cu operația respectivă în fișierul secvențial. Ordinea fizică și logică a înregistrărilor nu este aceeași. Deci, pentru a realiza eficient lectura exhaustivă a fișierului, este necesară definirea unui index asupra acestui fișier care ar determina ordinea logică de parcurgere a înregistrărilor. Dacă există cel puțin un index exhaustiv asupra fișierului, această operație poate fi efectuată prin accesul serial la fișierul index, urmată de accesări directe ale fișierului de date.

Reorganizarea fișierelor. Reorganizarea presupune o lectură exhaustivă și transferarea tuturor înregistrărilor într-un fișier nou de date. De fapt, toate înregistrările sunt transferate în fișierul principal și spațiul de tranzacții este eliberat. În continuare, se generează indexul nou.

Reorganizarea unui fișier indexat secvențial devine întotdeauna necesară, dacă: spațiul de tranzacții devine foarte mare, diminuând, astfel, eficiența executării accesărilor de fișiere. În afară de aceasta, are loc eliminarea înregistrărilor excluse logic (dacă aceste sectoare nu au fost folosite de algoritmul de inserare a înregistrărilor). Dacă trebuie, se restructurează fișierele index pentru îmbunătățirea accesului la înregistrări.

Astfel, reorganizarea fișierelor constă în lectura tuturor înregistrărilor și scrierea lor în alt fișier de date. Totodată, se eliberează spațiile ocupate de înregistrările excluse și se generează o structură nouă de indecși.

Operația de reorganizare trebuie făcută periodic când sistemul nu este utilizat (cel puțin când nu sunt făcute inserări, modificări și suprimări), deoarece reorganizarea necesită o analiză exhaustivă a înregistrărilor în două fișiere (principal și de tranzacții) și crearea unui fișier auxiliar cu aceeași structură, ca a fișierului principal.

28.6 Fișiere indexate

În fișierele indexate, înregistrările sunt identificate de un index, numit cheie de acces, sau cheie a înregistrării, sau cheie principală. Fiecare cheie reprezintă valoarea cu care celelalte date din înregistrare sunt identificate în fișier. Această cheie de acces trebuie să fie unică pentru fiecare înregistrare, adică, nu există două înregistrări cu aceeași valoare a cheii. Înregistrările sunt accesate întotdeauna de unu sau mai mulți indecși, care nu au nici o legătură cu ordinea de stocare fizică a înregistrărilor.

Păstrarea și respectarea ordinii fizice a înregistrărilor (pentru un acces serial eficient) în fișierele indexat secvențiale provoacă mai multe probleme, în principal, probleme care țin de inserarea înregistrărilor noi. Pe măsură ce necesitatea de accesări seriale scade, în raport cu numărul de accesări aleatorii, păstrarea secvenței fizice a fișierului afectează negativ eficiența accesării bazei de date, adică organizarea fișierelor devine contraproductivă.

Prin urmare, devine mai convenabilă utilizarea unui fișier indexat, în care înregistrările sunt accesate întotdeauna prin unu sau mai mulți indecși,

fără a ține cont de ordinea fizică a înregistrărilor. Indecșii pot fi structurați ca în organizarea indexat secvențială, ceea ce face mai eficientă căutarea înregistrării.

Structuri și operații în fișierele indexate

Fișierele indexate, ca și fișierele indexat secvențiale, utilizează structuri auxiliare (indecși) pentru accelerarea accesului la înregistrările cu date. Spre deosebire de precedentele organizări, în fișierele indexate există posibilitatea de utilizare a mai multor indecși (respectiv, mai multor chei) și lipsește restricția că înregistrările în fișierul principal trebuie să fie ordonate de o cheie. Ultima caracteristică exclude utilizarea spațiilor de tranzacții.

Astfel, fișierul cu organizare indexată este constituit din două componente:

- indexul, care conține cheile de acces și adresele fiecărei înregistrări din fișierul de date; toate accesările de înregistrări sunt făcute prin această componentă, care este păstrată întotdeauna ordonat, după cheia de acces;
- fișierul de date, care stochează celelalte câmpuri ale tuturor înregistrărilor fișierului, independente de orice ordonare sau secvență.

Ca și fișierele indexat secvențiale, fișierele indexate susțin două metode de acces: secvențial și aleatoriu.

Modul secvențial de acces reprezintă o lectură secvențială a indexului unui fișier indexat. Pentru fiecare intrare în index, este obținută adresa înregistrării din index în fișierul de date și cu această adresă se citește înregistrarea. În recuperarea secvențială a unui fișier, datele sunt citite întotdeauna ordonate după cheie.

Modul aleatoriu de acces se caracterizează prin inexistența unei secvențe logice pentru recuperarea înregistrărilor, adică, după citirea ultimei înregistrări, poate fi citită a cincea sau prima, sau oricare alta - această secvență este determinată de programul aplicație. În procesul de lectură, sistemul localizează valoarea cheii dezirabile în index, identificând adresa înregistrării corespunzătoare în fișierul de date, și asigură accesul, cu această adresă, pentru recuperarea înregistrării.

Consultarea fișierelor indexate poate fi realizată, folosind mai multe chei. Dacă există un index pentru o cheie, căutarea este făcută, pornind de la indexul în chestiune (prelucrarea ca în organizarea indexat secvențială) până se obține adresa unei înregistrări în fișierul de date. Nu este nevoie de multe liste (nu mai sunt accesări). Însă, dacă excluderile înregistrărilor sunt logice, este necesară examinarea existenței etichetelor de excludere. Dacă nu există vreun index pentru cheia cerută, se realizează căutarea exhaustivă a cheii în fișierul de date (analiza tuturor înregistrărilor până ce se întâlnește cea căutată).

Inserările sunt făcute întotdeauna în fișierul de date, în orice spațiu liber, deoarece înregistrările în acest fișier nu trebuie să respecte o anumită ordine. Pentru a putea reexploata spațiile înregistrărilor suprimate, se construiește o listă de spații libere. Însă, fiecare inserare implică și actualizarea imediată (*on-line*) a indecșilor, fapt ce poate influența negativ executarea acestei operații.

Modificările sunt făcute direct în fișierul de date și ele, de asemenea, provoacă actualizarea indecșilor (numai în care se indică câmpurile modificate).

Excluderea înregistrărilor poate fi fizică, cu actualizarea *on-line* a indexului, sau logică, cu utilizarea etichetelor. Ultima variantă implică, ulterior, ștergeri fizice cu scopul reutilizării spațiului eliberat.

