

17. TABELE DE DISPERSIE

17.1 Structura de date de tip tabela de dispersie

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 este 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 ValoareVector(element* vector, int n, int cheie)
{
    if(cheie >= n) return -1;
    else if(vector [cheie].cheie == -1) return -1;
    else return vector [cheie].valoare;
}
```

Cu toate acestea, în practică este evitată această soluție directă care, deși minimizează timpul de regăsire prezintă numeroase dezavantaje generate de:

- dimensiunea memoriei ocupate – spațiul de memorie, MEMORIE, necesar implementării acestei structuri este dat de relația

$$MEMORIE = \text{maxim(valoare cheie căutare)} * \text{dimensiune(element)} \quad (17.1)$$

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

$$MEMORIE = \max(nrMatricol) * dimensiune(Student) = 55630 * 48 = 2,54 \text{ MB} \quad (17.2)$$

iar gradul de utilizare a acestui spațiu este egal cu:

$$GUM = \frac{MEM_{elemente}}{MEM} * 100 = \frac{1450 * 48}{55630 * 48} * 100 = 2,6\% \quad (17.3)$$

- 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 – în cazul aplicației de gestiune a studenților se observă că din totalul de 55630 de elemente, doar 1450 sunt utilizate; faptul se datorează concentrării valorilor cheii de căutare în intervalul [54130, 55630]; reducerea spațiului ocupat de tabela de dispersie se realizează în această situație prin implementarea unei funcții *hash*: [54130, 55630] → [0, 1500] care să reducă valoarea maximă a cheii de căutare prin conversia valorii din intervalul [54130, 55630] într-o valoare din intervalul [0, 1500]; o astfel de funcție *hash* cu un nivel de complexitate scăzut este

$$hash(X) = X \text{ modulo } 1500, \text{ cu } X \in [54130, 55630] \quad (17.4)$$

- implementarea de chei alfanumerice – prin utilizarea unei funcții care să prelucreză 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 transla valoarea alfanumerică într-o valoare întreagă pozitivă; de exemplu, în limbajul Java este implementată funcția *hash* pentru un string

$$\text{hash}(S) = s[0]*31^{n-1} + s[1]*31^{n-2} + \dots + s[n-2]*31 + s[n-1] \quad (17.5)$$

unde:

- $s[i]$ – codul ASCII al caracterului i din șir;
- n – dimensiunea șirului de caractere.

Pentru cheia alfanumerică cu valoarea *salut* aplicarea funcției *hash* conduce la obținerea valorii

$$\text{hash}(\text{"salut"}) = 115*31^4 + 97*31^3 + 108*31^2 + 117*31 + 116 = 109202173 \quad (17.6)$$

valoarea obținută este introdusă din nou într-o funcție *hash* care să identifice poziția corespundătoare din mulțime pentru această valoare numerică foarte mare.

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 tablei 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; printre cele mai utilizate tehnici de evitare a coliziunilor se numără *chaining*, *re-hashing*, *linear probing*, *quadratic probing* și *overflow area*, [Will96], [Seng94].

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 17.1.

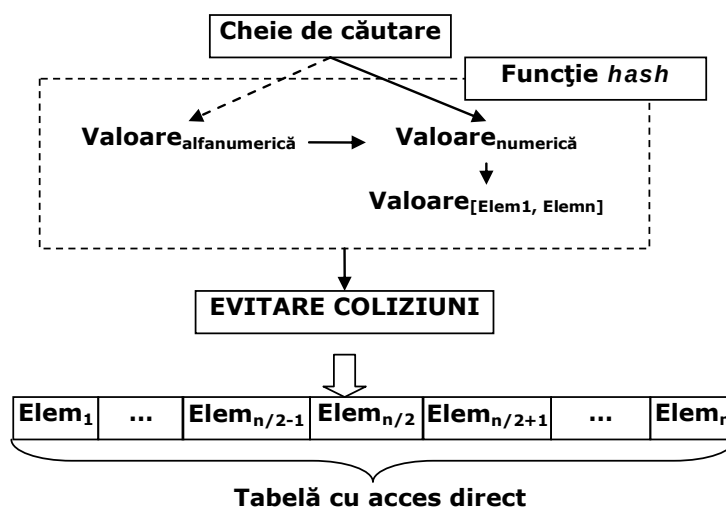


Figura 17.1 Tabela de dispersie

17.2 Chei si funcții *hash*

Pentru a identifica în mod unic fiecare înregistrare sunt utilizate valori dintr-un domeniu bine definit. Asocierea dintre aceste valori unice și înregistrările pe care le reprezintă este de unu la unu. În funcție de tipul valorii cu rol de cheie se identifică:

