

29. STRUCTURI FOLOSITE ÎN PROCESE DE CĂUTARE

29.1 Tipologii de aplicații

Aplicațiile informatice diferă funcție de obiectivele pentru care au fost elaborate.

Tipologiile aplicațiilor sunt considerate în raport cu criterii de clasificare elaborate atât de elaboratori, cât și de dealeri, dar mai ales de către clienți.

După criteriul structurii, aplicațiile informatice sunt:

- *seriale*, în care componentele sunt referite una câte una; calitatea execuției componentei precedente determină referirea componentei următoare în secvență și deci, execuția întregului lanț;
- *arborescentă*, caz în care componentele se dispun pe niveluri și referirea depinde de evaluarea unei expresii relaționale; sunt situații în care pe o ramură a arborescenței prelucrările conduc la rezultate corecte, în schimb pe alte ramuri rezultatele sunt incerte;
- *rețea*, care presupune atât execuții în secvență cât și evaluări de expresii relaționale; referirile sunt uneori comune pentru mai multe componente, atât pentru componente următoare, cât și pentru componente precedente ca poziție, care devin următoare în execuție.

Abordările moderne impun clasificări spre constituirea de clase, de module, de interacțiuni, ceea ce conduce la construcții simple, medii sau complexe.

De asemenea, în raport cu evoluțiile ulterioare se caută accentuarea laturilor de mentenanță sau portabilitate. Aplicațiile informatice sunt sau nu sunt înzestrate cu anumite proprietăți, caz în care sunt incluse într-una din categoriile aparținând criteriilor care se definesc ad-hoc. Din punct de vedere al utilizatorilor, aplicațiile informatice sunt:

- pentru un singur client, de regulă specializat în realizarea de prelucrări complexe și este vorba de aplicații particulare, cu volum mare de decizii care presupun nivele de agregare ridicate; aplicația presupune numeroase convenții referitoare la tipurile de operații, procedurile care se activează și ramurile care se traversează în rețeaua asociată; există numeroase puncte de intervenție pentru a actualiza parametrii și chiar proceduri de calcul și de selecție a datelor din bazele de date pentru a urmări dinamica sistemului economic din care aplicația este parte;
- cu număr restrâns de utilizatori, destinate lucrului cu baze de date al căror volum de actualizare este ridicat; aplicațiile sunt specializate iar operatorii sunt profesioniști în domeniu; se definesc tipologii de operații cărora le corespund coduri de selecție și date de intrare specifice; se urmărește rezolvarea de probleme într-o interval dat de timp, cu controlul riguros asupra volumului de date; se elaborează chei de control pentru garantarea calității datelor de intrare care garantează la rândul lor calitatea rezultatelor; aceste aplicații determină plăți, acceptarea unor stări de fapt, selecția de candidați și alte forme de tranzacții

sau de construire a argumentațiilor; operatorii sunt instruiți atât în utilizarea produselor software, cât mai ales în semnalarea de erori în efectuarea de corecții acolo unde este posibil;

- număr de utilizatori foarte mare, de masă, cu grad de neomogenitate ridicat; interfețele sunt astfel construite încât toți utilizatorii să obțină maximum de satisfacție; se impune ca aceste aplicații să fie atent elaborate, să țină seama de cerințele reale ale utilizatorilor și să nu necesite procese de instruire speciale; în cazul în care numărul de pași care se parcurg este indicat, trecerea de la un pas la altul este marcată prin comenzi de același tip; problema rescrierii la pasul inițial (de start) se va realiza de la oricare din etapele interacțiunii; neomogenitatea clienților impune utilizarea de interfețe grafice; clientul trebuie să obțină maximum de beneficii cu minimum de date introduse de la tastatură; se efectuează o cercetare sistematică pentru a evidenția care sunt situațiile destinate care trebuie tratați pentru cazurile cu frecvența cea mai mare se definesc fluxuri de opțiuni care prin tastări simple să conducă la traversarea de pași și obținerea de efecte în concordanță cu așteptările utilizatorilor; clienții efectuează un număr limitat de selecții, iar modul de reprezentare a datelor reduce repetitivitatea de operații generată de erori de interpretare; simplitatea dialogului om-mașină, lipsa de rigiditate și luarea în considerare a tipurilor de clienți cu comportamentul specific are rolul de a genera aplicații informatice destinate de succes, viabile.

29.2 Chei de regăsire

Aplicațiile informatice cele mai frecvente sunt caracterizate prin utilizarea în procesul de regăsire a unei singure chei.

De regulă, în procesul de analiză a problemei de rezolvat se parcurg următorii pași:

- se identifică mulțimile cu care se operează precum: mulțimea persoanelor, mulțimea materialelor, mulțimea documentelor, mulțimea operațiilor etc;
- se stabilesc numărul maxim de componente care alcătuiesc mulțimile;
- se alege algoritmul de căutare a cheilor așa fel încât să se asigure unicitatea în concordanță cu apartenența elementelor colectivității la eventuale submulțimi;
- se generează mecanisme de stare și de căutare care să reducă duratele de prelucrare.

În cazul în care dinamica înregistrată de colectivitate permite să construim chei numerice.

Pentru salariații unei întreprinderi, marca salariatului – unică – este un număr. Salariatul X are marca 7250, ceea ce înseamnă că a fost a 7250 a persoană care s-a angajat. Cei dinaintea sa mai lucrează, s-au transferat sau au plecat la pensie.

Aplicațiile bazate pe marca salariatului impun manipularea unei legitimații în mod obligatoriu.

În cazul în care se dorește diferențierea muncitorilor pe secții se face o codificare a secțiilor și o codificare a muncitorilor și se concatenează cele două coduri, rezultând marca muncitorului. De exemplu, pentru o întreprindere cu 25 de secții numerotarea secțiilor se face cu 01, 02, ..., 25, iar pentru codul muncitorului se constituie secvențe formate din patru cifre.

Marca salariatului 140014 arată că este vorba de salariatul de la secția 14, având codul 0014 în codul secției respective. Aplicațiile bazate pe o singură cheie sunt proiectate astfel încât să se poată gestiona cu ușurință toate datele. Simplitatea acestor aplicații impune un nivel ridicat de rigiditate. Generarea codurilor în mod serial corespunzătoare mulțimilor omogene face dificil procesul de regăsire atunci când lipsește suportul pe care este indicat codul elementului ce trebuie căutat. Elaborarea de structuri care se asociază cheilor și a mecanismelor de continuare a elementelor elimină parțial acest impediment.

Este important ca părțile ce formează structura să fie alcătuit din mulțimi deja existente, cunoscute, iar numărul de elemente generate să fie cât mai restrâns.

Dacă o firmă comercializează produse electrocasnice, autoturisme și materiale de construcții către persoane fizice codul contractului este o construcție rezultată din concatenarea secvențelor de numere următoare:

- secvența corespunzătoare tipului de produs achiziționat;
- secvența de identificare a cumpărătorului.

Se stabilește un algoritm de construcție a codului contractului. Fiecare literă din alfabet are o poziție: litera a are poziția 01, litera b are poziția 02, litera i are poziția 09, litera r are poziția 17 etc.

Lunile anului se numerează 01 (ianuarie) sau 07 (iulie) și 12 (decembrie). Zilele lunii se numerează de la 01 la 31. Anul de naștere se numerează cu patru cifre pentru a elimina ambiguitățile de tipul celei generate de trecerea la anul 2000.

Dacă sunt luate în considerare și părți din numele persoanei fizice care achiziționează produse, cu certitudine este asigurată unicitatea codului contractului.

Dacă structura codului de control este descrisă prin:

- poziția 1: tipul produsului E – electrocasnice, A – autoturism, C – material de construcții;
- pozițiile 2, 3, 4, 5: anul nașterii cumpărătorului;
- poziția 6,7: luna nașterii cumpărătorului;
- poziția 8,9: ziua de naștere a cumpărătorului;
- pozițiile 10, 11, 12: primele 3 litere din numele cumpărătorului la completare se asigură unicitatea codului contractului ca cheie de regăsire.

Întrucât un rol esențial în evidențe îl are codul numeric personal, în dezvoltarea aplicațiilor moderne va fi inclus în structura cheilor în vederea regăsirii.

În operațiile cu clienții posesivi de posturi telefonice numărul postului și codul personal deja formează un cod redundant.

La proiectarea cheilor de regăsire este preferabil să fie utilizate secvențe care descriu elemente deja existente precum:

- două litere ce reprezintă prescurtările numelui de județ;
- literă M, F prin care se indică sexul persoanei;
- data nașterii prezentată prin an, lună, zi;
- alte componente din structura codului personal;

- cuvinte prin care se desemnează în clar produse (telefon, frigider, cuier, tablă, cablu, geam);
- denumiri de localități, străzi, instituții.

Schimbarea concepției privind definirea cheilor rezidă din trecerea de la aplicațiile bazate pe informații simple utilizate în programul de regăsire, oferite complet, după reguli specificate, la aplicații cu clienți diferiți, neomogeni ca pregătire în raport cu o aplicație informatică, unde regăsirea trebuie să se facă dacă informațiile pe care clientul le are.

Cheile de regăsire sunt numeroase și la proiectarea aplicațiilor se impune comutarea de arborescență ale cor frunze sunt de fapt codurile unice.

În cazul informațiilor incomplete sunt extrase submulțimi dintre care este identificat elementul căutat.

De exemplu, se ia în considerare codul numeric personal, căruia i se asociază arborescența dată în figura 29.1.

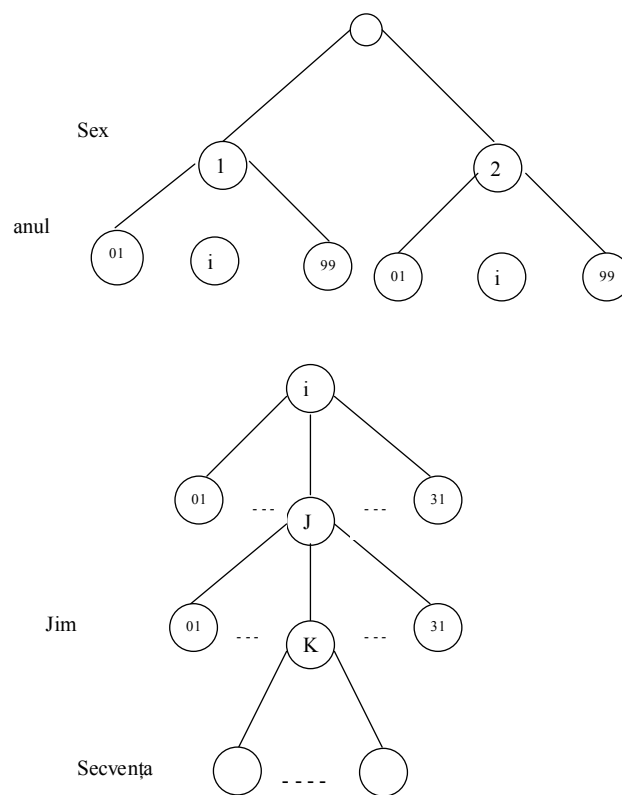


Figura 29.1 Structura arborescentă asociată CNP

Este important ca la un volum cât mai restrâns de date de intrare să se obțină o finețe de selecție cât mai bună

Pentru efectuarea de studii statistice se ajunge la forma de structură arborescentă care să genereze un număr cât mai restrâns de elemente.

