

23. COMPACTAREA DATELOR

23.1 Parametrii stocării datelor

Volumul datelor este dat în mai multe feluri. Pentru o colectivitate ale cărei indivizi se descriu cu același șablon, volumul informațiilor este prezent prin numărul indivizilor. Avem o imagine suficient de clară, despre un fișier care conține informații referitoare la 30.000 persoane sau despre o matrice are 50 linii și 80 coloane, din punct de vedere al volumului de date.

Pentru o mai bună precizare, se vor lua în considerare:

- N - numărul de indivizi ai colectivității;
- L - lungimea structurii de date asociate unui individ;
- B - factorul de blocare;
- R - lungimea informațiilor reziduale;

Volumul de date V este dat de relația:

$$V = f(N * L, B) + g(R) \quad (23.1)$$

unde f și g sunt forme analitice ale dependenței dintre factorii considerați, forme ce se determină pentru tipuri de suport extern de date, în mod corespunzător.

Volumul de date astfel calculat este exprimat în număr de baiți. Documentațiile tehnice consemnează capacitatea de memorare pentru fiecare tip de suport extern, ca număr maxim de baiți.

Fie suportul extern i , având capacitatea C_i . Raportul:

$$p_i = \frac{V}{C_i} * 100 \quad (23.2)$$

reprezintă ponderea pe care o au datele memorate în volumul V , față de capacitatea suportului.

De exemplu, pentru un suport de 1200 ko capacitate, un fișier care ocupă 400 ko, ocupă 33% din suport.

Dacă un programator trebuie să stocheze informații pe un suport de capacitate C_i , sub forma unor volume $V_1, V_2, \dots, V_n$, el stochează numai k volume, întrucât:

$$\sum_{j=1}^k V_j \leq C_i \quad (23.3)$$

Apare însă problema existenței unei diferențe, care descrie funcția obiectiv:

$$[\min] C_i - \sum_{j=1}^k V_j \quad (23.4)$$

a unui model de optimizare a combinației de volume, care să conducă la acest obiectiv.

De exemplu, se consideră:

$$C_i = 1000, V_1 = 200, V_2 = 500, V_3 = 400, V_4 = 300 \quad (23.5)$$

Se calculează sumele:

$$S_1 = V_1 + V_2 + V_3 + V_4 = 1400 > C_i \quad (23.6)$$

$$S_2 = 200 + 500 + 400 = 1100 > C_i \quad (23.7)$$

$$S_3 = 200 + 500 + 300 = 1000 = C_i \quad (23.8)$$

$$S_4 = 500 + 400 = 900 < C_i \quad (23.9)$$

$$S_5 = 500 + 400 + 300 = 1200 > C_i \quad (23.10)$$

$$S_6 = 200 + 500 = 700 < C_i \quad (23.11)$$

$$S_7 = 500 + 300 = 800 < C_i \quad (23.12)$$

$$S_8 = 200 + 300 = 500 < C_i \quad (23.13)$$

$$S_9 = 200 + 400 = 600 < C_i \quad (23.14)$$

$$S_{10} = 400 + 300 = 700 < C_i \quad (23.15)$$

Din toate combinațiile, S_3 reprezintă varianta care conduce la o bună umplere a capacității cu date.

Apar deci ca parametri de descriere ai utilizării unui suport, următorii:

- gradul de ocupare;
- numărul de fișiere;
- volumul de informații stocat în fișier exprimat prin intermediul numărului de indivizi pentru care se face stocarea sau numărul de cuvinte, depinzând de unitățile de măsură, de natura fișierelor.

Problema maximizării, apare pentru:

- creșterea gradului de umplere;
- creșterea numărului de fișiere ce se stochează;
- creșterea volumului de informații.

Există modalități specifice, care vin să amelioreze unul sau altul dintre parametrii considerați, ceea ce influențează însă asupra tuturor parametrilor, este compactarea datelor.

Prin compactarea datelor, înțelegem totalitatea metodelor care conduc la reducerea lungimii exprimate în baiți, a datelor. Fiind dată o mulțime de cuvinte:

$$A = \{a_1, a_2, \dots, a_n\}, \text{ de lungime } l_1, l_2, \dots, l_n \quad (23.16)$$

compactarea datelor revine la a găsi funcțiile:

$$f : A \rightarrow K \quad \text{și} \quad g : K \rightarrow A \quad (23.17)$$

unde:

$$K = \{k_1, k_2, \dots, k_m\} \quad (23.18)$$

este mulțimea cuvintelor compactate, așa fel încât există o pereche (i, j) , pentru care:

$$k_j = f(a_i) \quad (23.19)$$

$$a_i = g(k_j) \quad (23.20)$$

Funcția $f()$ se numește funcție de compactare, iar funcția $g()$ este funcția de decompactare.

Deci, orice metodă de compactare este complet definită, dacă s-a identificat și modalitatea de a reduce setul de date în forma inițială, prin decompactare.