- chei numerice reprezentate prin intermediul tipurilor fundamentale definite de limbajul de programare utilizat;
- chei alfanumerice reprezentate prin șiruri de caractere;
- chei compuse definite pe baza a mai multor atribute;

Rolul funcției *hash* este de a prelucra cheia asociată fiecărei înregistrări și de a determina poziția în cadrul tabelii de dispersie a elementului respectiv. Deoarece nu există o funcție *hash* generală, iar alegerea acesteia se face în funcție de caracteristicile mulțimii de valori chei, sunt utilizate modele matematice bazate pe:

- împărțire în modul; această metodă are un grad ridicat de utilizare în diferitele modele de implementare a tabelilor de dispersie prin prisma complexității scăzute a modelului și a ușurinței în implementare; indiferent de forma cheii de căutare, aceasta este transformată într-o valoare numerică și apoi transpusă în mulțimea $[0...nht-1]$, unde *nht* reprezintă dimensiunea tabelii de dispersie; acest tip de funcție are asociat modelul:

$$pozitie_tabela = val_cheie \% val_baza \quad (17.7)$$

unde:

- *pozitie_tabela* – valoarea *hash* obținută;
- *val_cheie* – valoarea numerică ce identifică unic fiecare înregistrare;
- *val_baza* – valoarea numerică stabilită la definirea modelului;

Valoarea utilizată ca bază în modelul funcției *hash* este dată de dimensiunea tabelii de dispersie; obiectivul definirii funcției *hash* este de a obține un model care să minimizeze mulțimea de chei diferite ce conduc la aceeași valoare *pozitie_tabela*; în acest sens se utilizează numere prime apropiate de numărul total de înregistrări; astfel, dacă cheile înregistrărilor definesc o mulțime continuă de valori unice, valorile *hash* determinate vor acoperi toată mulțimea $[0..val_baza-1]$; pentru a stabili dimensiunea inițială a tabelii de dispersii sunt alese numere prime particulare; de exemplu limbajul Java definește o clasă *HashTable* ce descrie o tabelă de dispersie cu dimensiunea inițială egală cu 101; valoarea nu este aleasă la întâmplare deoarece permite redimensionarea tabelii la o nouă valoare primă; pentru cazuri generale, este indicat să se folosească un număr prim determinat prin relația:

$$dimensiune_hash = (4*i+3) \text{ cu } i=0, 1, 2, 3,... \quad (17.8)$$

- înmulțirea cu un număr real aleatoriu și prelucrarea ulterioară a părții zecimale – valoarea cheii de căutare este înmulțită cu un număr din intervalul $[0;1)$ ce este ales aleatoriu; în urma

Înmulțirii se obține o valoare temporară, *temp*, a cărei parte zecimală are valori cuprinse în intervalul $[0;1)$; înmulțind *temp* cu dimensiunea tabelului de dispersie, *dimensiune_hash*, se obține o valoare în intervalul $[0, dimensiune_hash-1]$ ce reprezintă poziția elementului căutat în tabelă; relația 17.9 descrie procesul de determinare a valorii *hash*:

$$val_hash = ((val_cheie * random_{[0;1)}) - [(val_cheie * random_{[0;1)})]) * nth \quad (17.9)$$

unde:

- *val_hash* – valoarea *hash* calculată;
- *val_cheie* – valoarea cheii de căutare;
- *random_{[0;1)}* – număr aleatoriu din mulțimea $[0;1)$;
- *nth* – dimensiunea tabelului de dispersie.

- utilizarea unui bloc de biți obținut prin prelucrarea cheii de căutare – de exemplu se definește funcția care multiplică valoarea codului și care determină valoarea *hash* pe baza celor 10 biți din interiorul acesteia; dacă se consideră cheile reprezentate pe 32 de biți atunci relația este:

$$val_hash = ((val_chei * val_cheie) >> 11) \% nhash \quad (17.10)$$

- prelucrarea codurilor ASCII ale caracterelor alfanumerice – în cazul în care cheia de căutare este dată de o secvență de caractere atunci este necesar ca aceasta să fie prelucrată pentru a se genera un număr întreg care să indice poziția elementului în tabelă; prelucrarea se face pe baza codului ASCII asociat caracterului; pe baza primului caracter din cheie se definește relația:

$$val_hashs_1 = string_cheie[0] \% 255 \quad (17.11)$$

unde:

- *val_hashs₁* – valoarea *hash* calculată;
- *string_cheie* – valoarea cheii de căutare;

Relația anterioară definește un model cu un nivel de complexitate scăzut ce este destinat gestiunii unei colectivități mici de elemente; în cazul în care valoarea cheii reprezintă numele unei persoane, modelul este ineficient deoarece generează multe coliziuni pentru nume diferite dar care încep cu același caracter; rafinarea acestui model se poate face prin preluarea mai multor caractere din șirul respectiv; de exemplu se definește o nouă relație

$$val_hashs_2 = (string_cheie[0] + string_cheie[lungime_string_cheie]) \% dim_has \quad (17.12)$$

unde:

- *val_hashs₂* – valoarea *hash* calculată;
- *string_cheie* – valoarea cheii de căutare;
- *lungime_{string_cheie}* – dimensiunea șirului de caractere;

- *dim_hash* – dimensiunea tabelii de dispersie;
ce ia în calcul primul și ultimul caracter; de asemenea, pentru a nu reduce dimensiunea tabelii la maxim 255 elemente, se utilizează un număr prim, *dim_hash* suficient de mare; alte funcții *hash* de prelucrare a cheilor alfanumerice analizează toate caracterele din șir; de exemplu relația

$$val_hashs_3 = \sum_{i=1}^{lungime_cheie} ASCII(string_cheie[i]) \% dim_hash \quad (17.13)$$

unde:

- *val_hashs₃* – valoarea *hash* calculată;
 - *string_cheie* – valoarea cheii de căutare;
 - *lungime_cheie* – dimensiunea șirului de caractere;
 - *dim_hash* – dimensiunea tabelii de dispersie;
- determină valoarea *hash* pe baza sumei codurilor ASCII asociate tuturor caracterelor din cheie;

- înmulțirea cu un număr real definit; această valoare este determinată prin evaluarea expresiei $m = \frac{(\sqrt{5}-1)}{2}$ sau a expresiei

$$m = 1 - \frac{(\sqrt{5}-1)}{2}; \text{ funcția } hash \text{ definită pe baza acestui model este}$$

$$f(val_cheie) = [val_cheie * m] \quad (17.14)$$

acest model este implementat în cadrul funcțiilor de *rehashing* ce sunt utilizate în cadrul procedurilor de rezolvare a coliziunilor.

Dificultatea în definirea funcțiilor *hash* constă în identificare acelui model care să minimizeze numărul coliziunilor și care să genereze o distribuție uniformă a valorilor pe întreg intervalul.

17.3 Evitarea coliziunilor

Funcțiile de evitare a coliziunilor 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:

- *chaining* implementează lucrul cu liste – fiecare poziției din cadrul tabelii de dispersie conține capătul unei liste de elemente cu valori *hash* egale; regăsirea unui element presupune determinarea poziției în cadrul tabelii prin calcularea valorii *hash* și parcurgerea secvențială a listei atașate poziției respective după valoarea cheii de căutare, figura 17.2;

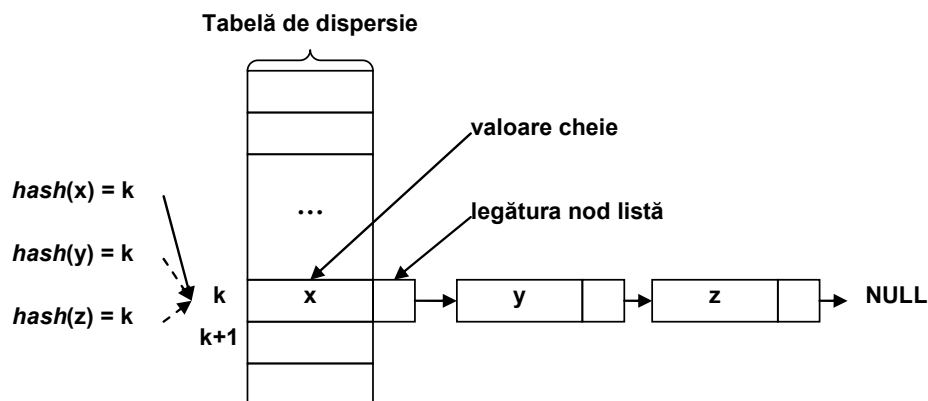