În cazul persoanelor încorporabile anul nașterii conduce la submulțimi formate din multe elemente.

În schimb ziua de naștere urmată de luna de naștere conduce la o mulțime de elemente din care se extrage cu ușurință un element pentru a efectua prelucrări solicitate.

Problema cheilor de regăsire trebuie rezolvată adecvat pentru a genera echilibrul aplicațiilor informatice și pentru a reflecta particularitățile

evoluției fiecărei aplicații în parte, în raport cu clienții cărora este destinată. Se au în vedere aplicațiile distribuite cu accesul direct al clienților.

Este important să se utilizeze pentru regăsire *chei grafice* date sub forma unor simboluri, liste de cuvinte care corespund numelor de localități, denumiri de produse, mărci de produse, categorii de produse, valori numerice, denumiri de evenimente etc.

Este important ca fiecare client să tasteze cât mai puține informații de intrare, iar produsul software să aibă o serie de algoritmi de normalizare care să conducă la realizarea acelui nivel de flexibilitate cerut aplicațiilor distribuite care se adresează publicului larg.

Astfel, dacă se dorește un număr de telefon al unei persoane vor fi solicitate următoarele date:

Numele: x x x...x

Prenumele: x x x ...x

Orașul: lista de orașe

Localitatea în ordine alfabetică sau posibilitatea de a tasta altă localitate

Strada: x.....x

Bloc: xxx

Apartament: xxx

Număr: xxxx

Atunci când aplicația este destinată utilizatorilor specializați sau care au efectuat un stagiul de instruire, introducerea datelor este clară, întrucât meniul este suficient de riguros.

În cazul utilizatorilor cu grad de instruire, cu cunoștințe incomplete asupra modului de formulare a cererii apar următoarele situații:

- numele este ortografiat astfel decât forma stocată în baze de date; de exemplu Brâncuși din baza de date poate fi tastat BRINCUS sau Brâncuș sau Brancusi sau Brancus;
- numelui îi lipsește o vocală sau o consoană dublă; de exemplu, cuvântul BELLER este ortografiat Beler sau Belăr, Johnson este ortografiat Jonson
- numele este interschimbat cu prenumele; Bălcescu Nicolae este ortografiat Nicolae Bălcescu
- prenumele este înlocuit cu o literă sau grup de litere; Nicolae e scris N Gheorghe e scris Gh. sau este înlocuit cu un formular sau diminutiv. Ion este înlocuit cu Nelu, Dan este înlocuit cu Daniel, Gheorghe este înlocuit cu Gică sau Gicu sau Gigel. Aceleași probleme sunt puse și în cazul ortografiei cu litere mari. În cazul în care numele și prenumele este introdus într-o singură rubrică cea a numelui se impune o operație opusă concatenării.

Se consideră colectivitatea $A = \{a_1, a_2, \dots, a_n\}$ pentru care un element a_i este descris prin șablonul Sc definit de caracteristicile $sc_1, sc_2, \dots, sc_m$ în mulțimea $Sc = \{sc_1, sc_2, \dots, sc_m\}$. Pentru elementul a_i se generează șirul de caractere c_{ij} pentru a specifica șirul caracteristicii Sc_j din șablon.

Descrierea este completă atunci când pentru toate elementele colectivității A au fost măsurate sau generate șiruri de caractere pentru toate caracteristicile ce definesc șablonul. O astfel de definire este dată în tabelul 29.1.

Tabelul nr. 29.1 Măsurarea elementelor colectivității materialelor

Denumire	Cod	U.M.	Stoc inițial	Intrări	Ieșiri	Preț unitar
Var	100	kg	100	10	50	5
Ciment	122	kg	500	30	200	100
Căramidă	400	buc.	10000	1000	5000	1
Cuie	310	kg.	100	50	25	25
Rigips	133	buc.	200	20	50	60

Sunt situații în care nu sunt efectuate definiri complete sau măsurători datorită neluării în considerare a caracteristicilor sau datorită costurilor foarte ridicate de a măsura, așa cum se prezintă în tabelul 12.2.

Tabelul nr. 29.2 Înregistrări incomplete pentru colectivitatea elevilor

Nume	Sex	Data nașterii	Înălțime (cm)	Școală
Popescu Ion	M	10.08.1985	173	174
Alexandrescu	M		180	174
Gheorghe Ion	M	11.05.1984		3
Pana Marius		14.03.1984	165	
Popescu Alina	F	10.07.1985	167	3
Jderu Ioana	F	22.05.1984	165	

Pentru fiecare dintre elementele mulțimii A se stabilesc șiruri de caractere care se constituie în vocabularul asociat colectivității A .
Se iau în considerare frecvențele de apariție ale cuvintelor.
Pentru o colectivitate B definită în tabelul se măsoară frecvențele de apariție a cuvintelor în vocabular.

Tabelul nr. 29.3 Nivelurile generate pentru caracteristicile C_1 , C_2 , C_3 și C_4 care definesc colectivitatea B

C_1	C_2	C_3	C_4
aa	x	u	www
bbb	yy	u	zzz
cc	x	v	zzz
ddd	x	v	zzz
e	yy	u	zzz
fff	x	u	zzz
g	yy	u	www
h	yy	v	zzz
iii	yy	v	zzz
jj	yy	v	zzz

Vocabularul asociat colectivității A , V_A este definit de șirurile de caractere $V_A = \{aa\ bbb\ cc\ ddd\ e\ fff\ g\ h\ iii\ jj\ x\ yy\ u\ vv\ www\ zzz\}$.

Frecvențele de apariție a cuvintelor sunt date pentru caracteristici în tabelul 29.4

Tabelul nr. 29.4 Frecvențele pentru caracteristicile C_1 , C_2 , C_3 și C_4 care definesc colectivitatea B

Cuvânt	fC_1	fC_2	fC_3	fC_4
aa	0	0	0	0
bbb	0	0	0	0
cc	0	0	0	0
dddd	0	0	0	0
e	0	0	0	0
ffff	0	0	0	0
g	0	0	0	0
h	0	0	0	0
iii	0	0	0	0
jj	0	0	0	0
x	0	0	0	0
yy	0	0	0	0
u	0	0	0	0
vvv	0	0	0	0
www	0	0	0	0
zzz	0	0	0	0
TOTAL:	10	10	10	10

Analizând frecvențele rezultă că subvocabularul VC_1 este format din cuvinte ce constituie chei unice în timp ce vocabularele VC_2 , VC_3 și VC_4 au un șir de cuvinte mult mai restrâns și un șir se atribuie mai multor elemente din colectivitatea B.

Compararea șirurilor de caractere este o operație uzuală care evidențiază măsura în care șirurile diferă sau nu unele de celelalte.

Dacă se consideră mulțimea șirurilor de caractere $S = \{C_1, C_2, \dots, C_{nsir}\}$

Se calculează indicatorul de asemănare AS după relația:

$$AS = \frac{KS}{C_{nsir}^2} \quad (29.1)$$

unde:

- KS – numărul de șiruri identice;
- C_{nsir}^2 – numărul combinațiilor din mulțimea S de șiruri luate câte două.

În mulțimea $S_0 = \{aa, bb, cc, dd\}$ perechile care se construiesc sunt în număr de $C_4^2 = 6$. Acestea sunt:

$$(aa, bb), (aa, cc), (aa, dd), (bb, cc), (bb, dd) \text{ și } (cc, dd) \quad (29.2)$$

Aceste perechi de șiruri sunt date diferite, rezultând $KS = 0$ și indicatorul $AS_{(S_0)} = 0$.

În cazul mulțimii de șiruri:

$$S_1 = \{aa, bb, cc, aa\} \quad (29.3)$$

rezultă perechile:

$$(aa, bb), (aa, cc), (aa, aa), (bb, cc), (bb, aa) \text{ și } (cc, aa) \quad (29.4)$$

Cum operația de comparare este comutativă rezultă că operațiile de comparare aplicate perechilor de șiruri (aa, bb) și (bb,aa) conduc la același rezultat. O situație identică se regăsește și în cazul perechilor (aa, cc) și (cc, aa).

Indicatorul $AS_{(S_1)}$ este:

$$AS_{(S_1)} = \frac{3}{6} = \frac{1}{2} \quad (29.5)$$

Se definește indicatorul:

$$AS_{(S_i)} = \frac{C_m^2}{C_n^2} \quad (29.6)$$

unde:

- m – numărul de cuvinte diferite din mulțime;
- n – numărul total de cuvinte;

$$AS_{(S_1)} = \frac{C_3^2}{C_4^2} = \frac{1}{2} \quad (29.7)$$

Cu cât valoarea indicatorului AS tinde către 1, cu atât rezultă că șirurile datorită asemănării dintre nu se constituie într-o mulțime de chei de regăsire a informației. Concatenarea șirurilor este operația prin care două șiruri C_i și C_j se alipesc rezultând un nou șir $C_{ij} = C_i || C_j$. Dacă se consideră mulțimile de șiruri S_1 și S_2 cu același număr de componente

$$S_1 = \{C_1, C_2, \dots, C_{nsir}\} \quad (29.8)$$

$$S_2 = \{d_1, d_2, \dots, d_{nsir}\} \quad (29.9)$$

prin concatenarea elementelor $l_i = c_i || d_i$ rezultă mulțimea:

$$S_3 = \{l_1, l_2, \dots, l_{nsir}\} \quad (29.10)$$

Indicele de asemănare pentru mulțimile S_1 , S_2 și S_3 trebuie să difere astfel încât:

$$AS(S_1) \geq AS(S_3) \text{ și } AS(S_2) \geq AS(S_3) \quad (29.11)$$

numai în cazul în care mulțimile S_1 și S_2 sunt identice aceste condiții nu sunt îndeplinite.

Pentru $S_1 = \{a, b, c, d\}$ și $S_2 = \{x, y, z, w\}$ rezultă valorile:

$$AS(S_1) = 0; \quad (29.12)$$

$$AS(S_2) = 0; \quad (29.13)$$

iar mulțimea S_3 generată, are valorile

$$S_3 = \{ax, by, cz, dw\} \text{ și } AS(S_3) = 0 \quad (29.14)$$

Se observă că în ipoteza $AS(S_1) = 0$, $AS(S_2) = 0$ rezultă $AS(S_3) = 0$.
Pentru $S_1 = \{a, b, c, a\}$ și $S_2 = \{x, y, z, x\}$ rezultă:

$$AS(S_1) = \frac{1}{2}; AS(S_2) = \frac{1}{2}; AS(S_3) = \frac{1}{2}; S_3 = \{ax, by, cz, ax\} \quad (29.15)$$

Pentru $S_1 = \{a, b, c, a\}$ și $S_2 = \{x, y, y, w\}$ rezultă:

$$AS(S_1) = \frac{1}{2}; AS(S_2) = \frac{1}{2}; AS(S_3) = 0; S_3 = \{ax, by, cy, aw\} \quad (29.16)$$

Concatenarea șirurilor corespondente din mulțimi presupune ca:

- mulțimile să aibă același număr de componente;
- operația să se efectueze pe elemente de pe aceeași poziție din fiecare mulțime.