Parcursul secvențial este puțin ineficient, dar nu influențează executarea sistemului în general. Ordinea înregistrărilor depinde de cheia aleasă și este dictată de informațiile din index (nivelul cel mai intern conține ordinea înregistrărilor). Problema este că fiecare pereche <cheie, adresă>, în ultimul nivel direcționează accesul la fișierul de date (deoarece înregistrările sunt dispersate fără ordine și nu sunt utilizate blocuri de înregistrări continue). Pentru parcursul secvențial, este de ajuns analiza directă a ultimului nivel al indexului (nu este necesară navigarea, pornind de la nivelurile externe). Parcursul secvențial după o cheie care nu este index necesită un fișier auxiliar și tehnici de ordonare.

Un mare avantaj al fișierelor indexate este posibilitatea de consultare după mai multe chei, ceea ce deosebește această organizare de celelalte. Consultările aleatorii (pentru o singură înregistrare), de asemenea, sunt făcute rapid. Acum operația de parcurs secvențial, cu toate că nu poate fi făcută în mod optimal, nu este foarte lentă ca operația respectivă în organizarea indexată secvențială, întrucât nu trebuie să se parcurgă liste înlanțuite.

Un dezavantaj este necesitatea de actualizare *on-line* a indecșilor, ceea ce poate influența executarea operațiilor, dar, pe de altă parte, permite aplicarea lor în sistemele cu utilizare neîntreruptă, cum sunt bazele de date. Însă, o dată cu implementarea indecșilor eficienți și moderni (precum se va vedea în continuare), se poate soluționa această problemă. Alt dezavantaj este că organizarea indexată cere mult spațiu pentru stocarea tuturor indecșilor. Cu toate acestea, fișierele indexate sunt larg utilizate în majoritatea SGBD-urilor.

Fișiere inversate

Nu toate căutările în bazele de date se fac după cheia primară. Dacă un atribut sau un grup de atribut apare des în cereri, atunci, pentru acest atribut sau grup de atribute, se construiește un index secundar, ce permite accesul rapid la înregistrările corespunzătoare valorilor date.

Un fișier cu un index secundar, corespunzător unui atribut sau grup de atribute *X*, se spune că este *fișier inversat* în raport cu *X*. În indexul secundar, înregistrările sunt indicate fie prin pointeri la ele, fie prin valorile cheilor primare corespunzătoare lor.

Referirea la înregistrări prin pointeri (listă inversată) are avantajul accesului mai rapid la date, dar produce constrângeri din cauza fixării înregistrărilor în locul unde au fost introduse, la nivel de bloc sau la nivel de înregistrări. Referirea prin cheia principală asociată are dezavantajul unui acces lent, dar nu mai fixează înregistrările.

De exemplu, o listă inversată după câmpul *Departament* al unui fișier de funcționari ai unei întreprinderi va conține toate valorile acestui câmp, și fiecareia din aceste valori îi va fi asociată o listă cu adrese sau pointeri ale înregistrărilor funcționarilor care conțin valoarea specificată în câmpul

Departament. Astfel, lungimea fiecărei liste va fi numărul de funcționari care lucrează în departamentul respectiv.

De fiecare dată când o înregistrare este inclusă, exclusă sau modificată, fișierul inversat, de asemenea, trebuie să fie modificat. O alternativă, precum s-a menționat, este stocarea nu a adresei, ci a cheii primare a fiecărei înregistrări. În acest mod, se evită actualizarea indexului când are loc reorganizarea fizică a fișierului (ordinea înregistrărilor). Însă, această opțiune frânează executarea consultărilor, o dată ce nu se cunoaște adresa, este necesară localizarea ei în alt mod (sau de alt index, sau secvențial).

O problemă ce diminuează calitățile necesare pentru viabilitatea fișierelor inversate apare când valorile câmpului pe care e definit indexul secundar nu posedă repetări de valori. Altă problemă ce trebuie să fie considerată este că listele, în general, nu au lungime fixă (numărul de înregistrări pentru fiecare valoare se poate schimba cu utilizarea sistemului).

Indecși Bitmap

Un index *bitmap* este o variantă a indexului secundar care asociază fiecărei valori a atributului indexat o secvență de biți, în loc de o listă de adrese sau chei. Fiecare bit identifică o înregistrare concretă a fișierului. Dacă înregistrarea în câmpul atributului indexat are valoarea specificată, atunci bitul respectiv conține valoarea 1. Un bit cu valoarea 0 semnifică faptul că înregistrarea nu conține valoarea specificată în câmpul indexat.

Lungimea secvenței de biți este numărul maximal de înregistrări în fișier. Dacă fișierul crește mult sau are loc reorganizarea fizică a înregistrărilor, indexul *bitmap* trebuie să fie refăcut.

Dacă se caută înregistrările cu valoarea v a atributului indexat, este suficient să fie examinat *bitmap*-ul asociat valorii v , căutați toți biții cu 1 și să fie accesate înregistrările respective. Un index *bitmap* este foarte eficient, dacă numărul de valori posibile ale atributului indexat este mic.

Unul din avantajele indecșilor *bitmap* este că nu ocupă mult spațiu. Un index *bitmap* este foarte mic în dimensiuni, în raport cu un B-arbore construit pe același atribut. Este foarte util în aplicații de tipul înmagazinărilor de date care conțin volume mari și clasifică informațiile conform unor atribute definite pe domenii mici de valori. Unele interogări pot fi executate foarte eficient, uneori fără a apela la fișierul de date.

O dată ce o înregistrare este inclusă, exclusă sau modificată (în câmpul respectiv), *bitmap*-ul, de asemenea, trebuie modificat. Acest tip de index poate fi folosit eficient și pentru câmpurile atributelor multivaloare.

Arborii binari

Arborii, în general, formează o clasă de indecși destul de eficienți și mult utilizați, astăzi, în bazele de date. Fiecare nod trebuie să posede o cheie și pointeri (vizi sau îndreptați spre nodurile fii).

Un arbore important este arborele binar. În el, fiecare nod poate avea cel mult doi fii. În calitate de index arborele binar poate fi organizat precum urmează. Dat fiind faptul că fiecare nod posedă o cheie, nodul fiu din stânga trebuie să conțină o cheie mai mică decât a nodului părinte, în timp ce

nodul fiu din dreapta trebuie să conțină o cheie mai mare (și așa succesiv până la frunze).