Figura 17.2 Metoda chaining de evitare a coliziunilor în tabele de dispersie.

- *re-hashing* presupune aplicarea în cascadă a aceleiași funcții *hash* sau a altui model dintr-o mulțime de funcții până când valoarea obținută reprezintă o poziție liberă din cadrul tabelului de dispersie; la fiecare pas al procesului de regăsire, valoarea cheii de căutare este introdusă într-o listă de funcții *hash* până când se identifică elementul cu valoarea căutată sau nu mai există alte posibilități de a recalcula valoarea *hash*; tipul și numărul de funcții *hash* aplicate valorilor intermediare descriu o procedură bine definită de căutare, respectiv inițializare element nou, pentru a conduce de fiecare dată la aceleași rezultate, figura 17.3;

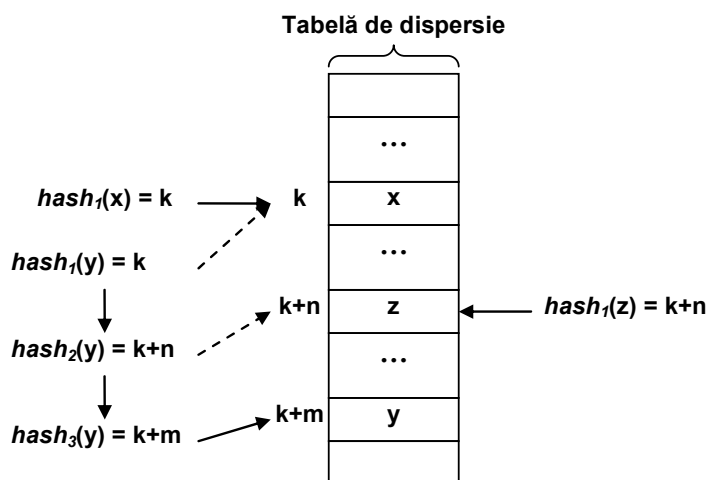


Figura 17.3 Metoda rehashing de evitare a coliziunilor în tabele de dispersie

- *linear probing* se bazează pe căutarea secvențială a primei poziții libere în care să fie inserat elementul nou, poziție aflată la stânga sau la dreapta coliziunii; în cazul căutării, procesul presupune verificarea elementelor adiacente poziției indicate de valoarea *hash*; cu toate că are un grad scăzut de complexitate, această metodă de evitare a coliziunilor are ca efecte secundare gruparea coliziunilor de același tip în aceeași zonă, *cluster*, fapt care conduce la creșterea probabilității de apariție a coliziunilor pentru valorile *hash* adiacente, figura 17.4;

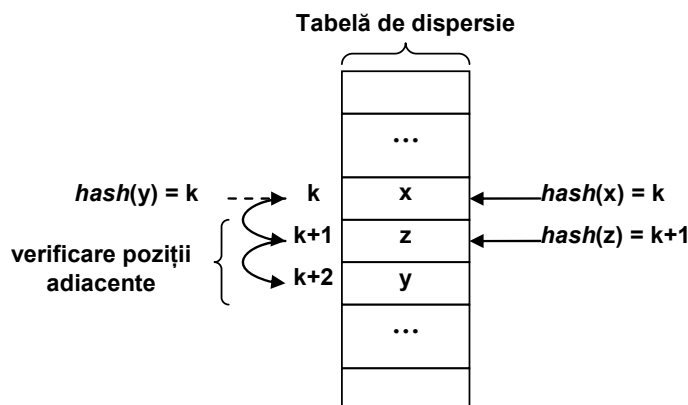


Figura 17.4 Metoda linear probing de evitare a coliziunilor în tabele de dispersie

- *quadratic probing* este o metodă de tipul *linear probing* care evită crearea grupurilor de coliziuni prin utilizarea unui pas de regăsire a următoarei poziții libere diferit de 1; astfel, în caz de coliziuni au loc salturi în tabela de dispersie din două în două poziții sau din patru în patru; în general, pentru a determina următoare poziție în care să se insereze un element nou sau în care să se caute un element existent este dat de funcția:

$$\text{poziție} = \text{hash}(X) + c * i^2 \quad (17.15)$$

unde:

- *poziție* – noua poziție din tabela de dispersie în care se inserează sau se caută un element;
- X – valoarea cheii asociate elementului;
- $\text{hash}(X)$ – poziția indicată de valoarea *hash* a elementului care se adaugă sau se caută;
- c – valoare constantă definită în mulțimea $\{1, 2, 4\}$;
- i – numărul operației de *rehash* sau numărul de poziții verificate.

Prin utilizarea de repositionări pe poziții neadiacente se evită gruparea coliziunilor în aceeași zonă din tabela de dispersie, figura 17.5;

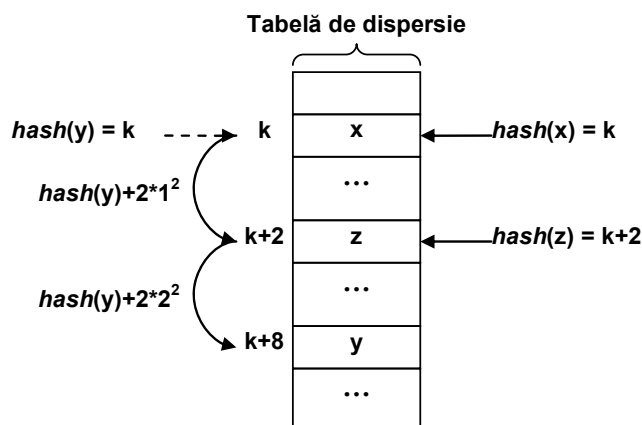


Figura 17.5 Metoda quadratic probing de evitare a coliziunilor în tabele de dispersie

- *overflow area* presupune o abordare ce împarte tabela de dispersie în două zone, primară pentru reținerea elementelor inițiale și secundară alocată elementelor ce generează coliziuni; la apariția unei coliziuni, fie la operația de regăsire sau la adăugarea unui element nou, se utilizează un element al zonei secundare pentru a reține noua valoare sau pentru a continua căutarea; accesul la zona secundară se realizează doar prin intermediul unui pointer din zona primară, funcția hash negenerând poziții în acest interval de valori; această metodă, descrisă în figura 17.6, permite o regăsire mai rapidă a informațiilor decât metoda *chaining* care presupune parcurgerea de liste;