Mulțimile de șiruri se dispun sub forma unor coloane într-un tabel cu MS coloane și $NSIR$ linii dacă numărul de mulțimi de șiruri este MS și numărul de elemente într-un șir este $NSIR$, tabelul 29.5.

Tabelul nr.29.5 Agregarea șirurilor

Poziție	S_1	S_2	...	S_i	...	S_{MS}
E_1	X_{11}	X_{12}	...	X_{1i}	...	X_{1MS}
E_2	X_{21}	X_{22}	...	X_{2i}	...	X_{2MS}
...	...	...	...	...	...	...
E_i	X_{i1}	X_{i2}	...	X_{ij}	...	X_{iMS}
...	...	...	...	...	...	...
E_{NSIR}	X_{NSIR1}	X_{NSIR2}	...	X_{NSIRi}	...	X_{NSIRMS}

în care:

- E_i – poziția elementului în șir;
- S_j – mulțimea j de șiruri;
- X_{ij} – șirul de pe poziția i din mulțime S_j .

Se calculează indicatorii $AS(S_i)$, $i = 1, 2, \dots, MS$. Dacă există mulțimile $S_{i1}, S_{i2}, \dots, S_{ik}$ pentru care indicatorii $AS() = 0$ rezultă că oricare din aceste chei se folosesc drept cheie unică de regăsire a informației.

Se construiesc cele C_{MS}^2 mulțimi $S_j = S_i || S_j$ ale căror elemente rezultă prin concatenarea de elemente corespondente.

Se calculează indicatorii:

$$AS(S_{11}), AS(S_{12}), \dots, AS(S_{ij})_{i>j}, \dots, AS(S_{MS-1 MS}) \quad (29.17)$$

Dacă există indicatori cu valorile $AS = 0$, acele mulțimi se utilizează ca chei unice.

Procedeul de concatenare se extinde la trei, patru mulțimi până rezultă mulțimea cu șirurile cele mai lungi obținute din concatenarea celor MS mulțimi.

$$S_{1, 2, \dots, MS} = \prod_{i=1}^{MS} S_i \quad (29.18)$$

De exemplu, se consideră mulțimile:

$$S_1 = \{a, b, c, d\} \quad (29.19)$$

$$S_2 = \{x, y, z, x\} \quad (29.20)$$

$$S_3 = \{1, 2, 3, 1\} \quad (29.21)$$

$$S_{123} = S_1 || S_2 || S_3 = \{ax1, by2, cz3, ax1\} \quad (29.22)$$

pentru care se determină valorile:

$$AS(S_1) = \frac{1}{2}; AS(S_2) = \frac{1}{2}; AS(S_{123}) = \frac{1}{2} \quad (29.23)$$

Rezultă că operația de concatenare nu a condus la obținerea unor chei unice.

Bordarea este operația prin care unui tablou i se adaugă coloane sau linii. Bordarea unei coloane cu valori pentru care indicatorul AS are valoarea 0 înseamnă adăugarea la tabloul dat a unei coloane de chei unice.

Se obține în acest fel un nou tablou în care elementele se identifică fără dificultate prin chei unice.

Atunci când cheia unică este adresa fizică a articolului sau poziția pe care articolul o ocupă în fișier, baza de date sau în depozitul de date, procesul de regăsire a informației este accelerat.

Dacă cheia unică este invariabilă atunci când rezultă din măsurătorile caracteristicilor, își pierde acest caracter când se produce trecerea pe un alt suport dacă cheia este rezultatul bordării cu o coloană de adrese.

Când cheia unică indică poziția elementului prin inserare se produce schimbarea de poziție ceea ce impune existența unui algoritm recalculare a adresei relative.

29.3 Algoritmi de căutare

Orice tip de aplicație informatică are componente de prelucrare a datelor indiferent de tipul sau sursa de proveniență a acestora. Obiectivul este de a obține rezultate, de a atinge obiectivul final pentru care a fost dezvoltată aplicația, rezolvarea unei probleme bine definite.

Datele prelucrate de aplicație sunt stocate în colecții conform și sunt aranjate conform unor metode considerate de programator importante și eficiente. Prelucrarea acestor date, presupune iterații în care seturi de una sau mai multe valori sunt preluate din mulțime și sunt introduse ca operanzi

în secvențe de instrucțiuni. Pentru aceasta programatorul utilizează rutine de căutare pentru a avea acces la valorile utilizate.

Algoritmii de regăsire au menirea de a identifica în mulțimea de valori, articolul a cărui cheie este specificată.

În funcție de complexitatea operației de căutare, există următoarele tipuri de algoritmi de regăsire:

- algoritmi bazați pe operații de comparare a cheii specificate cu chei din mulțimea de articole; procesul se încheie atunci când cheia specificată este localizat sau când a fost epuizată mulțimea articolelor; performanța algoritmilor este dată de numărul de comparații;
- algoritmi în care se calculează pe baza cheii adresa relativă și adresa fizică a articolului; performanța este dată de numărul de operații și de densitatea poziționării articolelor.

În funcție de locația datelor, operația de căutare are loc:

- în memoria internă a calculatorului; dimensiunea setului de date este redusă însă procesul este caracterizat de viteza mare de execuție datorită efortului redus al procesorului de acces la date; pentru a putea rula rutinele de căutare, la fiecare execuție a aplicației trebuie construit setul de date deoarece acesta nu există în afara memoriei rezervate aplicației;
- în memoria externă, datele fiind stocate în fișiere pe diverse suporturi de stocare; dimensiunea setului de date este foarte mare, iar procesul este caracterizat de viteza redusă de prelucrare datorită efortului procesor ridicat dat de citirea datelor și aducerea acestora în memoria internă; datele sunt independente de aplicație din punctul de vedere al existenței acestora, nefiind nevoie să fie recreate de fiecare dată când aplicația este lansată în execuție;
- în memoria internă și externă; se combină tehnicile de căutare în memoria internă, aplicate unor structuri auxiliare, cu metodele de acces la date stocate pe disc.

Legătura de interdependență dintre procesul de căutare și structurile de date are la baza necesitatea de a stoca informația fie în memoria virtuală sau în cea externă într-o formă care să asigure integritatea datelor și care să permită ulterior citirea acestora.

Căutarea secvențială este procesul prin care datele sunt regăsite prin parcurgerea lineară a tuturor elementelor din mulțimea de date pornind de la începutul acesteia. Procesul se finalizează cu regăsirea elementului căutat sau cu atingerea sfârșitului colecției de date, lucru care indică terminarea operației de căutare. Figura 29.2 descrie modul în care este aplicată căutarea secvențială în cadrul unei colecții de date formată din elementele $Elem_1, Elem_2, \dots, Elem_n$ pentru a găsi valoarea $XVal$.

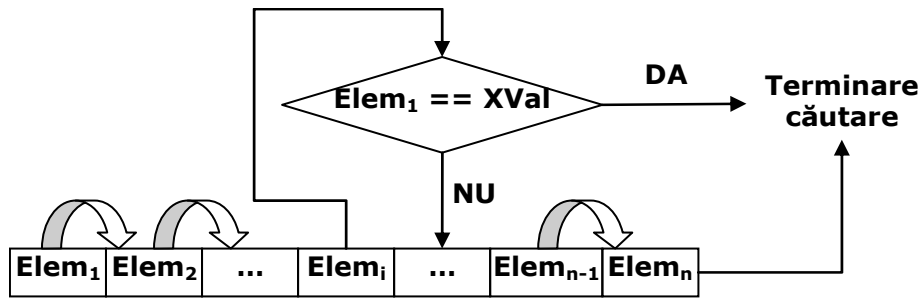


Figura 29.2 Căutarea secvențială a unei valori într-o colecție

Pentru a implementa într-un limbaj de programare această metodă se ține seama că aceeași rutină de verificare a valorii căutate se aplică pentru fiecare element în parte.

```

int SequentialSearch(int * array, int arrayLength, int searchvalue){
    int i;
    for( i = 0; i < arrayLength; i++){
        if(array[i] == searchvalue)
            return i;
    }
    return -1;
}
  
```

În cazul în care structura de date utilizată pentru a stoca valorile mulțimii este de tip listă simplu înlănțuită atunci rutina utilizată pentru căutarea secvențială a unei valori este *ListSequentialSearch*.

```

struct SLList
{
    int value;
    SLList *next;
};
SLList* ListSequentialSearch(SLList * head, int searchvalue){
    SLList *temp;
    for( temp = head; temp!=NULL; temp = temp->next ){
        if(temp->value == searchvalue)
            return temp;
    }
    return NULL;
}
  
```

Această metodă de căutare este îmbunătățită prin reducerea numărului de comparații realizate în cadrul ciclului repetitiv, [Knut73], tehnică specifică optimizării la nivel de text sursă, [Ivan07a].

```

int QuickSequentialSearch(int * array, int arrayLength,
int searchvalue){
    int i;
    array[arrayLength] = searchvalue;
    for( i = 0;; i++){
        if(array[i] == searchvalue)
            if(i==arrayLength)
                return -1;
            else
                return i; }}
  
```

Se observă că această metodă diferă prin:

- utilizarea unui element suplimentar pe poziția *arrayLength* în cadrul vectorului pentru a fi utilizat la verificarea condiției de terminare a căutării;
- reducerea la unu a numărului de condiții ce sunt evaluate la fiecare iterație a ciclului repetitiv;
- inserarea unei noi condiții de verificare a terminării căutării; procesul se încheie de fiecare dată cu găsirea elementului căutat, însă dacă acesta se găsește pe poziția *arrayLength* rezultă că elementul nu se regăsește în mulțime.

Această simplă modificare conduce la reducerea numărului de cicluri mașină necesare realizării căutării. În ciuda faptului că a doua metodă necesită un element suplimentar în cadrul vectorului, volumul redus de cicluri mașină reprezintă un efect mult mai important al procesului de optimizare.

Pe baza metodelor de optimizare a textului sursă, [Ivan07a] se maximizează viteza de căutare prin desfășurarea ciclului repetitiv și multiplicarea numărului de elemente verificate. De exemplu secvența

```
int FastSecquentialSearch(int * array, int arrayLength,
int searchvalue){
    int i;
    array[arrayLength] = searchvalue;
    for( i = 0;; i+=3 ){
        if(array[i] == searchvalue)
            if(i==arrayLength)
                return -1;
            else
                return i;
        if(array[i+1] == searchvalue)
            if(i==arrayLength)
                return -1;
            else
                return i+1;
        if(array[i+1] == searchvalue)
            if(i==arrayLength)
                return -1;
            else
                return i;
        if(array[i+2] == searchvalue)
            if(i==arrayLength)
                return -1;
            else
                return i+2;
    }
}
```

crește la trei numărul de elemente verificate în cadrul fiecărei iterații. Deoarece ciclul repetitiv se termină cu regăsirea elementului căutat pe o poziție $0 \leq i \leq \text{arrayLength}$ sau pe poziția *arrayLength* nu are loc depășirea acestuia și nu se generează o eroare de citire în afara vectorului.