Dacă se caută o înregistrare, localizarea ei se face, analizând nodurile fii. Astfel, pentru găsirea unei chei, se parcurge arborele de la rădăcină, mergând spre stânga, dacă cheia căutată este mai mică decât cheia nodului curent, sau spre dreapta dacă cheia căutată este mai mare (până se întâlnește cheia căutată sau până când nu mai sunt noduri fii).

Inserarea cheii în arbore este făcută după o anumită frunză (în calitate de nod fiu al unui nod frunză), în locul respectiv, aplicând aceleași proceduri ca în procesul de căutare (explicate anterior). Excluderea, însă, a unei chei presupune substituirea ei cu alta (mai mică sau mai mare decât următoarea, care se găsește în nodul frunză respectiv). Modificarea valorii cheii indexului este realizată în arbore de o excludere a cheii vechi și o includere a cheii noi.

Sunt două soluții de stocare a adreselor înregistrărilor:

- adresa se plasează în fiecare nod al arborelui, formând perechea <cheie, adresă>;
- adresa se păstrează numai în nodurile frunze, adică unele chei sunt repetate și în calitate de noduri frunze.

Soluția când fiecărei chei îi este asociată o adresă permite accesarea înregistrării în momentul în care s-a găsit cheia ei în arbore și, prin urmare, nu este necesară parcurgerea arborelui până la nodurile frunze. A doua variantă, când adresele sunt asociate numai nodurilor frunze, facilitează consultarea secvențială a înregistrărilor (după cheia indexului), deoarece nodurile frunze pot fi înlănțuite, formând o listă.

Arborele binar este puțin folosit în calitate de index în SGBD-uri, deoarece după o serie de operații de actualizare poate deveni nebalansat. Chiar dacă cheia unui nod rădăcină are o valoare intermediară, arborele poate avea unele ramuri mai lungi sau mai adânci decât altele, fapt ce poate genera un dezechilibru în executarea operațiilor. Evident că acest lucru poate fi depășit, dacă se va recurge permanent la reorganizarea arborelui. Dar această operație este laborioasă și sunt antrenate în ea toate nodurile arborelui. Cu atât mai mult că baza de date este un sistem care este interogată încontinuu și nu există răgaz pentru a face astfel de restructurări.

B-arbori

B-arborele este un tip de index în formă de arbore, care întotdeauna este balansat. Căutarea înregistrărilor este similară pentru orice cheie, indiferent de localizarea ei în arbore (fie în nodul rădăcină, sau în nodul intermediar sau frunză).

Fiecărui *B-arbore* i se asociază un număr numit ordinul arborelui, care determină numărul maximal și numărul minimal de noduri frunze pe care le poate conține. În afară de aceasta, un *B-arbore* (de ordinul m) trebuie să respecte următoarele reguli:

- fiecare nod posedă m sau mai puțini subarbori;
- fiecare nod, cu excepția rădăcinii, posedă $m/2$ (ia o valoare întreagă) sau mai mulți subarbori;
- nodul rădăcină posedă, cel puțin, doi subarbori nevizi, cu excepția când este un nod frunză;

- orice nod frunză se găsește la aceeași distanță de nodul rădăcină și toți subarborii lor sunt vizibili;
- într-un nod cu k subarbori se stochează $k-1$ înregistrări;
- orice nod de derivare (care nu este frunză) posedă exclusiv subarbori nevizi.

Consultarea după cheie este făcută ca în orice arbore: se parcurge subarborii stânga, dacă cheia căutată este mai mică decât cheia analizată (menționăm că un nod poate avea mai mult de o cheie) sau se parcurge subarborii dreapta, dacă cheia este mai mare. Acest proces continuă până se întâlnește cheia căutată sau se întâlnește subarborii vid).

Pentru a insera o cheie într-un B-arbore, se verifică dacă argumentul de căutare nu coincide cu o cheie deja existentă. Adică, mai întâi, se face consultarea arborelui.

Includerea unei chei este făcută într-un nod frunză, dacă acesta are spațiu suficient (nodul se găsește în limitele date de ordinul arborelui). Dacă nu este spațiu suficient pentru inserarea cheii, trebuie să se producă divizarea nodului (*splitt*) și formarea a două noduri noi. Cheia valorii intermediare trebuie să urce, adică trebuie să fie inclusă în nodul părinte, cu corecțiile corespunzătoare în pointerii precedent și următor, din această cheie, în nodul părinte (acum pointerii devin indicatori spre nodurile create). Nodul nou din stânga va conține cheile mai mici decât cea intermediară și nodul din dreapta, cheile mai mari.

Dacă în nodul părinte nu este spațiu pentru includerea cheii intermediare, care "urcă", nodul părinte trebuie divizat (aceeași procedură). Acest proces poate provoca divizări succesive, până la crearea unei noi rădăcini (din inexistența unui nod părinte, cheia care urcă formează un nou nod rădăcină).

Excluderea unei chei este un proces mai complicat. Mai întâi se localizează cheia respectivă. Acțiunile de mai departe depind de faptul dacă cheia se găsește într-un nod frunză sau nu.

Dacă cheia se găsește într-un nod frunză, ea este exclusă. Rămâne să se verifice dacă nodul frunză se găsește în interiorul limitelor admise. Dacă numărul de chei este mai mic de limita permisă de ordinul arborelui, nodul afectat trebuie să "solicite" o cheie de la nodul frate din stânga sau nodul frate din dreapta (dacă din unul din noduri poate fi extrasă o cheie, iar nodul respectiv să se încadreze în limitele admise). Cheia solicitată de la un nod frate trebuie să fie una de vârf (cea mai mare în nodul frate din stânga sau cea mai mică în nodul frate din dreapta). Atunci, această cheie se ridică la nodul părinte, substituind-o pe cea care se găsea între pointerii celor doi frați afectați. Cheia substituită este coborâtă în nodul de unde a fost exclusă cheia specificată (cheia coborâtă este poziționată în ordinea respectivă).

Dacă, însă, nici un nod frate nu poate "elibera" o cheie pentru nodul din care s-a exclus cheia specificată, se produce concatenarea (*unsplitt*) între nodul afectat și un nod frate. Atunci, cheia nodului părinte care se găsește între identificatorii fraților coboară pentru a fi cheie intermediară între cheile nodurilor frați.