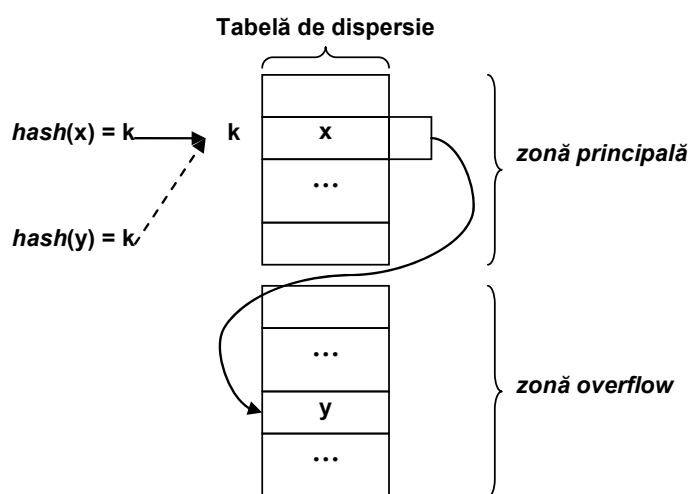


Figura 17.6 Metoda overflow de evitare a coliziunilor în tabele de dispersie

Se observă în cazul tabelelor de dispersie că probabilitatea de apariție a coliziunilor la inserare sau la căutare crește proporțional cu gradul de utilizare a tabelului. În cele mai multe cazuri, funcțiile *hash* cu un grad redus de complexitate nu conduc la rezultate unice pentru valori de intrare distincte. Din acest motiv, cu cât tabela pune la dispoziție un număr din ce în ce mai mic de poziții disponibile, crește riscul de a avea numeroase elemente cu chei de căutare diferite care ar trebui să se regăsească pe poziții identice. Metodele de evitare a coliziunilor rezolvă parțial problema deoarece ele plasează elementele noi pe alte poziții definind o nouă cauză pentru coliziuni viitoare.

Pentru a rezolva această situație și pentru a menține eficiența operației de căutare la un nivel acceptabil, astfel încât să nu se parcurgă toată mulțimea de elemente la fiecare operație de regăsire de informație, se analizează gradul de utilizare a tabelului *hash*, GU_{hash} .

$$GU_{hash} = \frac{NPO}{NTP} * 100 \quad (17.16)$$

unde:

- *NPO* – numărul de poziții ocupate;
- *NTP* – numărul total de poziții.

care este menținut la o valoare mai mică de 50%. În ciuda faptului că are loc o gestiune ineficientă a spațiului, eficiența abordării se bazează pe o viteză de regăsire mare. Pentru a menține valoarea indicatorului GU_{hash} sub limită impusă, tabela se redimensionează prin dublarea numărului de poziții. În cazul în care modelul funcției hash este bazat pe împărțirea în modul la dimensiunea acesteia, redimensionarea trebuie să conducă la o nouă dimensiune care să reprezinte tot un număr prim. Având în vedere importanța acestor tipuri de numere, lucru evidențiat la descrierea tipurilor de funcții hash, redimensionarea se face prin intermediul relației

$$dimensiune_noua = dimensiune_curenta * 2 + 1 \quad (17.17)$$

În cazul limbajului Java, unde dimensiunea inițială a tabelului de dispersie este 101, redimensionarea ulterioară a acesteia utilizând relația descrisă conduce la o serie de valori care reprezintă tot valori prime. De exemplu 203 și 407. În cazul în care relația utilizată pentru a determina numărul prim este

$$dimensiune_hash = (4*i+3) \text{ cu } i=0, 1, 2, 3, \dots \quad (17.18)$$

redimensionarea se realizează prin înlocuirea lui i cu $2*i$. Astfel, noua dimensiune va fi reprezentată tot de un număr prim.

17.4 Implementare

În continuare se exemplifică printr-un program scris în limbajul de programare C++ implementarea unei tabeli de dispersie pentru înregistrări de tip *Student*, căutările realizându-se după numărul matricolei. Ca funcție de dispersie a fost ales un exemplu simplu cu ajutorul operatorului modulo (%). Dimensiunea implicită a tabelului a fost aleasă 101, o valoare cuprinzătoare pentru exemplul ales.

Tabela de dispersie este implementată sub forma unui masiv unidimensional alocat dinamic, ale cărui elemente sunt liste care conțin datele propriu zise. Acest mod de implementare permite tratarea coliziunilor, în cazul în care pentru indexul generat există o înregistrare în tabelă, noua înregistrare va fi adăugată ca element la lista simplu înlănțuită de la nodul respectiv.