Pe baza metodei parcurgerii secvențiale a unei colecții se definește căutarea cu elemente sortate după valoarea cheii.

```

void SortSecventialSearch(int * array, int arrayLength,
int searchvalue){
    int i;
    for( i = 0; i < arrayLength; i++){
        if(array[i] >= searchvalue)
            break;
    }
    if(array[i]==searchvalue)
        return i;
    else
        return -1;
}

```

În cazul în care valoarea căutată se regăsește în mulțimea de valori, această metodă nu îmbunătățește nivelul de eficiență al căutării secvențiale, funcția *Secventia/Search*. Rezultatele, minimizarea efortului prin reducerea numărului de valori verificate, sunt evidente în cazul în care valoarea căutată nu se află în mulțime. În această situație, ciclul repetitiv este întrerupt dacă se întâlnește un element cu valoarea mai mare decât cel căutat și nu se mai parcurge tot vectorul.

O altă metodă de căutare secvențială se bazează pe frecvența de referire a elementelor. Cu cât frecvența unui element este mai ridicată, cu atât acestea sunt mai la început și astfel numărul de comparații regăsirii lor scade.

De exemplu, se consideră șirul cheilor de căutare pentru care se înregistrează frecvențele

$$\begin{array}{cccc}
 a & b & c & d \\
 1 & 1 & 3 & 7
 \end{array} \quad (29.24)$$

Datele sunt aranjate astfel încât, cuvintele să fie sortate în ordinea descrescătoare a frecvențelor

$$7d \quad 3c \quad 1a \quad 1b \quad (29.25)$$

Utilizarea unei astfel de structuri presupune implementarea unui câmp în structura de date utilizată pentru a înregistra numărul de căutări. De exemplu, se consideră problema stocării studenților dintr-o facultate pentru care se implementează structura

```

struct Student
{
    char nume[20];
    int varsta;
    int cod
    int frecventa;
};

```

Metoda utilizată la căutarea unui student după codul acestuia este identică cu rutina *Secventia/Search* din punctul de vedere al algoritmului implementat, însă ia în considerare actualizarea câmpului *frecventa* pentru elementele căutate și resortarea mulțimii.

Eficiența acestei metode în raport cu o căutare secvențială pe o mulțime neordonată sau sortată se bazează pe faptul că în practică, fiecărui

element din mulțime i se asociază o probabilitate de utilizare în funcție problema ce trebuie rezolvată.

Există numeroase studii în acest domeniu, [Knut73], însă două reguli stabilesc importanța unui set restrâns de elemente, cu frecvență ridicată de utilizare, în cadrul operațiilor de căutare:

- conform legii lui Zipf, [Wiki07a], probabilitate de utilizare a elementelor dintr-o mulțime respectă o distribuție dată de relația

$$p_1 = \frac{c}{1}, p_2 = \frac{c}{2}, \dots, p_N = \frac{c}{N}, \text{ cu } c = \frac{1}{\sum_{i=1}^N \frac{1}{i}} \quad (29.26)$$

acesta a observat în că frecvența celui mai utilizat k cuvânt este invers proporțională cu poziția sa în tabela de frecvențe; astfel cel mai utilizat cuvânt va apărea în texte de aproximativ de două ori mai mult decât al doilea cel mai utilizat cuvânt; acesta, la rândul său are de două ori mai multe apariții decât al patrulea cuvânt din lista de frecvențe; spre deosebire de alte modele de distribuție artificiale, această teorie a fost stabilită pe cale empirică și de multe ori există situații reale în care se respectă;

- legea 80-20, sau principiul Pareto, [Wiki07a], este un model empiric definit de economistul Vilfredo Pareto pe baza analizei veniturilor Italiei; acesta a observat că 80% din veniturile Italiei provin de la 20% din populație; principiul a fost extins pe baza de cercetări și în domeniu optimizării aplicațiilor informatice prin faptul că 80% din resurse sunt utilizate de către 20% din componentele aplicației; la nivel de viteză de execuție, se observă că 90% din prelucrările efectuate de program au loc la nivelul a doar 10% din codul sursă;

Pentru a evita definirea de câmpuri suplimentare în interiorul structurii de date cu scopul de a gestiona frecvența de utilizare, este definită o metodă care se bazează pe principiul repoziționării la începutul colecției de date a fiecărui element căutat. Astfel, se obține același rezultat, prin poziționarea elementelor căutate des la începutul mulțimii.

Repoziționarea unui element în cadrul unui vector implică recopierea elementelor, anterioare elementului mutat, pe poziții următoare. Acest lucru implică un efort suplimentar de prelucrare a datelor. Pentru a fi evitat, colecția de date este reprezentată printr-o listă simplă înlănțuită. În această situație, repoziționarea este echivalentă cu inițializarea unui set restrâns de pointeri.

Căutarea secvențială cu elemente sortate după valoarea cheii este îmbunătățită prin algoritmul de căutare binară, ce presupune începerea procesului dintr-un punct mai apropiat de valoarea căutată decât de la început. Primul reper în procesul de căutare este definit de poziția din mijloc a șirului de valori sortate. În funcție de mărimea relativă a valorii căutate față de valoarea mediană, căutarea continuă în șirul elementelor aflate la stânga sau la dreapta acesteia prin repetarea procesului.

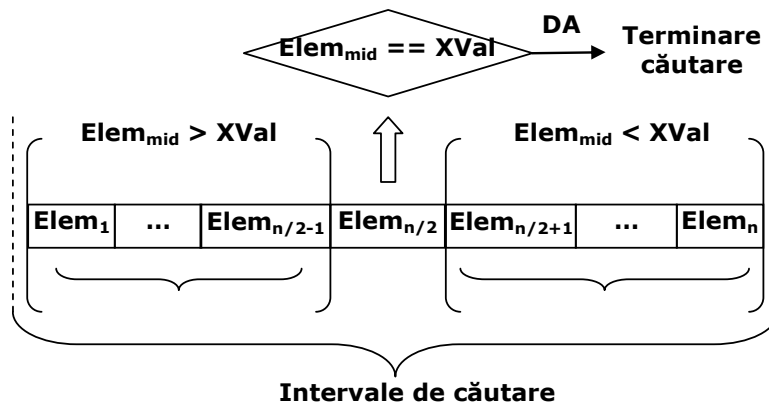


Figura 29.3 Căutarea binară a unei valori într-o colecție sortată crescător

Presupunând că procesul începe cu un vector de lungime $LV = n$ atunci poziția mediană este determinată prin relația

$$mediana = (poz_{start} + poz_{final})/2 \quad (29.27)$$

unde:

- *mediana* – poziția mediană;
- *poz_{start}* – poziția de start a intervalului de căutare; inițial aceasta este egală cu 0;
- *poz_{final}* – poziția finală a intervalului de căutare; inițial aceasta este egală cu *n*.

Prin compararea valorii căutate cu valoarea de pe poziția mediană se determină:

- momentul găsirii, dacă cele două valori sunt egale;
- continuarea căutării în intervalul inferior, dacă valoarea căutată este mai mică;
- continuarea căutării în intervalul superior, dacă valoarea căutată este mai mică;
- încetarea procesului dacă valorile nu sunt egale și nu există alte intervale de valori.

```
void BinaraSearch(int * array, int arrayLength, int searchvalue){
    int middlevalue;
    int start = 0;
    int end = arrayLength-1;
    int middle;
    while(start <= end){
        middle = (start + end)/2;
        middlevalue = array[middle];
        if(searchvalue == middlevalue){
            printf("\n SEARCH VALUE FOUND !");
            return;
        }
        else
            if(searchvalue < middlevalue)
                end = middle -1;
            else
                start = middle +1;
    }
}
```


Față de căutarea secvențială, căutarea binară reduce cu mult numărul de comparații, deoarece la fiecare iterație se elimină din aria de căutare jumătate din valorile neverificate.

Între tipurile de algoritmi de regăsire și tipurile de structuri de date utilizate în programare există o legătură strânsă de interdependență. Semnificația legăturii este evidențiată și de faptul că au fost definite tipuri particulare de structuri de date pentru a permite minimizarea efortului de căutare și regăsire a informațiilor. Astfel de exemple, sunt definite de structurile arborescente de căutare:

- arborii binari în care fiecare nod numit părinte are maxim două noduri numite fiu între care există relația

$$Val_{fiu_st\u00e2nga} < Val_{parinte} < Val_{fiu_dreapta} \quad (29.28)$$

regăsirea informației în arborii binari de căutare se face prin parcurgerea structurii și prin verificarea în fiecare nod a egalității dintre valoarea căutată și cheia nodului curent; în cazul în care nu există egalitate, se continuă procesul pe fiul din stânga sau pe cel din dreapta în funcție de sensul relației dintre valorile comparate

```
struct BST
{
    int info;
    BST * right;
    BST * left;
}

void BinarySearchTreeSearch(BST * root, int info)
{
    if(! root)
        return;
    else
    {
        if(info == root->info)
        {
            printf("\n SEARCH VALUE FOUND !");
            return;
        }
        else
        {
            if(info > root->info)
                BinarySearchTreeSearch(root->right, info);
            else
                BinarySearchTreeSearch(root->left, info);
        }
    }
}
```

- pentru a crește eficiența operației de căutare pe acest tip de structură, se iau în considerare tipuri particulare de arbori binari, ce sunt echilibrați, AVL sau Roșu & Negru; structurile echilibrate contribuie la creșterea eficienței algoritmului de căutare prin rearanjarea valorilor astfel încât să fie minimizat efectul situației cel mai puțin favorabile; în cazul arborilor binari de căutare, situația cel mai puțin favorabilă este dată de obținerea unei structuri arborescente cu înălțimea egală cu numărul de

elemente, figura 29.4 (a); reechilibrarea acestei structuri, figura 3.4 (b), minimizează timpul de căutare deoarece înălțimea structurii este minimizată și astfel scade numărul maxim de comparații ce trebuie efectuate pentru a găsi valoarea căutată;

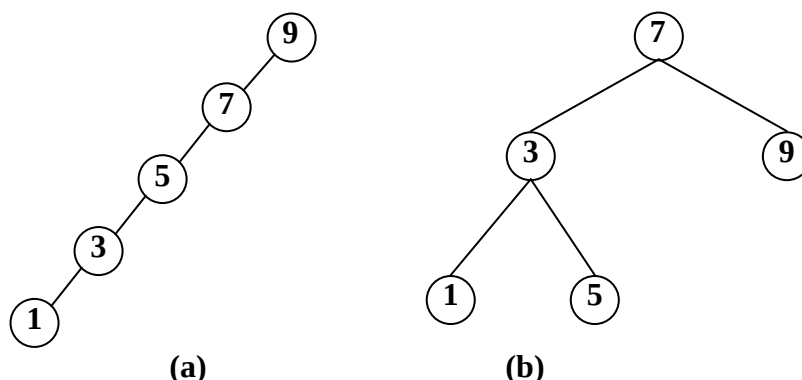


Figura 29.4 Tipuri de structuri arborescente binare de căutare.

echilibrarea arborilor de căutare nu influențează rutina de căutare, deoarece relația dintre valoarea nodului părinte și valorile nodurilor fiu se menține; diferența se observă la nivelul secvențelor de cod de inserare, respectiv de ștergere a unei valori, deoarece structura trebuie reechilibrată după fiecare operație de acest tip;

- arborii B sunt structuri arborescente de ordin N echilibrate, în care un nod poate conține maxim $2 \cdot N$ chei și minimum N chei, cu excepția nodului rădăcină care poate avea mai puțin de N chei; această situație este întâlnită în momentul construirii arborelui B, atunci când au fost inserate mai puțin de N chei; pentru fiecare pereche de chei $KValue_i$ și $KValue_{i+1}$ cu $i = 1..n$ există un nod fiu NF_{i+1} pentru care:

$$KValue_i < Valori_chei\{NF_{i+1}\} < KValue_{i+1} \quad (29.29)$$

Figura 29.5 descrie modul în care este definit un nod al arborelui B de ordin N.