Dacă nodul părinte, după pierderea unei chei, conține mai puține chei decât limita admisă, el trebuie, de asemenea, unit la unul din frații săi, aplicând aceeași procedură de concatenare. Acest proces poate provoca o concatenare în cascadă, până este eliminat nodul rădăcină (rădăcina nouă fiind nodul obținut din concatenarea nodurilor fiu).

În cazul când cheia ce trebuie exclusă se găsește într-un nod de derivare, aceasta va trebui să fie substituită de cheia imediat mai mare sau mai mică. Or, pentru aceasta e destul de parcurs unul din subarbori până la nodul frunză mai apropiat. Dacă nodul frunză, după pierderea cheii (pentru substituirea celei excluse), conține mai puține chei decât limita admisă, procedura de corecție este aceeași ca și în cazul excluderii cheii din nodul frunză.

Trebuie menționat că datele înregistrărilor (câmpurile) sunt păstrate într-un fișier separat. Așadar, fiecare nod are, în afară de pointeri spre nodurile fii, un pointer spre înregistrarea respectivă de date.

O alternativă este păstrarea în noduri și a datelor înregistrărilor. Avantajul acestei structuri este că nu mai e nevoie de încă o accesare pentru găsirea înregistrărilor. Pe de altă parte creșterea dimensiunii nodurilor face ca operațiile de divizare și concatenare să fie mai lente.

B⁺-arbori

Indecșii de tipul *B⁺-arbore* sunt asemănători celor de tipul *B-arbore*, cu deosebirea că toate cheile din nodurile derivate sunt, de asemenea, replicate și în nodurile frunze. Aceasta dă posibilitatea ca nodurile frunze să fie înlănțuite, formând o listă (fiecare nod are câte un pointer pentru nodurile precedent și următor). Această listă facilitează operația de parcurgere secvențială a înregistrărilor (după cheia indexului).

Câștigul în urma aplicării *B⁺-arborelui* este cu atât mai spectaculos, cu cât numărul înregistrărilor crește. Se poate observa o creștere foarte rapidă (exponențială) a numărului de înregistrări indexate în funcție de numărul nivelurilor indexului și invers, o creștere foarte mică (logaritmică) a numărului de niveluri în funcție de numărul înregistrărilor. Întrucât costul unei căutări după cheie este proporțional cu numărul nivelurilor și nu cu numărul înregistrărilor, indexarea permite micșorarea considerabilă a timpului de căutare.

B-arborii oferă o mai mare flexibilitate de găsire a adresei în fișierul de date, unde este stocată înregistrarea și asigură o mai mare eficiență, în principal, în operațiile de inserare și eliminare a înregistrărilor. În plus, *B⁺-arborii* nu exclud posibilitatea de parcurgere secvențială a datelor. Pentru aceasta, este suficient să se parcurgă mulțimea de noduri frunze, folosind lanțul de pointeri ce le leagă. Deci, accesul la fișierele cu un index *B⁺-arbore* poate fi făcut în mod secvențial și aleatoriu.

Totodată, se observă o simplificare în structura generală a fișierului, deoarece fișierele indexate sunt scutite de utilizarea spațiului de tranzacții. Dacă fișierele organizate secvențial și indexat secvențial trebuiau reorganizate periodic pentru a aduce ordinea logică la cea fizică, în fișierele indexate, ordinea fișierului de date este neesențială. Aceasta se datorează faptului că înregistrările sunt accesate prin index, fără a avea de-a face cu configurația fizică a acestora.

Cu toate acestea, o mare problemă întâlnită în fișierele cu organizare de tip indexat este necesitatea de actualizare a tuturor indecșilor când, de exemplu, o nouă inserare este făcută.

28.7 Fișiere cu dispersie

Într-un fișier cu dispersie există o relație predefinită între valoarea cheii utilizate pentru identificarea unei înregistrări și localizarea înregistrării în fișier. Înregistrările nu sunt, neapărat, ordonate fizic în acord cu valoarea cheilor. Înregistrările sunt stocate pe dispozitivele externe de stocare de acces direct și recuperate prin utilizarea acestei legături.

Acesta este tipul de organizare care conține numai un spațiu de date pentru stocarea înregistrărilor. Înregistrările sunt identificate de o cheie principală, a cărei adresă fizică, pentru stocare, este dată de valoarea cheii sau de o valoare calculată, pornind de la valoarea cheii. Într-un fișier cu organizare dispersată, de obicei, cheia este de tip numeric, pentru a calcula adresa și localiza înregistrarea în spațiul de date.

Caracteristicile fișierelor cu dispersie

Fișierele cu dispersie prezintă o organizare a datelor destul de des utilizată în tehnicile de proiectare fizică a bazelor de date, mai ales în modelele orientate pe obiecte. Ele necesită timp, relativ, redus de acces al înregistrărilor individuale, însă lectura secvențială a unui fișier cu dispersie este o sarcină anevoioasă. Există două forme de organizare a fișierelor cu dispersie:

- cu adresare directă;
- cu adresare indirectă.

Caracteristica principală a fișierelor cu dispersie este că adresa fiecărei înregistrări e determinată în funcție de valoarea cheii primare a înregistrării. Astfel, este utilizată o *funcție de dispersie (hash function)*, care, pentru fiecare valoare a cheii C , returnează adresa $h(C)$ într-un spațiu de stocare, unde înregistrarea respectivă trebuie plasată.

Aceasta este cea mai rapidă metodă de acces la date, numai dacă căutarea se petrece cu condiția de egalitate asupra câmpului de dispersie.

Funcția de dispersie trebuie să țină cont de spațiul alocat fișierului, pentru a realiza o distribuire uniformă a înregistrărilor în spațiul rezervat. Există funcții care păstrează ordinea (cheii) și care nu păstrează ordinea (cele mai populare). Între ultimele există funcțiile deterministe (o singură adresă poate fi generată de o cheie) și funcții nedeterministe (când mai mult de o cheie poate genera aceeași adresă, fapt care se numește coliziune).

Pentru includerea în fișier a unei înregistrări, este suficient să se calculeze adresa ei, aplicând valoarea cheii în funcția de dispersie aleasă. Dacă apare o coliziune, adică, dacă există o altă înregistrare pe aceeași poziție, trebuie să fie utilizate unele proceduri de tratare a coliziunilor.

Pentru excluderea unei înregistrări, este de ajuns să se găsească înregistrarea prin aceeași metodă și să se excludă datele ei sau să se marcheze că e exclusă. Procesul de modificare se petrece în mod analog. Pentru modificări în cheia primară se apelează la operația de excludere (a cheii vechi) urmată de o operație de includere (a cheii noi).