Pentru perechea (i, j) :

$$lg(k_j) < lg(a_i) \quad (23.21)$$

Rezultă că efectul compactării, pentru un text format din cuvintele:

$$a_1 \ a_2 \ a_3 \ a_3 \ a_3 \ a_3 \quad (23.22)$$

de lungime inițială:

$$L_1 = lg(a_1 \ a_2 \ a_3 \ a_3 \ a_3 \ a_3) = lg(a_1) + lg(a_2) + 4*lg(a_3) \quad (23.23)$$

prin compactare este transformat în textul:

$$k_1 \ k_2 \ k_3 \ k_3 \ k_3 \ k_3 \quad (23.24)$$

de lungime:

$$L_2 = lg(k_1) + lg(k_2) + 4*lg(k_3) \quad (23.25)$$

cu indicile de eficiență a compactării:

$$p = \frac{L_1 - L_2}{L_1} * 100 \quad (23.26)$$

23.2 Compactarea la nivel de caracter

Sistemele de coduri asociate caracterelor pornesc de la următoarele aspecte:

- mulțimea caracterelor ce sunt reprezentate este finită; de exemplu, codul ASCII permite reprezentarea unei mulțimi de caractere formate din 256 elemente;
- lungimea exprimată în biți a unui element al mulțimii, este constantă; codul ASCII asociat unui caracter are 8 biți.

Pentru un șir de n biți, se asociază o mulțime formată din 2^n elemente distincte, ce sunt puse în corespondență cu simboluri sau cuvinte, ale unei mulțimi cunoscute.

Vom considera un roman, în care apar numai litere mici și litere mari, semne de punctuație, separatorul blank și liniuța corespunzătoare semnului de dialog.

Analiza textului ce formează romanul, pune în evidență următoarele aspecte:

- dintre literele mari sunt utilizate 15;
- dintre literele mici sunt utilizate 24;
- semnele de punctuație cu liniuța de dialog, sunt în număr de 8.

Alfabetul nou cu care operăm, este format din $24+15+8 = 47$ simboluri. Fiecare simbol are reprezentare pe 6 biți, întrucât cel mai mic număr natural pentru care $2^n > 47$ este $n = 6$.

Se vor pune în corespondență cele 47 de caractere, cu coduri de câte 6 biți și textul romanului care avea o lungime inițială L_i :

$$L_i = 8 * m \quad (23.27)$$

unde m reprezintă numărul de caractere al textului.

După compactarea la nivel de caracter, textul compactat are o lungime finală L_f :

$$L_f = 6 * m \quad (23.28)$$

Indicele de eficiență al acestei compactări este:

$$p = \frac{L_i - L_f}{L_f} \cdot 100 = 25\% \quad (23.29)$$

Decompactarea, presupune interpretarea succesiunilor de 6 biți și înlocuirea lor cu caractere ASCII corespunzătoare din stiva caracterelor utilizate în text.

23.3 Compactarea la nivel de cuvânt

Prin analiza textului, înțelegem construirea unei stive a cuvintelor diferite din text și înregistrarea frecvenței lor de apariție.

Dacă menținând codul ASCII pentru caracterele uzuale ale textului, punem în corespondență cuvintele având frecvențele cele mai mari $C_1, C_2, \dots, C_k$, cu coduri asociate unor caractere ce nu apar în text, indicele de eficiență al compactării este:

$$p = \frac{L_i - \sum_{j=1}^k f_i * 1g(C_j) + \sum_{j=1}^k f_j}{L_i} * 100 \quad (23.30)$$

Dacă stiva cu cuvintele definite $C_1, C_2, \dots, C_k$, având frecvențe ridicate $f_1, f_2, \dots, f_k$, este suficient de mare, și mulțimea G a simbolurilor neutilizate în text, este insuficientă, este necesară construirea cuvintelor $g_1, g_2, \dots, g_k$ formate din $1, 2, \dots, n_g$ simboluri neutilizate, atunci performanța compactării este:

$$p = \frac{L_i - \sum_{j=1}^k f_j * [1g(C_j) - 1g_j]}{L_i} * 100 \quad (23.31)$$

Decompactarea revine la a înlocui cuvintele g_j din text, cu cuvintele c_j , întrucât atât algoritmul de compactare cât și cel de decompactare presupun existența stivei cuvintelor diferite ale textului inițial și cuvintele

din mulțimea $\{g_1, g_2, \dots, g_k\}$, a cuvintelor formate din simboluri neutilizate în text.