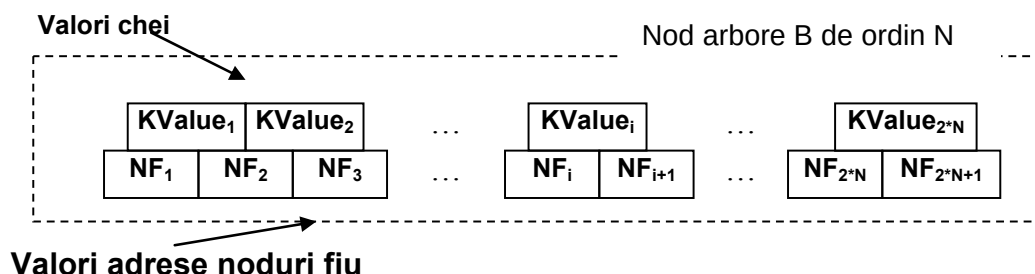


Figura 29.5 Structura unui nod al arborelui B

Dacă se consideră implementarea structurii arborelui B prin intermediul listelor de chei și de noduri fiu atunci secvența de cod de căutarea a unei valori este

```

typedef struct NodArboreB_liste;

struct ListaChei
{
    int valCheie;
    ListaChei * next;
};

struct ListaFii
{
    NodArboreB_liste * Fiu;
    ListaFii * next;
};

struct NodArboreB_liste
{
    int NrChei;
    ListaChei *capListaChei;
    NodArboreB_liste *NodParinte;
    ListaFii *capListaFii;
};

struct NodArboreB
{
    int chei[MAX_CHEI];
    int NrChei;
    NodArboreB * NodParinte;
    NodArboreB * NodFiu[MAX_CHEI+1];
};

void cautaValoare(NodArboreB *NodStart,int ValCheie, NodArboreB *&
NodGasit, int &pozCheie)
{
    if(NodStart){
        int nrCheiNod = NodStart->NrChei;
        if((ValCheie<NodStart->chei[0])&&(ValCheie>NodStart->
chei[nrCheiNod-1]))
            if((NodStart->NodFiu[0] == NULL) && (NodStart->
NodFiu[nrCheiNod])){
                Printf("\n Cheie inexistentă !");
                return;
            }
        for(int i=0;i<nrCheiNod;i++){
            if(ValCheie==NodStart->chei[i]){
                NodGasit=NodStart;
                pozCheie=i;
            }
            else
                if(ValCheie>NodStart->chei[i])
                    break;
        }
        if((ValCheie<NodStart->chei[i]) && (NodStart->NodFiu[i]!=NULL))
            cautaValoare(NodStart->NodFiu[i], ValCheie, NodGasit, pozCheie);
        if((ValCheie>NodStart->chei[i-1]) && (NodStart-> NodFiu[i]!=NULL))
            cautaValoare(NodStart->NodFiu[i], ValCheie, NodGasit, pozCheie);
    }
};

```

Pentru a gestiona informații stocate pe suport extern, limbajele de programare pun la dispoziția programatorilor funcții de lucru cu fișiere. În funcție de tehnicile de stocare și regăsire a datelor pe suport extern, sunt definite:

- fișiere cu acces secvențial în care adăugarea de noi înregistrări se face întotdeauna la sfârșit, neexistând o procedură specială de aranjare a datelor; fișierele cu acces secvențial sunt asemănătoare ca modalitate de lucru cu masivele unidimensionale și din acest motiv căutarea se realizează prin parcurgerea fiecărei înregistrări până când se găsește informația căutată sau până când se ajunge la sfârșitul fișierului;

```
struct Student{
int ID;
char Name[10];
int Age;
}
void FILESearch(char * FileName, int SearchValue){
    Student value;
    FILE *pFile = fopen(FileName,"rb");
    while(!feof(pFile)){
        fread(&value,sizeof(Student),1, pFile);
        if(SearchValue == value){
            printf("\n FILE SEARCH VALUE FOUND !");
            return;
        }
    }
}
```

- fișierele cu acces direct sunt structuri de date externe, ce se regăsesc pe disc, în care există o corespondență directă între poziția în fișier a unui element și valoarea cheii de căutare; deoarece această legătură se bazează pe unicitatea cheii și pe ipoteza că numărul de articole din fișier este egal cu valoarea maximă a cheii de căutare, fiecare înregistrare din fișier se găsește în una din stările, validă sau ștearsă logică;
- fișierele de tip invers permit căutarea de informații pe bază de chei compuse; sunt utilizate în situațiile în care se caută submulțimi de articole ce corespund unei filtrări multicriteriale; pentru a minimiza spațiul ocupat de fișierul invers se vor implementa structuri pe biți ce corespund criteriilor de căutare, [Smeu04]; se consideră problema gestiunii elevilor dintr-un județ; pentru fiecare se înregistrează date descrise de structura *elev*;

```
struct elev
{
    char sex;
    char naționalitate;
    char cetatenie;
    char mediu;
    char nume[20];
    char note[10];
    int nrmatricol;
}
```

Deoarece o parte din câmpuri au valori fixe:

- câmpul *sex* ia valori în mulțimea {0 – masculin, 1 – feminin};
- câmpul *nationalitate* ia valori în mulțimea {0 – română, 1 – romă, 2 – maghiară, 3 – alta};
- câmpul *cetatenie* are valorile {0 – română, 1 – alta};
- câmpul *mediu* are valorile {0 – rural, 1 – urban};

Pentru că aplicația trebuie să permită căutări după aceste criterii se definește o structură pe biți:

```
struct elevI
{
    unsigned sex:1;
    unsigned nationalitate:2;
    unsigned cetatenie:1;
    unsigned mediu:1;
    unsigned :3;
};
```

ce definește elementele fișierului de tip invers utilizat în procesul de căutare; motivul implementării unei structuri pe biți este dat de minimizarea spațiului; setul de date al elevilor este gestionat prin intermediul a două fișiere, unul ce conține toate informațiile și unul ce conține valorile elementelor de tip *elevI*; legătura dintre cele două structuri de date externe este definită de corespondența ce există între elementele de pe poziții identice; rolul fișierului invers, ce conține doar valorile criteriilor utilizate la căutare, este de a permite implementarea unui proces de căutare mult mai eficient; operația de căutare multicriterială se realizează prin intermediul variabilelor de tip *mască* și *filtru*; masca este utilizată pentru a testa dacă elementul analizat îndeplinește criteriul de selecție prin intermediul operației de sau exclusiv pe biți; filtrul este utilizat pentru a aduce în prim plan doar acele câmpuri care compun condiția multicriterială de selecție; figura 29.6 descrie procesul de selecție a elevilor care au naționalitatea română și provin din mediul urban;

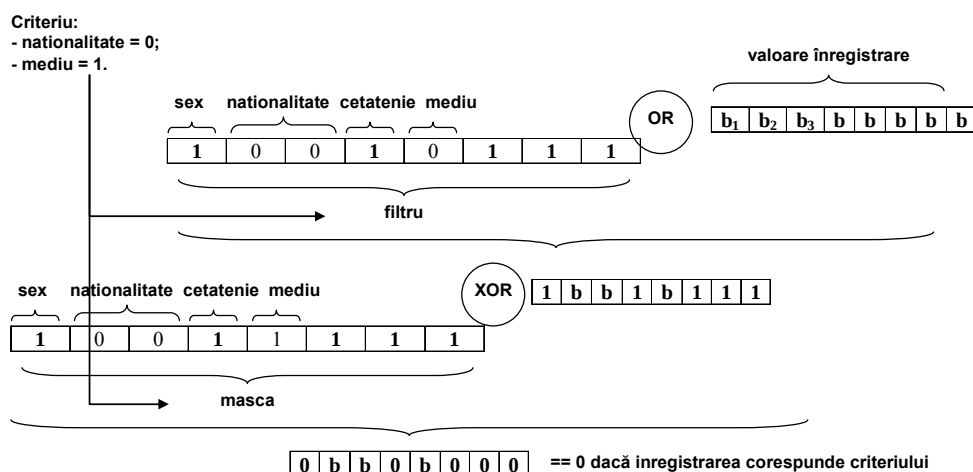


Figura 29.6 Selecția multicriterială prin intermediul fișierelor de tip invers

- fișierele indexate sunt fișiere secvențiale pentru care se construiesc o structură auxiliară de tip arbore care să eficientizeze

procesul de căutare; spre deosebire de fișiere secvențiale, căutarea înregistrării după o chei unică se face în totalitate în structura arborescentă asociată fișierului; deoarece această structură are forma:

```
struct IndexFisier
{
    int ValoareCheie;
    int PozitieFisier;
};
```

odată ce este găsită valoarea căutată, este determinată poziția în fișier a înregistrării respective; evitând o serie de citiri secvențiale din fișier prin secvența de instrucțiuni

```
int PozInregistrare = PozitieFisier - 1;
fseek(Fisier, PozInregistrare * sizeof(Inregistrare), SEEK_SET);
```

este citită înregistrarea căutată;

Cel mai eficient algoritm de căutare este acela în care fiecărei valori din colecția de elemente *i* se asociază o poziție unică în cadrul mulțimii. Astfel, pe baza valorii căutate se determină poziția acesteia în cadrul colecției.

În cadrul vectorilor, această abordare se implementează cu ușurință deoarece asocierea dintre valoarea unei chei numerice întreagă și poziția acesteia în vector se realizează prin:

- definirea unui vector cu număr de elemente egală cu valoarea maximă posibilă a cheii de căutare;
- se stabilește o valoare neutilizată în cadrul problemei de rezolvat pentru a indica dacă elementul cu cheia căutată există; deoarece vectorul reprezintă o zonă de memorie continuă, se implementează principiul conform căruia elementele există sau sunt șterse logic; pentru o cheie de căutare ce ia valori în intervalul [0...255], valoarea -1 poate fi utilizată pentru a indica inexistența elementului căutat în colecție.

Principalul avantaj al căutării bazate pe acces direct este dat de rapiditate cu care se găsește elementul căutat sau este indicată inexistența acestuia.