Există eventualitatea ca funcția de dispersie să producă, pentru două sau mai multe valori diferite ale cheii, aceeași adresă. În acest caz, prima înregistrare inclusă rămâne în domeniul de date, iar pentru celelalte se

apelează la o procedură de prelucrare a coliziunilor pentru a fi incluse în spațiul de coliziuni. Există mai multe proceduri de tratare a coliziunilor:

- utilizarea unui fișier de tranzații (overflow), unde sunt plasate înregistrările excedentare (care nu pot fi incluse în fișierul principal, deoarece, deja, există altă înregistrare pe aceeași adresă), formând o listă înlănțuită;
- alocarea înregistrării în poziția următoare liberă (fapt ce poate mări probabilitatea apariției coliziunilor);
- înregistrările care se ciocnesc sunt alocate în locurile libere ale aceluiași fișier, formând liste înlănțuite;
- utilizarea unei a doua funcții de dispersie, pentru poziționarea înregistrării pe altă adresă (de asemenea, mărește probabilitatea coliziunilor care, la rândul lor, necesită alt mod de tratament al coliziunilor);
- utilizarea blocurilor – adresa generată de funcția de dispersie aparține unui bloc unde înregistrările pot fi plasate în formă serială.

Consultarea înregistrărilor (după cheia primară) este mai eficientă în acest tip de organizare, dar trebuie să se țină cont de procedura aleasă pentru tratarea coliziunilor.

Dezavantajele acestui tip de organizare sunt: imposibilitatea de consultare după mai multe chei (trebuie să fie utilizate împreună alte mecanisme) și necesitatea de rearanjare (reorganizare), când fișierul crește prea mult (multe includeri). În acest caz, trebuie aleasă o altă funcție de dispersie și toate înregistrările repositionate de funcția nouă.

Pentru a diminua probabilitatea de apariție a coliziunilor, se poate alocă un spațiu mai mare pentru fișier (un spațiu suplimentar de 20%-30%). O bună funcție de dispersie, de asemenea, poate diminua coliziunile (dar nu le poate evita). Acest tip de organizare este optimal pentru aplicațiile care au nevoie de consultări aleatorii rapide (ale unor înregistrări specifice) și când fișierul nu crește mult.

Dacă fișierul crește mult, pot fi utilizate alte tehnici. Tehnicile cunoscute ca *dispersia dinamică* și *dispersia extensibilă* utilizează blocuri pentru alocarea înregistrărilor. Dacă blocul a crescut prea mult (până aproape de limita sa), el trebuie să fie divizat și înregistrările lui redistribuite în blocurile nou-produse.

Cu dispersia dinamică, realocarea înregistrărilor în blocurile noi este dictată de o funcție nouă. Acum, pentru găsirea unei înregistrări, este necesară examinarea istoriei realocării (retrecând toate funcțiile deja utilizate).

În cazul dispersiei extensibile, realocarea este determinată de creșterea cheii. Într-adevăr, cheia produsă de funcția de dispersie este de lungime fixă, dar codificată de o secvență de biți de lungime suficient de mare. Numărul de biți ce este analizat pentru identificarea blocului înregistrării se modifică (se majorează când blocurile sunt divizate și se diminuează în cazul fuzionării blocurilor).

Aceste două tehnici implică reorganizarea unei părți a fișierului, de aceea nu influențează mult funcționarea generală a sistemului.

Dat fiind faptul că înregistrările nu sunt ordonate, ci distribuite în formă aleatorie în fișier, structura fișierului cu dispersie nu funcționează bine pentru consultări secvențiale (parcursarea secvențială a înregistrărilor în ordinea cheii primare).

Dacă există o așa necesitate, poate fi utilizată tehnica *dispersia indexată*. Aceasta este o combinaire între organizările de tip indexat și dispersat. În dispersia indexată, valoarea produsă de funcția de dispersie nu este adresa directă a înregistrării de date, ci o adresă în fișierul index (cu perechea care reprezintă cheia și adresa înregistrării).

Fișierul index funcționează ca cel mai intern nivel al organizărilor de tip indexat, unde cheile apar ordonate. Consultarea secvențială poate fi făcută asupra acestui fișier index, în timp ce, pentru cele de consultare aleatorie apare numai o accesare în plus. O altă alternativă poate fi înlănțuirea înregistrărilor cu utilizarea pointerilor.

Astfel, pentru fișierele cu dispersie există două moduri de acces:

- secvențial, care reprezintă o lectură secvențială a datelor, conform ordinii dictate de cheie.
- aleatoriu, care localizează o înregistrare după adresa calculată cu ajutorul unei funcții de dispersie, pornind de la valoarea cheii de acces.

Adresarea directă

Adresa directă, în forma sa mai simplă, este o adresă hardware, fie relativă, fie absolută, care permite sistemului să acceseze înregistrarea corespunzătoare. Adresa hardware a unei înregistrări este formată din numărul cilindrului, numărul pistei și din poziția înregistrării pe pistă.

O abordare mai simplă a acestui tip de organizare constă în tratarea valorii cheii ca un număr tradus de sistem într-o adresă hardware relativă.

De exemplu, fie un fișier cu organizare directă păstrează înregistrări cu valorile cheii primare între 0 și 9999. Dacă valoarea cheii este tratată ca număr relativ al înregistrării, atunci trebuie rezervat un spațiu pentru 1000 înregistrări, care corespunde unei valori posibile a cheii între 0 și 9999 inclusiv.

Adresarea directă semnifică faptul că sistemului îi este furnizată adresa fizică a înregistrării căutate, pentru a-i permite recuperarea directă a înregistrării. Adresa furnizată poate fi atât fizică a hardware-ului (numărul cilindrului, numărul pistei, numărul înregistrării), cât și un număr de adresă relativă (poziția într-un fișier). A doua opțiune este mai utilizată, și din faptul că adresa relativă poate fi automat tradusă într-o adresă fizică.

Tratarea unei valori a cheii ca număr de adresă relativă are un dezavantaj. De exemplu, dacă un fișier cu organizare directă are înregistrări cu valorile cheii între 0 și 9999, trebuie rezervate 10000 înregistrări, fiecare corespunzând unei valori posibile a cheii. Această rezervă trebuie făcută, chiar dacă în fișier există numai o înregistrare (cheia 9999, de exemplu).