23.4 Compactarea prin analiza caracteristicilor textului

Vom exemplifica modul de analiză al unui text, folosind reprezentarea în memorie a tabelului de mai jos ce trebuie imprimat:

Tabelul nr. 23.1 Situația materialelor

5	30	15
NR. CRT.	DENUMIRE	VALOARE
0	1	2
1	CUIE	100
2	TABLA	200
3	VAR	600
TOTAL		900

Pentru memorarea acestui tabel, se definește un articol de 74 bairți și în fișierul *TABEL.DAT*, vor fi memorate 20 de rânduri, incluzând și blaturile dintre antet și capul de tabel. În fișierul *TABEL.DAT* vor fi ocupați $20 \cdot 74 = 1480$ bairți cu acest tabel.

Prin convenție, întrucât asteriscul nu este utilizat, în continuare este folosit ca separator, iar două asteriscuri consecutive au semnificația de *CR*.

```

27      30b*SITUATIA MATERIALELOR
23      30b*22 - **74b**10b*
27      54 = **10b*!*5b!*30b*!*15b**
60      10b*!  NR.      !                DENUMIRE                !                VALOARE
!**
25      10b*!  CRT      !*30b*!*15b*!**
22      10b*!*5b*!*30b*!*15b**
69      10b*54 = **10b*!      0      !      1      !      2      !**
69      10b*54 = **10b*!      1      !  CUIE  !      100  !**
60      10b*!      2      !      TABLA  !      200  !**
60      10b*!      3      !      VAR      !      600  !**
9      10b*54 = **
65      10b*!                TOTAL                !                900                !**
10b*54 =
-----
523 baiți

```

Textul astfel codificat are lungimea de 523 bairi.
Indicele de performanță în acest caz este:

$$p = \frac{1480 - 523}{1480} * 100 = 64\% \quad (23.32)$$

Compactarea merge mai în profunzime, prin identificarea elementelor invariante. Apar în mod repetat:

10b*54 = **
 10b*!*5b*!*30b*!*15b*!**

Dacă aceste succesiuni vor fi înlocuite, prima cu caracterul & iar a doua cu @:

$$p' = \frac{1480 - (523 - 4 * 8 - 21)}{1480} * 100 = 68\% \quad (23.33)$$

Dacă se pune în corespondență construcția ****10b*** cu :, care are 10 apariții, încă se obține o ameliorare a indicelui de eficiență.

23.5 Compactare prin asocierea unor coduri ce permit eliminarea separatorilor

În textele de orice tip ar fi ele, se utilizează diverși separatori, care ocupă un număr de poziții, depinzând de regulile acceptate de utilizare, dintre care se enumeră:

- separatorii punct și virgulă sunt urmați de un spațiu;
- separatorul linie de dialog când este urmat de o literă mare sau precedat de punct sau două puncte, este precedat de 3-6 spații pentru a marca un dialog;
- cuvintele sunt separate prin cel puțin un spațiu; în procesarea de texte, variabilitatea numărului de spații depinde de lățimea textului și de dorința de a realiza o aliniere la extremitățile definite pentru un rând.

Se consideră de exemplu, înregistrarea profesiilor persoanelor ce alcătuiesc o colectivitate. Din înregistrările efectuate, rezultă mulțimea de profesii distincte:

forjor
strungar
frezor
economist
mecanic
supraveghetor

Fișierul ce se creează, conține înșiruirea acestor profesii, urmând ca programele pentru consultarea lui, să permită numărarea elementelor ce aparțin fiecărei meserii.

Înregistrarea brută a informațiilor, conduce la ocuparea unei zone de memorie:

$$Lb = n_1 * lg(forjor) + n_2 * lg(strungar) + n_3 * lg(frezor) + n_4 * lg(economist) + n_5 * lg(mecanic) + n_6 * lg(supraveghetor) \quad (23.34)$$

Dacă fiecare meserie este pusă în corespondență cu un mnemonic precum:

<i>fo</i>	<i>pentru</i>	<i>forjor</i>
<i>st</i>	<i>pentru</i>	<i>strungar</i>
<i>fr</i>	<i>pentru</i>	<i>frezor</i>
<i>ec</i>	<i>pentru</i>	<i>economist</i>
<i>me</i>	<i>pentru</i>	<i>mecanic</i>
<i>su</i>	<i>pentru</i>	<i>supraveghetor</i>

în mod sever, lungimea ocupată se reduce.

$$L_m = 2 * (n_1 + n_2 + n_3 + n_4 + n_5 + n_6) \quad (23.35)$$

Pentru eliminarea ambiguității generate de reducerea lungimii mnemonicelor de la 2 caractere la un singur caracter, se efectuează punerea în corespondență a meseriilor cu caracterele:

<i>f</i>	<i>forjor</i>
<i>s</i>	<i>strungar</i>
<i>r</i>	<i>frezor</i>
<i>e</i>	<i>economist</i>
<i>m</i>	<i>mecanic</i>
<i>g</i>	<i>supraveghetor</i>