```
#include "stdafx.h"
#include <iostream>

using namespace std;

//structura de date stocata in tabela de dispersie
//pentru simplificarea codului s-a utilizat o structura
struct Student
{
    int matricula;
    char nume[20];
    char adresa[30];
    char localitate[20];
    char e_mail[20];
    char telefon[15];
};
```

```

        void Afisare()
        {
            cout<<"Matricola: "<< matricola <<endl;
            cout<<"Matricola: "<< nume <<endl;
            cout<<"Matricola: "<< adresa<<endl;
            cout<<"Matricola: "<< localitate<<endl;
            cout<<"Matricola: "<< e_mail <<endl;
            cout<<"Matricola: "<< telefon <<endl;
        }
    };

//elementele tabelii de dispersie
//pentru simplificarea codului s-a utilizat o structura
struct NodDate
{
    NodDate *urm;
    Student data;
};

//Cheia este matricola studentului
//
class HashStudent
{
    int size;
    NodDate ** elem;

    //functie simpla de dispersie care utilizeaza operatorul modulo
    //pentru a genera pozitia in tabela
    int HashFunction(int cheie)
    {
        return cheie%size;
    };

    //aloca memorie pentru structura si initializeaza cu null
    void AlocaMemorie()
    {
        elem = new NodDate*[size];
        for (int i=0;i<size;i++)
            elem[i] = NULL;
    }
    void ElibereazaMemorieLista(NodDate *&);
public:
    HashStudent():size(101)
    {
        AlocaMemorie();
    }
    ~HashStudent();
    HashStudent(int iSize): size (iSize)
    {
        AlocaMemorie();
    }
    int Insereaza(Student data);
    Student Cauta(int matricola);
    int Sterge(int matricola);
    void HashStudent::Afiseaza();
};