Pentru structura *element* funcția de regăsire este descrisă în secvența de cod

```
struct element
{
    int cheie;
    int valoare;
};
int DirectAccess(element* array, int arrayLength, int searchvalue)
{
    if(searchvalue >= arrayLength) return -1;
    else if(array[searchvalue].cheie == -1) return -1;
    else return array[searchvalue].valoare;
}
```

Cu toate acestea, în practică este evitată această soluție directă care, deși minimizează timpul de regăsire are efecte negative asupra:

- memoriei ocupate; dimensiunea spațiului, MEM, necesar implementării acestei structuri este dat de relația

$$MEM = \max(\text{valoare cheie căutare}) * \text{dimensiune}(\text{element}) \quad (29.30)$$

iar în cazul în care valoarea maximă este foarte mare atunci trebuie să fie rezervat un spațiu considerabil;

- gradului de utilizare a spațiului; deoarece dimensiunea structurii este dependentă de valoarea maximă a cheii de căutare, nu se ține cont de numărul real de elemente utilizate; cea mai nefavorabilă situație este dată de un raport foarte mic dintre numărul de elemente și valoarea maximă a cheii; de exemplu, se consideră o aplicație informatică ce gestionează studenții din cadrul unei facultăți reprezentați de structura:

```
struct Student
{
    char nume[20];
    int varsta;
    char facultate[20];
    int nrMatricol;
}
```

pentru a minimiza timpul de regăsire, se implementează o structură cu acces direct în care numărul matricol al studentului este egal cu poziția în cadrul colecției a respectivului element; în cazul în care valoarea maximă a câmpului *nrMatricol* este egală 55630, iar numărul real de studenți este egal cu 1450 rezultă un spațiu ocupat egal cu

$$MEM = \max(nrMatricol) * dimensiune(Student) = 2,54 \text{ MB} \quad (29.31)$$

iar gradul de utilizare a acestui spațiu este egal cu:

$$GUM = \frac{MEM_{\text{elemente}}}{MEM} * 100 = \frac{1450 * 48}{55630 * 48} * 100 = 2,6\% \quad (29.32)$$

- tipului cheii de căutare; acesta trebuie să aibă un tip numeric întreg deoarece trebuie să corespundă indexului cu care se accesează elementele unui vector.

Pentru a remedia dezavantajele utilizării structurilor cu acces direct sunt definite tabelele de dispersie, *hash tables*, ce reprezintă colecții de date în care pe baza unei funcții *hash* cheia de căutare este pusă în corespondență cu poziția elementului în cadrul colecției. Avantajul acestei tip de structură de stocare și căutare este dat de:

- utilizarea mai eficientă a spațiului fără a se stoca memorie pentru elemente care nu sunt utilizate;
- implementarea de chei alfanumerice; prin utilizarea unei funcții care să prelucreze cheia de căutare este depășită bariera care limita pentru structurile cu acces direct tipul cheii de căutare la o

valoarea numerică întreagă; în această situație rolul funcției *hash* este de a translata valoarea alfanumerică într-o valoare întreagă pozitivă.

Dezavantajul tabelelor de dispersie în cadrul proceselor de căutare este descris de:

- efortul suplimentar de prelucrare dat de funcția *hash* care poate avea în unele situații un nivel de complexitate ridicat;
- apariția în cadrul tabelii a coliziunilor; acestea sunt generate de situația în care două valori X_h și Y_h ce aparțin domeniului de definiție a funcției *hash* conduc la obținerea de valori identice, $\text{hash}(X_h) = \text{hash}(Y_h)$; evitarea coliziunilor implică realizarea de operații suplimentare prin care să se identifice poziția elementului în cadrul mulțimii.

Din aceste puncte de vedere, tabela de dispersie utilizată în procesul de regăsire a informației se descrie ca o structură în care datele sunt stocate în tabele cu adresare directă, iar poziția acestora este determinată prin prelucrarea cheii de căutare cu o funcție *hash*, figura 29.7.

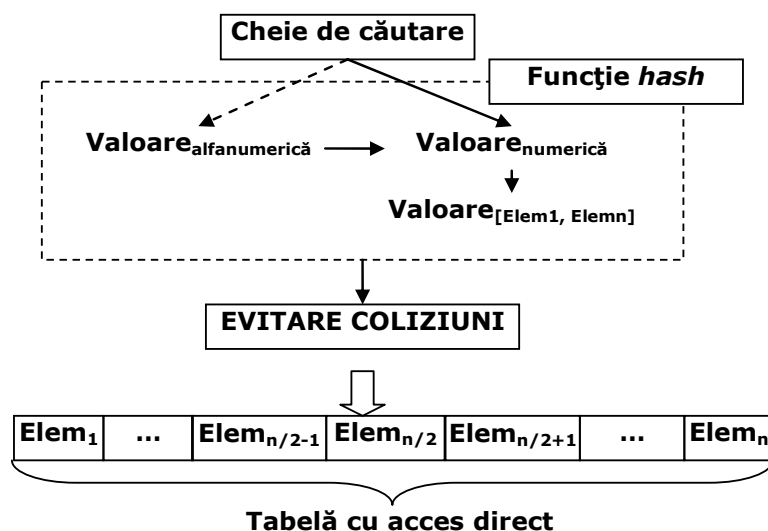


Figura 29.7 Căutarea în tabela de dispersie

Funcțiile de evitare a coliziunilor, descrise în cadrul capitolului dedicat tabelilor de dispersie, definesc metode de regăsire a elementelor ce sunt descrise de chei cu valori diferite dar care conduc la valori *hash* identice ceea ce implică poziții identice în cadrul structurii.

Analiza comparată a metodelor de căutare în vederea alegerii metodei care să dea cele mai bune rezultate se bazează pe:

- măsurarea timpului de execuție; metoda optimă este caracterizată de cea mai mică valoare înregistrată;
- măsurarea efortului procesor prin prisma numărului de cicluri mașină realizate pentru a prelucra secvența de cod asociată metodei; codul care va genera cel mai mic nivel de cicluri mașină evidențiază algoritmul care se va executa cel mai repede pentru datele de test și care este cel mai eficient din punctul de vedere al sistemului;

- măsurarea zonei de memorie utilizată; algoritmul care va necesita cel mai mic nivel de memorie internă pentru a prelucra datele este considerat cel mai eficient din punct de vedere a spațiului.
- analiza complexității algoritmului din punctul de vedere a calculelor efectuate și alegerea modelului care este cel mai eficient.

Primele trei tehnici presupun realizarea unui program care să implementeze metodele de căutare analizate și generarea unei proceduri de testare cu înregistrarea directă a valorilor indicatorilor. Ultima abordare ia în considerare o analiză a modelului asociat algoritmului de căutare înainte ca acesta să fie implementat.

Această tehnică se aplică în faza de analiză a problemei de rezolvat și din acest motiv face abstracție de mediul de implementare neluând în calcul:

- platforma hardware și software pe care va rula programul;
- memoria disponibilă;
- limbajul de programare utilizat;
- tipul și relațiile dintre datele de test.

Indicatorii utilizați pentru a determina complexitatea algoritmului au caracter abstract și descriu metoda prin prisma structurii interne, a numărului de instrucțiuni de comparare și a numărului de iterații.

Din punct de vedere matematic, fiecare algoritm de căutare are asociat un model bazat pe o funcție $f(n)$ care descrie modul în care acesta este aplicat. De exemplu, în cazul căutării secvențiale într-un vector *vector*, cu n elemente, modelul asociat are forma:

$$f_{\text{cautaresecv}}(n) = \begin{cases} \text{dacă vector}[0] = X\text{Val, element gasit} \\ \text{dacă vector}[1] = X\text{Val, element gasit} \\ \dots\dots\dots \\ \text{dacă vector}[n-1] = X\text{Val, element gasit} \\ \text{altfel, element inexistent} \end{cases}$$

Figura 29.8 Model asociat căutării secvențiale

Dacă se consideră că principala operație pe care e bazează metoda de căutare secvențială este compararea elementului curent din vector cu valoarea căutată, atunci complexitatea algoritmului este egală cu n , fiind notată $O(n)$.

În seturi reale de date de test, elementul poate fi găsit după o singură operație de căutare, după $n/2$ verificări sau acesta se poate găsi pe ultima poziție din vector. Din acest motiv și pentru a normaliza valorile înregistrate pentru complexitatea algoritmilor analizați se iau în calcul cazuri generale de execuție:

- cel mai bun caz descrie situația în care este necesar minimul de efort procesor pentru a regăsi valoarea; în practică are probabilitatea cea mai mică de apariție; în cazul căutării secvențiale a unei valori într-o listă simplu înlănțuită, cel mai bun caz este dat de găsirea elementului căutat pe prima poziție;

- cazul mediu descrie o analiză a unui set de operații de căutare pe baza căruia se determină un efort mediu necesar pentru a regăsi un element; în cazul algoritmilor complecși este dificil de determinat acest nivel mediu;
- cel mai ineficient caz reprezintă situația în care este necesar cel mai ridicat nivel de efort pentru a căuta o valoare; pentru o analiză obiectivă a algoritmilor de căutare se utilizează acest criteriu în vederea alegerii rutinei mai eficiente; în cazul căutării secvențiale a unei valori într-o listă simplu înlănțuită, cel mai ineficient caz este dat de găsirea elementului căutat pe ultima poziție.

Pe baza modelului asociat algoritmului de prelucrare se determină nivelul de complexitate al acestuia. Chiar dacă funcția matematică pe baza căreia este definit modelul conține numeroși factori, pentru a determina gradul de complexitate se ia în considerare dominantul relației.

Din punct de vedere matematic, [Rile87], o funcție, $f(x)$, domină o altă funcție, $g(x)$, dacă există o valoare constantă *const* astfel încât pentru orice valoare x este adevărată relația

$$const * f(x) > g(x) \quad (29.33)$$

În domeniul informatic, pentru a determina nivelul de complexitate al unui algoritm se utilizează o altă versiune a relației anterioare, numită dominare asimptotică. Conform acestei teoreme, o funcție $f(x)$ domină asimptotic o altă funcție $g(x)$ dacă $f(x)$ domină pe $g(x)$ pentru toate valorile mari ale lui x .

Dacă se consideră funcțiile

$$f(x) = 2 * x^3 + 1 \quad (29.34)$$

$$g(x) = 5 * x^2 + 3 * x + 10 \quad (29.35)$$

se observă din graficul 29.9 că pe intervalul $[1;100]$ ambele funcții cresc constant, însă de la valoarea 31 funcția $f(x)$ prezintă o creștere mult mai rapidă decât $g(x)$.