Spațiul neutilizat, în această organizare, poate fi foarte mare. Pentru aceasta administratorul bazei de date trebuie să analizeze necesitatea reală de adrese directe în ce privește spațiul real ocupat de fișier și spațiul irosit. Dacă spațiul liber, procentual, nu este mare, poate fi compensat de avantajele utilizării adresării directe.

Când numărul de chei posibile este, relativ, aproape de numărul real de chei, această formă de adresare este foarte eficientă. Dar, aproape întotdeauna, numărul real de chei este mult mai mic decât numărul de spații de înregistrări care trebuie păstrate. Când aceasta se întâmplă, pentru accesarea unui fișier cu dispersie, trebuie să fie utilizate consultarea dicționarului sau adresarea indirectă. Pentru ca sistemul să poată converti

această valoare într-o adresă hardware, trebuie cunoscute următoarele mărimi: valoarea cheii, numărul de cilindri pe dispozitivul extern de stocare, numărul de înregistrări pe pistă, numărul de piste pe cilindru, adresa hardware absolută a fișierului.

Pentru a găsi adresa relativă și absolută a unei înregistrări sunt trei etape trecute de sistem:

- valoarea cheii se împarte la numărul de înregistrări pe pistă, unde câtul este numărul pistei relative a înregistrării, iar restul este poziția relativă a înregistrării pe această pistă;
- numărul pistei relative a înregistrării este împărțit la numărul de piste pe cilindru și se obține numărul relativ al cilindrului și pista relativă a cilindrului (restul);
- adresa hardware absolută a înregistrării se găsește, dacă se adaugă adresa hardware relativă a înregistrării la adresa hardware absolută a fișierului.

Adresarea indirectă

Adresarea indirectă presupune că se aplică o funcție de transformare a valorii cheii într-o adresă directă, care poate fi utilizată de sistemul de gestiune pentru localizarea înregistrării dorite. Aceasta, spre deosebire de adresarea directă, unde nu se face nici o transformare a valorii cheii înainte de utilizarea ei de sistem.

În fișierele cu dispersie, cu adresare indirectă, înregistrările se grupează în clase. Apartenența unei înregistrări la una din clase este determinată, în funcție de valoarea pe care o are cheia înregistrării.

Spațiul necesar pentru stocarea claselor de înregistrări este divizat în blocuri de stocare. Un bloc de stocare poate conține una sau mai multe înregistrări logice, capacitatea lor fiind determinată de proiectant. Funcția de dispersie calculează, de fapt, adresa primului bloc al clasei, unde este stocată înregistrarea. În implementarea fișierelor cu dispersie se construiește o listă numită *directoriu* cu pointeri la blocul de începere corespunzător fiecărei clase, figura 28.12.

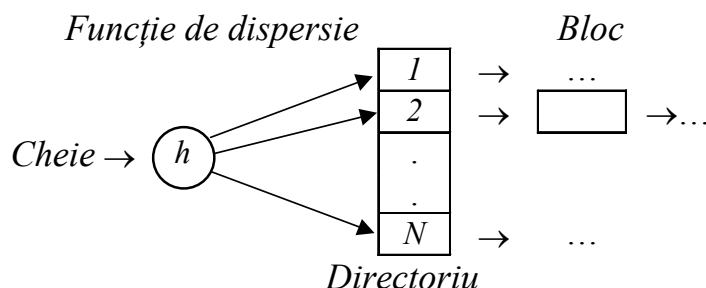


Figura 28.12 Un fișier cu dispersie

Funcția de dispersie, de obicei, poate mapa diferite valori ale cheii pe aceeași adresă a spațiului de stocare, provocând fenomenul numit *coliziune*. Evident că înregistrările care au aceeași adresă produsă de funcția de dispersie aparțin aceleiași clase. Stocarea înregistrărilor ce aparțin unei clase în blocuri depinde de metodele de soluționare a coliziunilor.

Principala problemă ce ține de adresarea indirectă este alegerea funcției h , care transformă valoarea C a cheii unei înregistrări în adresa A , care îi corespunde în fișier. O altă problemă este găsirea unei metode adecvate de soluționare a coliziunilor și de plasare a înregistrărilor excedentare în blocuri. Astfel, algoritmul de lucru cu fișierele cu dispersie presupune, printre altele, două activități: maparea (utilizând funcția de dispersie) valorii cheii într-o adresă absolută sau relativă și citirea înregistrărilor excedentare.

Dacă proiectantul alege un spațiu mic de stocare a primului bloc din fiecare clasă de înregistrări, apare un număr relativ înalt de coliziuni, fapt ce produce lecturi suplimentare pentru localizarea înregistrării dezirabile în spațiul de înregistrări excedentare. O capacitate de stocare mai mare provoacă mai puține coliziuni, dar mult spațiu rămâne nevalorificat. În practică, capacitatea spațiului de stocare depinde de caracteristicile dispozitivului secundar de stocare a fișierului cu date. Raportul dintre numărul de înregistrări existente în fișier și numărul total de locuri disponibile pentru înregistrări în fișier se numește *grad de încărcare* a fișierului. Pe măsură ce gradul de încărcare crește, probabilitatea unei înregistrări, ce trebuie adăugată, de a deveni una excedentară (care provoacă o coliziune), de asemenea, crește.

Eficiența unui fișier cu dispersie, cu utilizarea adresării indirecte, este influențată de mai mulți factori:

- dimensiunea spațiului de stocare;
- gradul de încărcare;
- valorile cheilor;
- funcția de dispersie;
- metoda de soluționare a coliziunilor;
- metoda de citire a înregistrărilor excedentare.

Calcularea adresei

Există diverse tehnici pentru transformarea cheii într-o adresă.

Pot fi considerate două tipuri de funcții de dispersie, primul fiind constituit din *funcții deterministe* care pentru două valori diferite ale cheii produc adrese diferite. Acest tip de funcții prezintă avantaje evidente. Dar, în practică, pentru un număr mare de înregistrări, este imposibil de găsit o funcție deterministă simplă. Acele funcții care pot fi folosite sunt așa de complexe, că elimină toate avantajele accesului direct. În afară de necesitatea de adaptare la fiecare inserare suferită de fișier, de obicei, produc adrese necontinue, ceea ce duce la utilizarea unui spațiu larg, similar adresării directe. Din acest motiv, funcțiile deterministe nu prezintă un mare interes practic.