În aceste condiții, lungimea textului este:

$$L_s = n_1 + n_2 + n_3 + n_4 + n_5 + n_6 \quad (23.36)$$

Forma brută a caracterelor, conduce la calculul lungimii fișierului în baiți. Lungimea efectivă, exprimată în biți, este în continuare redusă dacă meseriile sunt puse în corespondență cu șiruri de biți, ce se bucură de o serie de proprietăți suplimentare:

<i>forjor</i>	<i>1</i>
<i>strungar</i>	<i>101</i>
<i>frezor</i>	<i>1001</i>
<i>economist</i>	<i>10001</i>
<i>mecanic</i>	<i>10101</i>
<i>supraveghetor</i>	<i>100001</i>

Aceste construcții, se bucură de proprietatea ca prin concatenare, locul unde s-a produs această operație apar două cifre binare de 1. Astfel:

$$L_B = n_1 * 1 + n_2 * 3 + n_3 * 4 + n_4 * 5 + n_5 * 5 + n_6 * 6 \quad (23.37)$$

Observăm că:

$$8 * L_b > 8 * L_s > L_B \quad (23.38)$$

În multe cazuri, lungimea câmpului este dimensionată în așa fel încât, să poată cuprinde meserii cu cele mai multe caractere, fișierul având articole de lungime fixă.

În exemplul dat, lungimea este dată de cuvântul "supraveghetor", care are 13 caractere.

$$L_F = 8 \cdot 13 \cdot (n_1 + n_2 + n_3 + n_4 + n_5 + n_6) \quad (23.39)$$

Toți indicii de performanță, se calculează în raport cu lungimea L_F .
Dacă de exemplu, într-o colectivitate de 1000 persoane:

$$n_1=100, n_2=300, n_3=100, n_4=100, n_5=300, n_6=100 \quad (23.40)$$

$$L_F = 8 \cdot 13 \cdot 1000 = 104000 \text{ biți} \quad (23.41)$$

$$L_b = 100 \cdot 6 + 300 \cdot 8 + 100 \cdot 6 + 100 \cdot 9 + 300 \cdot 7 + 100 \cdot 13 = 7900 \text{ baiți} \quad (23.42)$$

$$L_B = 100 + 300 \cdot 3 + 100 \cdot 4 + 100 \cdot 5 + 300 \cdot 5 + 100 \cdot 6 = 4000 \text{ biți} \quad (23.43)$$

În cazul în care lungimea codurilor asociate, iau în considerare frecvențele așa fel încât, meseria cu cea mai mare frecvență să fie codul de lungime cel mai mic, se obține:

$$L_B' = 300 \cdot 1 + 300 \cdot 3 + 100 \cdot 4 + 100 \cdot 5 + 100 \cdot 5 + 100 \cdot 6 = 3000 \text{ biți} \quad (23.44)$$

$$p_1 = \frac{104000 - 3000}{104000} \cdot 100 = 97\% \quad (23.45)$$

$$p_2 = \frac{7900 \cdot 8 - 3000}{7900 \cdot 8} \cdot 100 = 95\% \quad (23.46)$$

$$p_3 = \frac{4000 - 3000}{4000} \cdot 100 = 25\% \quad (23.47)$$

23.6 Compactarea prin identificarea de subșiruri repetitive

Pentru cele 27 litere ale alfabetului, se construiește matricea frecvențelor de apariție a grupurilor de câte două litere, X . Astfel, x_{ij} reprezintă numărul de apariții al literei cu poziția i , urmată de literă cu poziția j din alfabet.

Construirea matricei X se realizează, prin parcurgerea unei diversități de texte și frecvențele depind de particularitățile fonetice ale fiecărei limbi.

Dintre grupurile de câte două litere, vor fi extrase acelea cu frecvențele cele mai mari și vor fi dispuse pe linii într-un tabel, ale cărui coloane conțin literele alfabetului.

Se construiește matricea Y , ale cărei elemente y_{ij} , conțin frecvențele de apariție ale grupului de două litere de pe linia i , urmat de litere de pe coloana j a tabelului.

Se are în vedere că totalitatea grupurilor de litere ce vor fi selectate în ordinea descrescătoare a frecvențelor de apariție, să nu depășească un număr K așa fel încât:

$$K + L < 256 \quad (23.48)$$

unde L reprezintă numărul de caractere considerat necesar pentru introducerea unui text într-un fișier.

În continuare, se vor considera cele 27 litere ale alfabetului, cele 10 simboluri ale cifrelor și caracterele: spațiu, plus, minus, egal, punct, virgulă, două puncte, punct și virgulă, semnul mirării, semnul întrebării și asteriscul. În total sunt 48 de caractere.

Din analizele statistice efectuate pe texte, se rețin grupurile de litere alcătuind o mulțime formată din 64 de elemente dispuse în tabloul de mai jos, care au în dreptul liniilor și coloanelor combinații de biți care alcătuiesc codul asociat fiecărui șir.