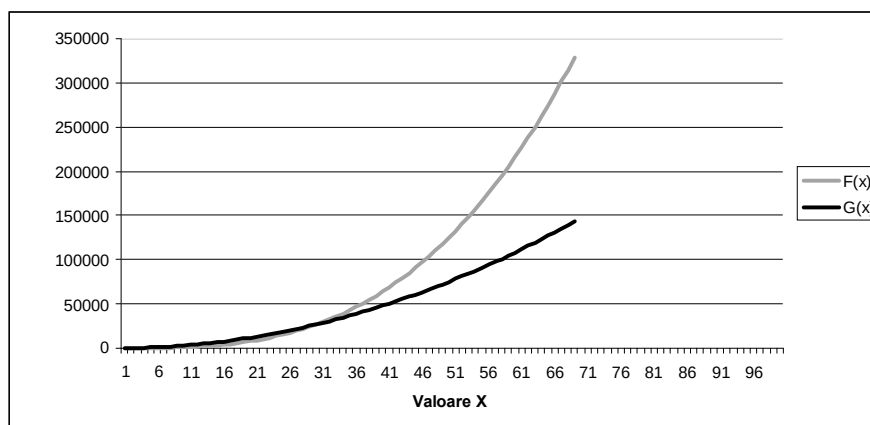


Figura 29.9 Exemplu funcții dominante

În [Wiki07b] se prezintă principalele caracteristici și reguli de determinare a nivelului complexității, indicator denumit *big-O*.

Principalele niveluri de complexitate regăsite la structurile de prelucrare a datelor, căutare sau sortare, sunt bazate pe funcții:

- constante, de tipul $f(n) = 1$;
- liniare, de tipul $f(n) = n$;
- logaritmice, de tipul $f(n) = \log_2 n$;
- pătratice, de tipul $f(n) = n^2$;
- cubice, de tipul $f(n) = n^3$;
- polinomial, de tipul $f(n) = n^c$, cu $c > 1$;
- exponențiale, de tipul $f(n) = 2^n$ sau $f(n) = a^n$, cu $a > 1$.
- factoriale, de tipul $f(n) = n!$

Tabelul 29.6 descrie comparativ influența pe care o au diferitele grade de complexitate asupra efortului de prelucrare

Tabelul nr. 29.6 Analiză comparată a nivelurilor de complexitate

Valoare n	$f(n) = 1$	$f(n) = n$	$f(n) = \log_2 n$	$f(n) = n^2$	$f(n) = 2^n$
10	1	10	3.32	100	1024
100	1	100	6.64	10000	$1,26 * 10^{30}$
1000	1	1000	9.97	1000000	-
10000	1	10000	13.29	100000000	-

Pentru a analiza complexitatea structurilor de căutare descrise se ia în considerare cazul cel mai puțin favorabil.

Tabelul 29.7 descrie complexitatea algoritmilor de căutare analizați.

Tabelul nr. 29.7 Complexitatea metodelor de căutare, cazul cel mai nefavorabil

	Complexitate
metode de căutare în memoria internă	
Căutare cu acces direct	$O(1)$
Căutare secvențială	$O(n)$
Căutare binară	$O(\log_2 n)$
Căutarea în tabele de dispersie	$O(GU_{hash})$
Căutare în arbori binari de căutare echilibrați	$O(\log_2 n)$
Căutare în arbori B	$1 + \log_N((n+1)/2)$, unde N este ordinul arborelui
metode de căutare în memoria externă	
Căutare secvențială în fișier	$O(n)$
Căutare cu acces direct în fișier	$O(1)$
Căutare în fișier indexate	$O(\log_2 n)$ pentru index implementat pe arbori binari de căutare echilibrați
Căutare în fișier inverse	$O(n)$

unde GU_{hash} descrie gradul de utilizare a tabelii de dispersie.

Se observă că din punctul de vedere al complexității utilizarea metodelor de căutare în vector și în fișier secvențial conduce la același nivel

de eficiență. În cazul implementării celor două metode, se observă că citirea unui element dintr-un vector este mult mai rapidă decât citirea unui element din fișier. De exemplu secvența

```
for( i = 0; i < arrayLength; i++ ){
    if(array[i] == searchvalue)
        printf("\n SEARCH VALUE FOUND !");
}
```

necesită pentru un vector de 30000 de elemente distincte NCMV = 835 cicluri mașină, fiind de două ori mai eficientă decât secvența

```
int value;
FILE *pFile = fopen("SearchTest.dat","rb");
while(!feof(pFile))
{
    fread(&value,sizeof(int),1, pFile);
    if(searchvalue == value)
    {
        printf("\n FILE SEARCH VALUE FOUND !");
        return;
    }
}
```

care necesită pentru același set de date NCMF = 17317 cicluri mașină.

Programatorul trebuie să acorde atenție acestui aspect și trebuie să-și fundamenteze decizia de a implementa o structură de căutare în defavoarea alteia pe imposibilitatea de a stoca volumul mare de date în memoria internă sau pe obiectivul de a maximiza viteza de prelucrare a unui set redus de date.

Dacă este considerată a colectivitate C formată din elementele $c_1, c_2, \dots, c_n$ și identificate prin cuvinte cu valoare de cheie de regăsire având grade de utilizare diferite, algoritmi de regăsire performanți revin la a:

- determina durata operației de regăsire stabilă; această caracteristică a procesului de regăsire este măsurată în funcție de valoarea dispersiei statistice; aceasta este calculată pentru șirul valorilor înregistrate la testarea performanței algoritmului; gradul de stabilitate a operației tinde către un nivel optim pe măsură ce valoarea dispersiei se apropie de nivelul zero; o altă abordare este dată de considerarea a două loturi de procese de regăsire; cele două serii de valori sunt analizate prin prisma mediilor m_1 și m_2 , în condițiile în care acestea nu diferă semnificativ; (se iau 2 utilizatori și se înregistrează tranzacțiile) procesul se extinde la k utilizatori ai algoritmului de regăsire;
- numărul de situații ambigue este controlat; numărul situațiilor în care se produceau ambiguități e redus;
- există o situație în care se determină un volum maxim prelucrări și o situație în care se determină un volum minim de prelucrări.

În cazul algoritmilor secvențiali pentru cheia primului articol este nevoie de o comparare, pentru a doua cheie sunt necesare două comparări, iar pentru ultimul articol n comparări. Numărul total de comparări necesar regăsirii tuturor elementelor este

$$NTC = \frac{n*(n+1)}{2} \quad (29.36)$$

Pentru cele n articole, numărul mediu de comparații este:

$$NMC = \frac{NTC}{n} = \frac{n+1}{2} \quad (29.37)$$

Dacă este dată o structură arborescentă de căutare:

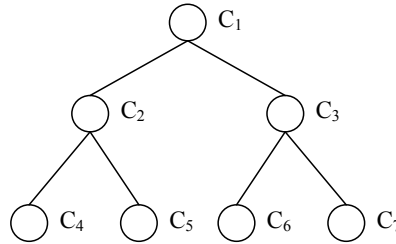


Figura 29.10 Structura arborescentă de căutare

Pentru regăsirea elementelor sunt necesare numărul de căutări:

C_1 – 1 comparare
 C_2 – 2 comparații
 C_3 – 2 comparații
 C_4 – 3 comparații
 C_5 – 3 comparații
 C_6 – 3 comparații
 C_7 – 3 comparații

În această situație, numărul total de comparații necesar identificării fiecărui element din arbore este

$$NTC = 1*1 + 2*2 + 3*4 = \sum_{k=1}^m k * 2^{k-1} \quad (29.38)$$

unde k reprezintă numărul de niveluri din arbore. Numărul mediu de comparații efectuate pentru oricare element este

$$NMC = \frac{\sum_{k=1}^{nv} k * 2^{k-1}}{2^{nv+1} - 1} \quad (29.39)$$

unde nv reprezintă numărul de niveluri din arbore și se consideră că $2^{nv+1} - 1$ reprezintă numărul de noduri dintr-o structură arborescentă binară completă.

În ipoteza că arborele este complet echilibrat, atunci acesta conține pe fiecare dintre cele $i = 1, 2, \dots, nv$ niveluri câte 2^{nv} noduri.

A optimiza o astfel de structură, înseamnă a găsi momentul de reechilibrare a arborelui. Fie V_i volumul de prelucrări necesar echilibrării i a

arborelui în care au fost adăugate noduri, prin adăugarea de frunze cu chei mai mari.

De exemplu, se consideră șirul valorilor 10, 15, 21, 43 și 67 ce formează arborele binar de căutare echilibrat din figura 29.11.

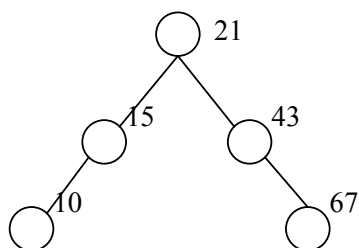


Figura 29.11 Arbore binar de căutare

Dacă se adaugă cheile 75, 101, 119, 400 se obține arborele de căutare:

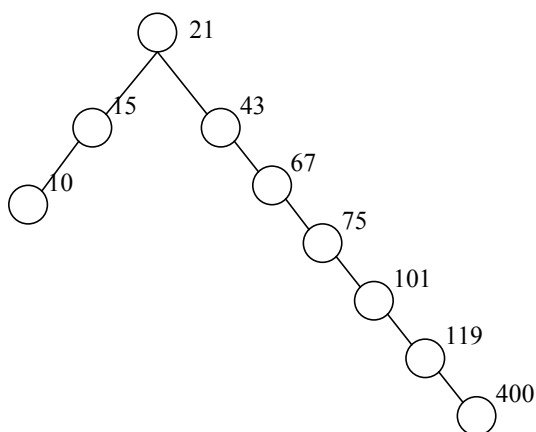


Figura 29.12 Arbore binar de căutare dezechilibrat

Volumul total de comparații pe arborele neechilibrat, din figura 29.12 este $VCA_1 = 33$. De la reorganizarea i a arborelui se trece la reorganizarea $i+1$ dacă și numai dacă

$$VCA_{k1} + VCA_{k2} + \dots + VCA_{kI} > V_{kl} \quad (29.40)$$

Dacă se consideră operația de reechilibrare a arborelui bazată pe o parcurgere în inordine și o reconstrucție a arborelui prin metoda *divide et impera* pentru a insera la fiecare iterație elementul din mijlocul intervalului, atunci se obține arborele echilibrat din figura 29.13.

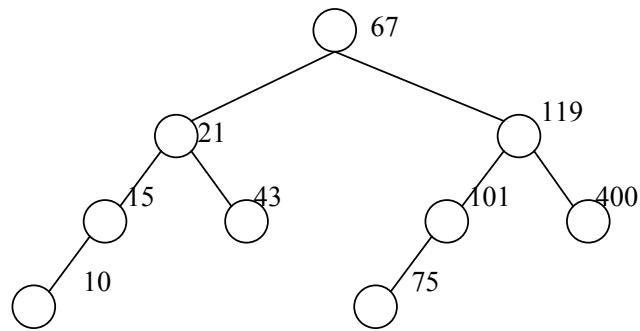


Figura 29.13 Arbore binar de căutare reechilibrat

Luând în considerare că operația de inserare a unui nod nou este echivalentă cu căutarea elementului respectiv, volumul de prelucrări al reechilibrării arborelui ca număr de comparații efectuate este $VRA = 25$. De asemenea, această valoare este egală cu volumul total de comparații pe arborele echilibrat obținut, $VCA_2 = 25$. Se observă, minimizarea volumul total obținut prin rearanjarea structurii, $VCA_2 < VCA_1$.