Al doilea tip formează funcțiile nedeterministe, care produc pentru fiecare valoare a cheii o singură valoare a adresei, iar, uneori, pentru valori distincte ale cheii, aceeași adresă, ceea ce constituie o coliziune. Funcțiile nedeterministe pot fi proiectate și cu scopul de a atinge unele obiective secundare, obiective care nu sunt incompatibile:

- păstrarea ordinii înregistrărilor după valoarea cheii de acces;
- mărirea gradului de unicitate a adreselor produse, adică, distribuirea înregistrărilor în fișier, cu o mare uniformitate posibilă.

O funcție de dispersie nedeterministă care nu ține cont de ordinea cheii se mai numește *funcție de randomizare* și are ca obiectiv distribuirea uniformă a înregistrărilor în fișier.

Funcția de dispersie "restul de la împărțire". Această funcție de dispersie este utilizată frecvent. În calitate de valoare a funcției de dispersie se folosește restul de la împărțirea unei chei la un număr întreg N : $h(C) = C \bmod N + 1$. Aici la rezultatul împărțirii se adaugă 1 pentru a obține o valoare a funcției între 1 și N .

Evident că eficiența distribuirii cheilor depinde mult de valoarea lui N . Astfel, contraindicat ca N să fie puterea bazei unui sistem de numerație, deoarece, în acest caz, valorile funcției de dispersie vor fi reprezentate de ultima cifră a cheii. Cele spuse sunt adevărate și pentru cheile alfanumerice. Dacă $N=2^8$, de exemplu, valoarea funcției de dispersie va fi ultimul caracter al cheii. Pentru a depăși concentrarea cheilor în jurul unor adrese, se recomandă ca N să fie un număr prim.

Deseori, mulțimea de chei reprezintă niște secvențe de progresii de tipul $XXA, XXB, \dots$ sau $FT01, FT02, \dots$. Funcția de dispersie *restul de la împărțire* transformă astfel de chei în adrese succesive. O asemenea situație poate fi considerată destul de satisfăcătoare.

Metoda înmulțirii. Fie cheia C este reprezentată în forma unui număr binar și fie $N=2^n$. Se înmulțește o fracție α cu C și se ia partea fracționară care este notată prin $\{C\alpha\}$, iar în calitate de valoare a funcției de dispersie se folosesc numai primii n octeți. Cu alte cuvinte, $h(C) = \lfloor N \times \{C\alpha\} \rfloor$, unde $\lfloor x \rfloor$ este cel mai mare număr întreg mai mic decât x .

Se recomandă a se folosi în calitate de valoare α un număr irațional, aproape de lungimea unui cuvânt. De exemplu, dacă α este reprezentată de $\alpha = (\sqrt{5} - 1)/2$, atunci segmentul $[0,1]$ se împarte foarte bine la $\{\alpha\}, \{2\alpha\}, \dots$. Cu alte cuvinte, oricât de mărunț s-ar diviza segmentul, lungimea segmentelor obținute după divizare nu se vor exprima prin mai mult de trei valori $\alpha^C, \alpha^{C+1}, \alpha^{C+2}$. În cazul divizării următoare segmentul de lungime α^C se împarte în segmente cu lungimile $\alpha^{C+1}, \alpha^{C+2}$. Această proprietate, în cazul aplicării metodei înmulțirii, permite a obține rezultate bune.

Cerința $N=2^n$ pentru calcularea funcției de dispersie $h(C)$ nu este obligatorie. Trebuie menționat că, dacă $\alpha=1/N$, atunci această metodă este echivalentă cu metoda *restul de la împărțire*.

Foarte aproape de această metodă de dispersie este metoda "*mijlocul pătratului*", în care, în calitate de valoare a funcției de dispersie $h(C)$, se folosesc n octeți din mijlocul numărului C^2 . Și aceasta este o metodă de dispersare, dar, după mulți parametri, cedează metodei înmulțirii.

Metoda transformarea sistemelor de numerație. Metoda de calculare a valorii funcției de dispersie presupune transformarea valorii cheii C , reprezentată în sistemul de numerație în baza p , $C = a_0 + a_1p + a_2p^2 + \dots$, într-o valoare reprezentată în sistemul de numerație în baza q cu constrângerea s asupra ordinii rezultatului: $h(C) = a_0 + a_1q + \dots + a_{s-1}q^{s-1}$. Aici p și q sunt numere pozitive, astfel, că $p < q$.

Pentru calcularea valorii funcției de dispersie $h(C)$, sunt necesare s operații de înmulțire sau împărțire. Prin urmare, complexitatea (numărul de operații) acestei metode este mai înaltă decât a metodei *înmulțirea*.

Metoda împărțirea polinoamelor. Fie cheia C , exprimată în sistemul binar de numerație, se scrie $C = 2^{m-1}b_{m-1} + \dots + 2b_1 + b_0$ și fie $N=2^n$. Cheia binară C se reprezintă în forma unui polinom $C(t) = b_mt^m + \dots + b_1t + b_0$, păstrând aceiași coeficienți. Acum se determină restul de la împărțirea acestui polinom la polinomul constant de forma $P(t) = t^n + p_{n-1}t^{n-1} + \dots + p_1t + p_0$. Restul obținut se consideră din nou în sistemul binar de numerație și se folosește în calitate de valoare a funcției de dispersie $h(C)$. Pentru calcularea restului de la împărțirea polinoamelor, se folosește aritmetica polinomială după modulul 2. Dacă în calitate de $P(t)$ este ales un polinom ireductibil simplu, atunci, pentru două chei foarte aproape, dar nu egale, $C_1 \neq C_2$, întotdeauna se va realiza condiția $h(C_1) \neq h(C_2)$. Această funcție de dispersie posedă proprietăți mai puternice de dispersare a mulțimii de chei decât precedentele.

Trebuie menționat că, în fișierele cu dispersie, transformările cheii în adresă sunt similare funcțiilor de dispersie utilizate pentru căutarea în tabelele stocate în memoria principală. Deosebirea constă în caracteristicile accesului la memoria secundară, care diferă de cele ale accesului în memoria internă accesibilă direct. Pentru funcțiile de dispersie în memoria principală timpul de calculare a adresei poate fi critic, însă, pentru memoria secundară, acest timp poate fi neglijat.

Tratarea coliziunilor

Tratarea coliziunilor este unul din aspectele mai importante în adresarea indirectă și este consecința utilizării funcțiilor nedeterminate pentru transformarea valorilor cheii în adrese ale fișierului. Se spune că există o coliziune când două valori diferite ale cheii de acces sunt transformate în aceeași adresă.