Tabelul nr. 23.2 Combinații de biți asociate șirurilor de caractere

	1000	1001	1010	1011	1100	1101	1110	1111
1000	u1	1e	pr	mb	mp	ni	in	lui
1001	lor	se	ut	re	tr	te	ta	nu
1010	it	la	ea	ta	ti	ca	oi	au
1011	am	ar	ei	ra	ne	un	ns	nt
1100	cr	sc	st	os	ti	at	ri	oa
1101	sa	ma	ne	tre	ist	tri	urile	ind
1110	nstr	eau	eam	esti	ndu	u-se	ati	nul
1111	asera	res	tit	ros	oasa	isem	tit	art

Aceste grupuri de litere au fost puse în corespondență cu codul unui caracter, altul decât cele 48 considerate.

Astfel, versurile eminesciene:

*Dintre sute de catarge
Care leagă malurile
Câte oare le vor sparge
Vânturile, valurile*

*A fost odată ca-n povești
A fost ca niciodată
Din rude mari împărătești
O prea frumoasă fată.*

care însumează 177 caractere incluzând și spațiile care separă cuvintele, vor fi compacte astfel:

*Dx x sux de x x rge
 17 64 26 36 27*

*x x x aga x lux x
 26 24 12 62 57 12*

Cix x x x vor spx ge
 26 58 24 12 42

Vix ux x , valux x
 48 57 12 57 12

A fox odată ca-n povex
 53 73

A fox x x ciodată
 53 26 16

Dx rude x x ix aratx
 17 62 57 15 74

o x x fata
 13 33 85

ceea ce conduce la un indice de performanță:

$$p = \frac{177 - 137}{177} * 100 = 22\% \quad (23.49)$$

Construirea matricei generale a subșirurilor, are avantajul că este unică pentru orice text care se compactează, dar există posibilitatea de apariție a situației ca frecvențele grupurilor în textul de compact să nu urmeze nivelurile de frecvență a textelor care au stat la baza obținerii ei.

Dacă pentru fiecare compactare se construiește o matrice de subșiruri proprie, performanța este cu totul alta.

Pentru textul analizat, subșirurile identificare se organizează într-un tabel, căruia i se atașează o matrice C a codurilor, ce conduce la un text de 116 caractere, având un indice de performanță:

$$p = \frac{177 - 116}{177} * 100 = 34\% \quad (23.50)$$

Dacă pornim de la ideea că acest text este memorat folosind succesiuni de 6 biți, pentru că el conține 25 de combinații de biți, corespunzătoare grupurilor de litere și cele 27 de litere, comparativ cu memorarea ca text având fiecare caractere câte 8 biți, indicele de performanță este:

$$p = \frac{177 * 8 - 116 * 6}{177 * 8} * 100 = 52\% \quad (23.51)$$

Se observă că algoritmi de compactare, vizează atât lungimea codului sub care se reprezintă un caracter din text, cât și modul în care se identifică elementele invariante în texte.

Spațiul ocupat de textul compactat, i se adaugă o zonă cu informații care permit reconstituirea textului inițial. Aceste informații conțin:

- lungimea codurilor asociate caracterelor;

- mulțimea subșirurilor și a codurilor cu care acestea se pun în corespondență.

În cazul în care se identifică pentru reguli complexe de scriere a textelor proceduri, acestea se pun în corespondență cu coduri și ori de câte ori apar codurile respective în text, vor fi activate procedurile care vor prelucra textul adecvat.

De exemplu, dacă un text este centrat pe un rând, blancurile vor fi eliminate, respectivul text fiind precedat de un cod care odată identificat, preia textul, îl centrează, reconstituind blancurile eliminate la compactare.

În cazul dialogurilor, începerii unui nou paragraf sau scrierii unei formule, reconstituirea poziției reale a textului, se efectuează cu ajutorul procedurilor activate odată cu apariția codurilor care semnalizează fiecare dintre situațiile menționate.

23.7 Compactarea programelor date în formă executabilă

Pentru utilizatori, prezintă importanță programele executabile. Obiectivele programelor care efectuează compactarea acestui tip de text, sunt găsirea unor modalități care să determine stocarea unui număr cât mai mare de programe executabile pe un suport. În același timp, se urmărește și găsirea posibilității de a efectua decompactarea textului transformat.

Programele de compactare a programelor executabile, iau în considerare următoarele aspecte:

- limbajul de asamblare are o mulțime finită de coduri de instrucțiuni, dintre care programatorii folosesc sub 40%, ca diversitate în programe;
- programele executabile sunt rezultate ale compilării și editării de legături; în cea mai mare parte, aceste operații conduc la o pondere ridicată a secvențelor construite mecanic, pe baza unor reguli tip;
- programele executabile, au ca entitate instrucțiunea și aceasta este plasată într-o zonă de memorie de lungime și structură fixă; toate analizele, vizează componentele în număr restrâns de pe o zonă restrânsă; frecvențele construite vin să ajute alături de celelalte considerații, la dimensionarea codurilor cu care se pun în corespondență, instrucțiunile programului executabil.