Atunci când este îndeplinită restricția, K_i reprezintă momentul care precede echilibrarea arborescenței, optimizând procesul de regăsire a informației. Această abordare a optimizării regăsirii în structuri arborescente urmărește nu numai reducerea efortului de căutare, ci și operațiile de gestiune a arborelui, inserarea, respectiv, ștergere de valori. Obiectivul este de a reechilibra atunci când eficiența căutării este afectată și nu după fiecare operație de adăugare sau ștergere noduri.

Căutarea informației sortate reprezintă un mod de optimizare. Se consideră șirul valorilor sortate crescător

$$10, 20, 25, 40, \dots, 200 \quad (29.41)$$

Dacă informația nu e sortată este dificil de stabilit un mod de oprire a căutării dacă s-a depășit o valoare.

În exemplul dat, dacă se caută elementul cu cheia 43 și s-a trecut de cheia 50 din fișier, se conchide că articolul cu cheia 43 nu există.

Căutarea din mai multe direcții oferă o soluție alternativă pentru optimizarea timpului de regăsire. În ciuda faptului că soluția implică o creșterea a nivelului de memorie ocupat, obiectivul minimizării timpului de prelucrare a datelor are prioritate din acest punct de vedere. Soluții practice care implementează această abordare definesc mai multe repere în cadrul colectivității țintă.

Există numeroase aplicații software ce utilizează articole cu cheie alfabetică. În funcție de poziționarea relativă a valorii cheii față de începutul setului sau de sfârșitul acestuia, căutarea are loc în direcții diferite.

Dacă sunt chei care încep cu caracterul 'a', atunci căutarea are loc de la stânga la dreapta. Dacă sunt articole căutate după cheia j ce are valori apropiate caracterului 'z' atunci căutarea se realizează de la sfârșitul seriei de valori.

Este necesar ca toate produsele software să fie înzestrate cu componente care înregistrează informații referitoare la modul în care s-au efectuat prelucrările.

De exemplu, se consideră programul P care lucrează cu un fișier în care se accesează articole după valori cheie, ce sunt sortate crescător.

Accesul la datele aplicației se realizează de la stânga spre dreapta, având ca punct de start începutul fișierului.

Programul contorizează numărul de citiri din fișier și numărul de comparații necesar identificării articolului căutat. Acesta va gestiona informațiile prin realizarea unei structuri asemenea celei descrise în tabelul 29.8.

Tabelul nr. 29.8 Analiza procesului de căutare și citire din fișier

Valoare cheie	Număr comparații	Număr citiri
k_1	a_1	cit_1
k_2	a_2	cit_2
...	...	...
k_i	a_i	cit_i
...	...	...
k_n	a_n	cit_n

În urma analizei datelor din tabel, rezultă $\bar{\alpha}$ numărul mediu de comparații și o dispersie σ^2 ce descrie omogenitatea înregistrărilor. Dacă există T valori vor exista $\bar{\alpha}_1, \bar{\alpha}_2, \dots, \bar{\alpha}_T$ și $\sigma_1^2, \sigma_2^2, \dots, \sigma_T^2$ indicatori.

Prin teste statistice se verifică egalitatea mediilor și egalitatea dispersiei. În ipoteza de egalitate, procesul este stabil și face obiectul procesului de optimizare. Se împarte fișierul în două zone de căutare, de la articolul art_1 la articolul $art_{\left[\frac{n}{2}\right]}$ și de la acest punct până la art_n . În această

situație, dacă valoarea căutată are o valoare mai mică decât $A = art_{\left[\frac{n}{2}\right]}$, căutarea are loc în prima jumătate. Dacă valoarea cheii este mai mare ca $A = art_{\left[\frac{n}{2}\right]}$, atunci căutarea este aplicată în a doua jumătate.

Pentru aceleași set de chei de căutare, $k_1, k_2, \dots, k_n$ se repetă procesul, rezultând șirul numărului de căutări efectuate $\beta_1, \beta_2, \dots, \beta_n$. Pentru alte R loturi se obțin indicatorii $\bar{\beta}_1, \bar{\beta}_2, \dots, \bar{\beta}_R$ și $\sigma_1^2, \sigma_2^2, \dots, \sigma_R^2$.

Dacă mediile și dispersiile sunt egale revine că procesul este stabil.

Se compară mediile și dispersiile rezultate în urma celor două teste pentru a identifica care dintre cele două procedee este mai bun.

Odată selectată abordarea care conduce la cele mai bune rezultate, se împarte fișierul în trei și procedeul continuă.

Se va găsi nivelul optim de împărțire a fișierului pentru ca numărul de comparații să fie, statistic, cel mai mic.

Căutarea este un proces ce implică mai mult decât operația de comparare. Când aceasta este precedată de citire, este important ca numărul de citiri să fie cât mai redus.

Regăsirea informației presupune asocierea unor chei textului. Se consideră o mulțime de fișiere $F_1, F_2, \dots, F_n$ și se pune problema la căutare:

- găsirea fișierului F_i ;
- găsirea în fișierul F_i a anumitor informații.

Se ia fiecare fișier $F_1, F_2, \dots, F_n$ și se construiește o structură informațională privind conținutul acestuia, în care sunt reținute:

C_1 – termeni generali;
 C_2 – termeni specifici;
 C_3 – grad de cuprindere mare;
 C_4 – grad de cuprindere redus, particular conținutului și domeniului de apartenență.

Pentru fiecare fișier se construiește vocabularul și se iau termenii de bază după frecvențe. Se asociază coeficienții de complexitate ($C_1 F_1$), ($C_2 F_2$), ..., ($C_i F_i$). Fișierele sunt împărțite în patru categorii de complexitate C_1 , C_2 , C_3 și C_4 . Se studiază elementele de vocabular și se identifică tipul de complexitate pentru fiecare. Se ierarhizează cuvintele și vor rezulta la motorul de căutare că pentru cuvântul cv_i sunt informații în fișiere de complexitate cuprinse între $[\alpha_i, \beta_i]$. Din mulțimea de fișiere se extrag numai acelea care au această complexitate după care se merge pe vocabular și se iau fișierele unde cv_i au frecvența cea mai mare de apariție.

Toate acestea presupun analiza comportamentului statistic pentru entități text, [Ivap05].

29.4 Mecanisme flexibile de regăsire

Spre deosebire de aplicațiile clasice, aplicațiile distribuite care se adresează unui număr foarte mare de clienți trebuie să fie caracterizate prin nivel de flexibilitate foarte ridicat.

În primul rând flexibilitatea se manifestă prin diversitatea modurilor de introducere a datelor. Datele se introduc de la tastatură, prin coduri bară, prin scanare sau de pe un suport extern. Sunt aplicații care preiau comenzi vocale.

În al doilea rând flexibilitatea se manifestă prin structura datelor de intrare. Sunt acceptate atât date complete, cât mai ales date incomplete, care permit totuși identificarea corectă a elementului căutat.

Expresiile care combină atribute se construiesc astfel încât să fie accesat din aproape în aproape acel conglomerat informațional care corespunde evenimentului, tranzacției produsului sau serviciului de care este interesat clientul.

Pentru a afla nota obținută la examen, studentul folosește adresa de Internet a universității la care acesta studiază. Primul meniu va fi:

Student
Profesor
TESA

Studentul va alege opțiunea STUDENT.

Meniul care se activează este referitor la facultate:

Facultăți:
Matematică
Chimie
Management
Electronică
Anul: 1 2 3 4 5

Studentul selectează MATEMATICĂ și cifra 4.
Meniul următor se referă la disciplinele din anul și facultatea selectată:

Disciplina
Geometrie
Analiză
Ecuații diferențiale
Algebra
Astronomie

Este selectată opțiunea ALGEBRĂ
Meniul următor vizează introducerea numelui, prenumelui parolei

Nume:
Prenume:
Parolă:

În cazul în care datele introduse sunt corecte, se afișează meniul:

Studentul - - - - -
A obținut la examenul de la
disciplina
- - - - -
nota: - - - - -

corectare ieșire

Selectarea opțiunii continuare permite apariția meniului disciplinelor și pot fi selectate 1, un grup sau toate și vor fi afișate note.

Se observă că accesarea presupune selecție. Datele de intrare sunt numele prenumele și parola care trebuie ortografiate strict după reguli convenite la înscrierea în facultate privind literele mari, literele mici și diacriticele.

În cazul achiziționării unui produs, clientul accesează un magazin virtual. Meniul de legătură este:

Magazinul x unde
Produse

- Cash
- În rate

continuare

Clientul alege opțiunea convenabilă. Meniul următor se referă la tipuri de produse:

Produs

- frigidere
- telefoane
- mașini de spălat
- autoturisme

Selectează televizoare:

Diagonala

40 cm
50 cm
60 cm
80 cm
100 cm

Selectează 60 cm:

Existent în stoc

Panasonic	8 500.000
LG	6 300.000
Sony	12 000.000
Siemens	11 000.000

Rate cash

Se tastează rate:

Venit net: 4 000.000
Sumă maximă 800.000
Rată:
Afișez număr rate:
Lista valorii nete:
Data plății în ziua:
Valori nete +dobândă

Problema căutării este definită printr-o mulțime de elemente $a_1, a_2, \dots, a_n$ înzestrate cu diferite proprietăți, a căror dispunere este de asemenea diferită. Trebuie localizat un element din mulțime în vederea utilizării. Localizarea este efectuată prin procese de căutare și identificare elementul găsit este cel căutat, are loc prelucrarea.

Problema căutării are numeroase aplicații și mai ales dezvoltări teoretice în domeniul valutar întrucât apar multe necunoscute în privința numărului elementelor $a_1, a_2, \dots, a_n$, poziției și proprietăților acestora.

Asemenea căutării în aplicațiile valutare, implicațiile Internet reprezintă numeroase necunoscute pentru clienți. Fiecare aplicație are specificul ei.

Modul în care este definită interfața are rolul de a determina sistemul selecțiilor de căutare a revoluționat întregul sistem al căutării.

Este important ca întregul sistem al căutării să se bucure de caracteristica de continuitate.

În primul rând se impune studiul interfeței aplicațiilor cele mai frecvent referite.

În al doilea rând, se utilizează caracteristicile de căutare deja cunoscute.

În al treilea rând se procedează la înzestrarea interfeței cu caracteristici de reorientare, care va da din aproape în aproape o caracteristică apropiată de cerințele materiale. Se înregistrează ferestrele de referire a opțiunilor după care se reorganizează structura mesajelor.

În al patrulea rând, se evaluează opțiunile în așa fel încât situația în care se tastează date se reduc cât mai mult posibil.

În al cincilea rând se are în vedere găsirea tuturor posibilităților de introducere date și de realizare de operații în așa fel încât să se asigure completitudinea operațiilor (mod de plată de prelucrare informații).

Căutarea clienților necesită algoritmi flexibili și un nou mod de abordare.

Se subordonează criteriile de definire chei de selecție și mai ales opțiuni care țin seama de cerințele clienților și nu de restricțiile echipelor care dezvoltă aplicațiile.