//insereaza un element pe o pozitie generata de functia de dispersi

```

```

//pe baza valorii matricolei din structura
int HashStudent::Insereaza(Student s)
{
    int index = -1;

    if (s.matricola<0)
        return index;

    if (elem!=NULL)
    {
        index = HashFunction(s.matricola);
        //prima inregistrare
        NodDate *nod_nou = new NodDate;;
        nod_nou->urm = NULL;
        nod_nou->data = s;

        if (elem[index]== NULL)
        {
            elem[index] = nod_nou;
        }
        else //coliziune
        {
            NodDate *t = elem[index];
            while(t->urm!=NULL)
                t = t->urm;
            t->urm = nod_nou;
        }
    }
    return index;
}

//cauta un element dupa chiea matricola
//in cazul in care nu este gasita nici o inregistrare este
//returnata o structura Student cu matricola -1
Student HashStudent::Cauta(int matricola)
{
    //valoare de retur pentru cazurile in care nu este gasit
    elementul
    Student s_negasit;
    s_negasit.matricola = -1;

    if (matricola<0)
        return s_negasit;

    if (elem!=NULL)
    {
        int index = HashFunction(matricola);
        if (elem[index]==NULL)
        {
            return s_negasit;
        }
        else
        {
            if (elem[index]->data.matricola ==matricola)
                return elem[index]->data;
            else//coliziune
            {
                NodDate *t = elem[index];
                while (t!= NULL && t->data.matricola!=matricola)
                    t = t->urm;
            }
        }
    }
}

```

```

        if (t==NULL)
            return s_negasit;
        else
            return t->data;
    }
}
//neinitializat sau nu exista inregistrarea
return s_negasit;
}

//functie recursiva care elibereaza memoria ocupata de nodurile listei
void HashStudent::ElibereazaMemorieLista(NodDate *&inceput)
{
    if (inceput == NULL)
        return;
    else
    {
        ElibereazaMemorieLista(inceput->urm);
        cout<<"Eliminat nodul cu codul: " <<inceput->data.matricola<<endl;
        delete inceput;
    }
}

//elibereaza memoria ocupa de tabela de dispesie
HashStudent::~HashStudent()
{
    if (elem != NULL)
    {
        for (int i=0;i<size;i++)
            ElibereazaMemorieLista(elem[i]);
        delete [] elem;
    }
}

//Sterge elementul a carui matricola coincide cu matricola trimisa ca parametru
//returneaza pozitia din tabela de a fost eliminat elementul
//sau -1 in cazul in care nu exista elementul cu cheia data
int HashStudent::Sterge(int matricola)
{
    if (matricola<0)
        return matricola;

    if (elem!=NULL)
    {
        int index = HashFunction(matricola);
        if (elem[index]==NULL)
        {
            return -1;
        }
        else
        {
            //este primul
            if (elem[index]->data.matricola ==matricola)
            {
                if(elem[index]->urm == NULL)
                {
                    delete elem[index];
                    elem[index]=NULL;
                }
            }
        }
    }
}

```

```

        }
        else//mai sint si alte elemente in lista
        {
            NodDate *t = elem[index];
            elem[index] = t->urm;
            delete t;
        }
    }
    else//coliziune, nu este primul, il cautam
    {
        NodDate *t = elem[index];
        while (t->urm!= NULL && t->urm->data.matricola
!=matricola)
        {
            t = t->urm;
            if (t->urm==NULL)
                return -1;
            else//t->urm.data.matricola == matricola
            {
                NodDate *p = t->urm;
                if (p->urm ==NULL)
                {
                    t->urm = NULL;
                    delete p;
                }
                else
                {
                    t->urm = p->urm;
                    delete p;
                }
            }
        }
    }
    return index;
}

return -1;
}

}

//Afiseaza continutul tabelii de dispersie
void HashStudent::Afiseaza()
{
    if (elem!=NULL)
    {
        for (int i=0;i<size;i++)
        {
            if (elem[i]!= NULL)
            {
                NodDate *t = elem[i];
                while (t!=NULL)
                {
                    cout<<"S-a inserat in pozitia: " << i <<
" cheia: "<<t->data.matricola<<endl;
                    t = t->urm;
                }
            }
        }
    }
}

```