Se consideră în continuare programul:

	ORG	420	H
	MOV	D , E	
	PUSH	B	
	XRA	A	
CICLU:	LDAX	B	
	ADC	M	
	DCR	E	
	JZ	BETA	
	STAX	B	
	INX	B	
	INX	H	
	JMP	CICLU	
BETA:	MOV		

```

                                LOAX  B
                                XRA  M
                                MOV  A , E
                                STAX  B
                                STC
                                JM   GAMA
                                MOV  A , M
                                XRA  E
                                STC
                                JM   DELTA
GAMA: CMC
DELTA: POP  B
        MOV  E , D
        RET
        END

```

Acestui program îi corespunde codul obiect:

```

0420 53
0421 C5
0422 AF
0423 0A
0424 8E
0425 1D
0426 CA 2F 04
0429 02
042A 03
042B 23
042C C3 23 04
042F 5F
0430 0A
0431 AE
0432 7B
0433 02
0434 37
0435 FA 3F 04
0438 7E
0439 AB
0440 37
043B FA 3F 04
043F C1
0440 5A
0441 C9
0000

```

Textul obiect generat are 31 baiți. Se vor scrie frecvențele de apariție a elementelor din acest text.

Tabelul nr. 23.3 Frecvențele de apariție a elementelor textului

Element	Frecvența
0	8
1	2
2	4
3	9
4	4
5	4
6	-
7	4
8	1
9	1
A	6
B	1
C	3
D	1
E	4
F	6

Împrăștierea cifrelor hexazecimale, reduce masiv posibilitatea efectuării unor prelucrări directe pe textul obiect.

Se procedează la efectuarea modificărilor:

a) se începe contorizarea de la zero și se memorează 0420 și textul devine:

0000	53		
0001	C5		
0002	AF		
0003	0A		
0004	8E		
0005	1D		
0006	CA	00	0F
0009	02		
000A	03		
000B	23		
000C	C3	00	03
000F	5F		
0010	0A		
0011	AE		
0012	7B		
0013	02		
0014	37		
0015	FA	00	1F
0018	7E		
0019	AB		
001A	37		
001B	FA	00	1F
001F	C1		

0020	5A
0021	C9

Cea mai mică valoare din prima cifră hexazecimală a părții de instrucțiune, este 0 și cea mai mare este F.

Tabelul nr. 23.4 Frecvențele de apariție a elementelor textului

Element	Frecvența
0	5
1	1
2	-
3	2
4	-
5	2
6	-
7	2
8	1
9	-
A	2
B	-
C	5
D	-
E	-
F	2

Din 15 simboluri lipsesc 7. Construim pe 2 baiți masca elementelor absente:

0	1	2	3	4	5	6	7	8	9	A	B	C	D	E	F
1	1	0	1	0	1	0	1	1	0	1	0	1	0	0	1

Se analizează frecvența de apariție a cifrelor hexazecimale pe al II-lea bait al instrucțiunii.

Tabelul nr. 23.5 Frecvențele de apariție a cifrelor hexazecimale pe baitul II al instrucțiunii

Element	Frecvența
0	-
1	1
2	2
3	4
4	-
5	1
6	-
7	2
8	-
9	1

A	6
B	2
C	-
D	1
E	3
F	2

Se construiește pe 2 baiți masca elementelor absente:

0	1	2	3	4	5	6	7	8	9	A	B	C	D	E	F
0	1	1	1	0	1	0	1	0	1	1	1	0	1	1	1

Se observă că pentru un text de lungime redusă, elementele repetitive nu se identifică și deci compactarea are efecte reduse, uneori nule sau chiar păgubitoare.

Limbajele de asamblare moderne, au în definire operații de tip *RR* (registru-registru) și întrucât numărul registrelor este redus, sunt asociate coduri diferite pentru combinația de cod operație *R1, R2* ceea ce reprezintă o compactare la nivel de limbaj. Tot astfel, în cazul instrucțiunilor memorate pe trei baiți, deplasarea operandului este reprezentată pe un singur baît, prin recalcularea ei.

În cazul unor texte mult mai lungi, dacă resursele sunt folosite convenabil și programatorii stabilesc reguli precise de realizare a unor anumite prelucrări, compactarea este realizabilă.

Astfel, în toate subprogramele, secvența de preluare a parametrilor este standardizată și deci este pusă în corespondență cu un cod și înlocuită cu acesta.

Aplicând principiile de analiză ale textului, compactarea se realizează în principal prin tratarea elementelor repetitive și în programele scrise în limbajul de asamblare sau generate de compilatoare, acestea nu au o pondere importantă.

23.8 Compactarea datelor numerice

Reprezentarea matricelor rare, este un exemplu de compactare a datelor. În cazul seriilor de date, problema compactării este importantă dacă seriile au variații mici sau dacă seriile au o lungime foarte mare.