Deoarece pe o adresă se poate stoca o singură înregistrare, este necesară stabilirea unei proceduri pentru a alege o altă adresă, unde va fi plasată înregistrarea care a provocat coliziunea. Această procedură se aplică și în operația de acces, dacă, pe adresa produsă în baza argumentului de căutare, se găsește o înregistrare cu o cheie diferită de cea a argumentului.

Soluțiile mai frecvent utilizate sunt *adresarea deschisă* și *înlănțuirea*.

Tratarea cu adresare deschisă. Această metodă de soluționare a coliziunilor presupune că se folosește o anumită regulă de parcurgere a înregistrărilor. O dată cu apariția unei coliziuni la operația de inserare, este făcută o căutare în fișier pentru localizarea unei adrese libere, unde va fi stocată înregistrarea. Procesul de inserare se supune unei secvențe predefinite, care este urmată și în timpul accesului la înregistrare. Regula de căutare trebuie să fie astfel, încât, chiar în lipsa pointerilor, aceeași secvență să poată fi parcursă. Procedura poate arăta astfel:

1. $i=0$
2. $A=h_i(C)$:
 - 3.1. Dacă adresa A este liberă, stop (dacă aceasta e procedura de inserare, algoritmul se termină cu succes, însă, dacă aceasta e procedura de căutare, înseamnă că așa cheie nu există);
 - 3.2. Dacă cheia de pe adresa A este egală cu C , stop (dacă este procedura de inserare, inserarea a doua nu se mai face, iar dacă e procedura de căutare, algoritmul se termină cu succes);

4. În caz contrar – coliziune; $i:=i+1$ și se trece la pasul 2.

În sensul larg, $\{h_i\}_{i=0}^{\infty}$ reprezintă o secvență de funcții de dispersie. Alegerea lor este o problemă destul de complicată. Considerând că h_0 este una din funcțiile examinate anterior, mai departe se examinează celelalte funcții din secvență.

Căutarea poziției libere poate fi realizată cu una din următoarele metode de adresare deschisă:

- cercetarea liniară;
- cercetarea pătratică;
- rerandomizarea.

Cercetarea liniară. Aceasta este cea mai simplă schemă de adresare deschisă, în care $h_i(C) = h_0(C) + ki$, unde k este o constantă. Regula este simplă, însă există pericolul unei concentrări a înregistrărilor la începutul procesului, unde foarte mult sunt mixate adresele. Pentru a evita aceasta, k și N trebuie să fie reciproc numere prime, iar k un număr nu prea mic. Dar în cazul când numărul cheilor este mic în raport cu numărul N , pericolul creșterii căii de căutare nu este mare, iar faptul că sinonimele se vor stoca pe adrese adiacente este un avantaj. În figura 28.13 numărul de adrese $N=10$, funcția de dispersie inițială este $h_0(C) = C \bmod 10 + 1$, iar $h_i(C) = h_0(C) + i$.

C	$h_0(C)$	$h_i(C)$	Adresă	Înregistr.
25	6		1	80
43	4		2	
56	7		3	
35	6	8	4	43
54	5		5	54
13	4	9	6	25
80	1		7	56
104	5	10	8	35
			9	13
			10	104

Figura 28.13 Cercetarea liniară

Cercetarea pătratică. Aceasta este schema de adresare deschisă, în care $h_i(C) = h_0(C) + ki + di^2$, unde k și d sunt constante. Datorită neliniarității acestei adresări se poate evita creșterea numărului de probe pentru mai mult de două sinonime. Însă, în cazul unui fișier aproape plin, nu se poate evita a doua concentrare de chei. Aceasta e condiționată de faptul că, în cazul unui număr mare de chei, sinonimele unui grup intră în coliziune cu alte chei.

De acest dezavantaj practic este absolută metoda următoare, în care secvențele de probe pentru astfel de chei se dovedesc a fi diferite.

Rerandomizarea. Rerandomizarea este schema de adresare deschisă în care, $h_i(C) = g(h_{i-1}(C))$, unde $g(C)$ este o funcție asemănătoare cu h_0 , dar nu echivalentă ei. Procesul este repetat până se găsește o adresă liberă. Avantajul acestei metode, în comparație cu cele anterioare, este de a

evita concentrarea excesivă a înregistrărilor în unele regiuni ale fișierului, mărind, deci, eficiența operațiilor și accesului. În figura 28.14, este reprezentat rezultatul aplicării funcțiilor de dispersie $h_0(C)=C \bmod 10+1$ și $h_i(C)=(h_{i-1}(C)+1) \bmod 10+1$.

C	$h_0(C)$	$h_1(C)$	$h_2(C)$	$h_3(C)$	Adresă	Înregistr.
25	6				1	80
43	4				2	
56	7				3	
35	6	8			4	43
54	5				5	54
13	4	6	8	10	6	25
80	1				7	56
104	5	7	9		8	35
					9	104
					10	13

Figura 28.14 Rerandomizarea

Tratarea prin înlănțuire. Tratarea prin înlănțuire a coliziunilor presupune că toate sau o parte de chei sinonime sunt adunate și înlănțuite cu ajutorul pointerilor, formând o listă. Fiecare listă se accesează prin adresa produsă de funcția de dispersie, figura 28.15. Listele pot fi păstrate atât împreună cu directoriul, cât și într-un spațiu de memorie separat.

C	$h(C)$	Adresă	Înregistr.
25	6	1	80
43	4	2	
56	7	3	
35	6	4	43
54	5	5	54
13	4	6	25
80	1	7	56
104	5	8	
		9	
		10	

13	
104	
35	

Figura 28.15 Tratarea prin înlănțuire

În cazul când listele de adrese sunt stocate aparte, consultarea lor este simplă. După calcularea valorii funcției de dispersie $A=h(C)$, problema se reduce la căutarea secvențială a listei cu adresa A (eventual stocată în memoria secundară). Memoria secundară, în acest caz, poate fi divizată în blocuri, astfel încât cheile sinonime să se păstreze într-un bloc. Deoarece directoriul va conține numai pointeri, numărul lor poate fi destul de mare pentru a micșora probabilitatea apariției coliziunilor. Această metodă este adecvată unui număr mare de înregistrări de lungime variabilă.

Nu apar probleme la includerea unor chei noi. Dacă, însă, numărul de sinonime devine prea mare, în loc de liste liniare, pot fi utilizate structuri arborescente.