```

}

int _tmain(int argc, _TCHAR* argv[])
{
    Student grupa[] ={
        {1000,"Ion Vlad", "Str. Norilor
20","Cluj","ion@server.ro", "0101010101"},
        {204,"Mihai Vlad", "Str. Norilor
20","Cluj","mihai@server.ro", "0101010101"},
        {406,"Dan Vlad", "Str. Norilor 20","Cluj","dan@server.ro",
"0101010101"},
        {305,"Lili Vlad", "Str. Norilor
20","Cluj","lili@server.ro", "0101010101"},
        {1022,"Silviu Vlad", "Str. Norilor
20","Cluj","silviu@server.ro", "0101010101"},
        {1021,"Alina Vlad", "Str. Norilor
20","Cluj","alina@server.ro", "0101010101"},
        {1030,"Sorin Vlad", "Str. Norilor
20","Cluj","sorin@server.ro", "0101010101"},
        {1032,"Titi Vlad", "Str. Norilor
20","Cluj","titi@server.ro", "0101010101"},
        {1200,"Gigel Vlad", "Str. Norilor
20","Cluj","gigel@server.ro", "0101010101"},
        {2021,"Anca Vlad", "Str. Norilor
20","Cluj","ioanca@server.ro", "0101010101"},
        {1230,"Maria Vlad", "Str. Norilor
20","Cluj","maria@server.ro", "0101010101"},
        {1008,"Elena Vlad", "Str. Norilor
20","Cluj","elena@server.ro", "0101010101"}
    };
    HashStudent hs;

    for (int i=0;i<sizeof(grupa)/sizeof(Student);i++)
    {
        hs.Insereaza(grupa[i]);
    }
    hs.Afiseaza();
    int m;
    cout<<"Matricola: ";
    cin>>m;

    while (m!=-1)
    {

        Student rez = hs.Cauta(m);
        if (rez.matricola != -1)
            rez.Afisare();
        cout<<"Matricola: ";
        cin>>m;

    }
    hs.Sterge(1008);
    hs.Sterge(406);
    hs.Sterge(305);
    hs.Afiseaza();
    cin.get();
    return 0;
}

```

După rulare, programul produce următoarele rezultate la dispozitivul standard de ieșire:

```

S-a inserat in pozitia: 1 cheia: 2021
S-a inserat in pozitia: 2 cheia: 204
S-a inserat in pozitia: 2 cheia: 406
S-a inserat in pozitia: 2 cheia: 305
S-a inserat in pozitia: 11 cheia: 1021
S-a inserat in pozitia: 12 cheia: 1022
S-a inserat in pozitia: 18 cheia: 1230
S-a inserat in pozitia: 20 cheia: 1030
S-a inserat in pozitia: 22 cheia: 1032
S-a inserat in pozitia: 89 cheia: 1200
S-a inserat in pozitia: 91 cheia: 1000
S-a inserat in pozitia: 99 cheia: 1008
Matricula: 10
Matricula: 1200
Matricula: 1200
Matricula: Gigel Ulad
Matricula: Str. Norilor 20
Matricula: Cluj
Matricula: gigel@server.ro
Matricula: 01010101
Matricula: 1008
Matricula: 1008
Matricula: Elena Ulad
Matricula: Str. Norilor 20
Matricula: Cluj
Matricula: elena@server.ro
Matricula: 01010101
Matricula: -1
S-a inserat in pozitia: 1 cheia: 2021
S-a inserat in pozitia: 2 cheia: 204
S-a inserat in pozitia: 11 cheia: 1021
S-a inserat in pozitia: 12 cheia: 1022
S-a inserat in pozitia: 18 cheia: 1230
S-a inserat in pozitia: 20 cheia: 1030
S-a inserat in pozitia: 22 cheia: 1032
S-a inserat in pozitia: 89 cheia: 1200
S-a inserat in pozitia: 91 cheia: 1000
Eliminat nodul cu codul: 2021
Eliminat nodul cu codul: 204
Eliminat nodul cu codul: 1021
Eliminat nodul cu codul: 1022
Eliminat nodul cu codul: 1230
Eliminat nodul cu codul: 1030
Eliminat nodul cu codul: 1032
Eliminat nodul cu codul: 1200
Eliminat nodul cu codul: 1000

```

Figura 17.7 Rezultatele programului HashStudent