Fie seria de date *X* având termenii:

11120	11130	11131	11136
11170	11135	11131	11109
11121	11122	11120	11104

Cei 12 termeni ai seriei de date, dacă sunt reprezentați binar, ocupă fiecare câte 2 baiți, deci în total 24 baiți. Observăm că termenul minim al seriei este 11104, iar termenul maxim este 11170.

$$D = X_{max} - X_{min} = 66 \quad (23.52)$$

iar:

$$\frac{D}{X_{min}} = \frac{60}{11108} = 0,005 \quad (23.53)$$

ceea ce arată că datele se află într-un interval foarte îngust în raport cu mărimea lor.

Considerăm $X_{min} = 11104$. Se calculează o nouă serie x' :

$$x'_i = x_i - x_{min} \quad (23.54)$$

al cărei termeni sunt:

$$\begin{array}{cccccc} 16 & 26 & 27 & 32 & 66 & 51 \\ 27 & 5 & 17 & 18 & 16 & 0 \end{array}$$

Numerele acestea se reprezintă pe 7 biți. Deci, pentru cei 12 termeni sunt necesari 84 de biți, iar pentru 11104 sunt necesari 2 biți.

$$p = \frac{84 + 2 * 8}{24 * 8} * 100 = \frac{100}{192} * 100 = 52\% \quad (23.55)$$

Compactarea la nivel de biți, se dovedește cu efecte mai importante dacă se combină cu alte procedee de compactare.

$$\begin{array}{cccccc} 111 & 111 & 111 & 112 & 112 & 112 \\ 113 & 113 & 111 & 112 & 112 & 112 \end{array}$$

Elementul minim este 111, iar elementul maxim este 113. În reprezentare binară a întregilor, pentru fiecare termen este necesar un байт. Seria aceasta necesită 12 байți.

Observăm că există elementele repetitive 111, 112 și 113, care vor fi puse în corespondență cu caracterele a, b, c, respectiv 00, 01, 10. Pentru cele trei numere sunt necesari 7 biți.

$$p = \frac{3 * 7 + 12 * 2}{12 * 8} * 100 = \frac{45}{96} * 100 = 46\% \quad (23.56)$$

23.9 Programe care efectuează compactarea

Un program P de lungime L , se zice că este compactat de programul X , dacă forma obținută P' , are o lungime L' cu peste 20% mai mică decât lungimea L .

Programul X execută complet compactarea, dacă încercând compactarea programului P' , se obține un program P'' de lungime L'' și dacă $L' = L''$ și P'' este identic cu P' .

Programul Y realizează decompactarea programului P' , dacă după executarea sa se obține programul $P' ' '$ de lungime $L' ' '$, așa fel încât $L' ' ' = L$ și $P' ' '$ este identic cu P' .

Programele de compactare sunt specializate să aibă ca intrare, fie texte sursă, fie texte obiect, fie componente executabile. Există și programe generale care efectuează compactarea și decompactarea. Exemple de astfel de programe sunt PKZIP . EXE, PKUNZIP . EXE, ARJ . EXE, PKARC . EXE. Fiecare dintre acestea folosesc algoritmi specifici de compactare, în funcție de tipul fișierelor supuse comprimării.

Formatul general al unei etichete de fișier arhivat, îl descriem prin structura următoare în C:

```
struct fisier_compactat
{
    int metoda_compactare[10];
    char nume_fisier[8];
    char extensie_fisier[3];
    unsigned char atribut_fisier;
    short int data;
    short int timp;
    short int cluster_start;
    long int lungime_fisier_initial;
    long int lungime_fis_compactat;
    int suma_control;
};
```

Se asociază fiecărei metode de compactare folosită un anumit cod, ca de exemplu:

- 0, dacă fișierul a fost memorat în forma sa inițială, compactarea nefiind eficientă;
- 1, dacă fișierul a fost comprimat printr-un algoritm nerepetitiv;
- 2, dacă fișierul a fost comprimat în mai multe etape.

Aceste probleme sunt înzestrate cu funcții specifice gestionării de fișiere arhivate, respectiv adăugare, ștergere, modificare, protecție.

Unul dintre algoritmii utilizați de aceste programe de compactare este cel al lui Huffman. Acest algoritm are drept scop de a minimiza numărul de biți necesar pentru reprezentarea oricărui simbol dintr-un alfabet. Astfel, se atribuie simbolurilor cu o frecvență de utilizare mare, coduri mai scurte, iar simbolurilor mai rar folosite, coduri mai lungi. În acest mod, media lungimii codului, este mai redusă.

Această tehnică a utilizat-o și Samuel Morse când a definit alfabetul Morse.

În codul Huffman, frecvența de apariție a caracterelor din textul sursă, trebuie cunoscută sau estimată. Apoi, se asociază fiecărui simbol o secvență de biți unică, care este definită utilizând un arbore binar.

Un cod Huffman este construit astfel:

- se ordonează descrescător, după frecvență sau probabilitatea de apariție, simbolurile din textul sursă;
- se consideră fiecare simbol un nod în arbore, fiecărui nod fiindu-i asociat o probabilitate (frecvența) de apariție;
- se leagă două noduri cu probabilitatea de apariție cea mai mică, într-un nod a cărui probabilitate este dată de suma probabilităților celor două noduri acest proces se încheie în momentul în care rămâne un nod nelegat.

Rezultatul acestei operații este un arbore binar, în care ramura stângă a unui nod părinte are eticheta 0, iar ramura din dreapta, eticheta 1.

Să presupunem că textul nostru inițial este format din 40 de caractere, utilizând un alfabet din 8 simboluri; frecvențele de apariție ale fiecărui simbol, sunt calculate ca raport între numărul de apariții a simbolului respectiv și numărul total de caractere ale textului.

În tabelul de mai jos, presupunem numărul de apariții ale simbolurilor și în funcție de acest număr calculăm numărul total de biți pe care îl ocupă simbolul respectiv, în reprezentarea normală (8 biți/caracter).

Tabelul nr. 23.6 Varianta I de construire a codului Huffman

Simbol	Nr. de apariții simbol	Total nr. de biți pentru un simbol (reprez. normală)	Total nr. de biți pentru un simbol (cod Huffman)
0	12	96	12
1	8	64	16
2	6	48	18
3	5	40	20
4	4	32	20
5	3	24	18
6	2	16	14
7	0	0	0
Total	40	320	118

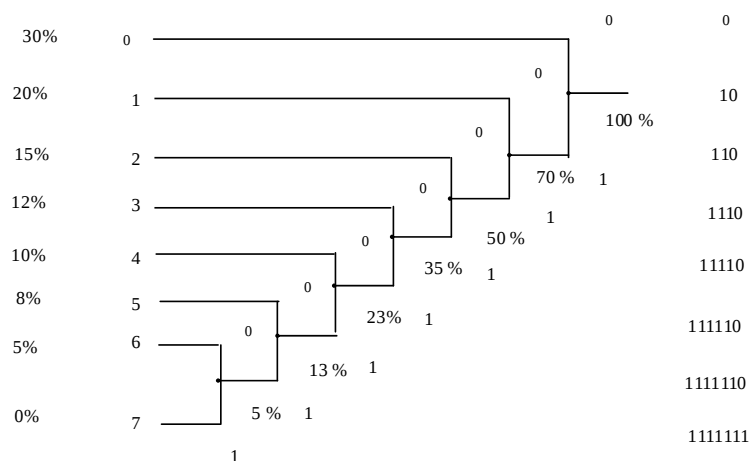


Figura 23.1 Structura arborescentă a codurilor Huffman

Dacă însă, cuplăm nodurile în felul următor:

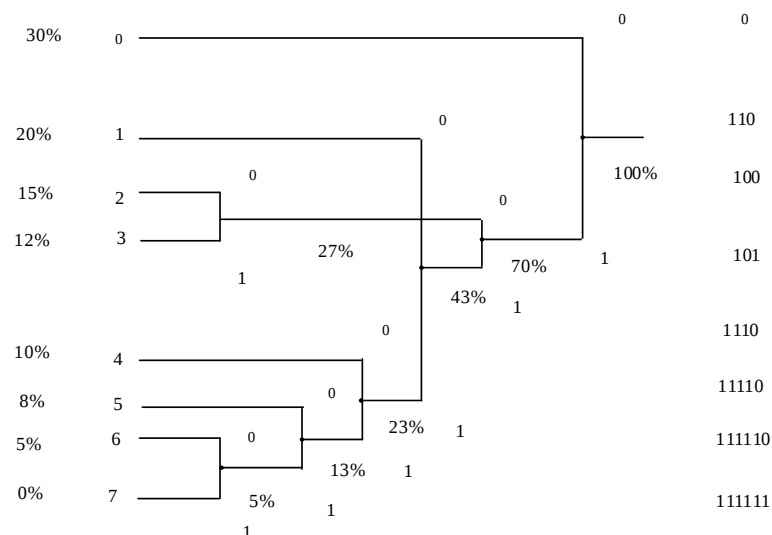


Figura 23.2 Structura arborescentă a codurilor Huffman după cuplarea nodurilor

vom obține:

Tabelul nr. 23.7 Varianta II de construire a codurilor Huffman

Simbol	Nr. de apariții simbol	Total nr. de biți pentru un simbol (reprez. normală)	Total nr. de biți pentru un simbol (cod Huffman)
0	12	96	12
1	8	64	24
2	6	48	18
3	5	40	15
4	4	32	16
5	3	24	15
6	2	16	12
7	0	0	0
Total	40	320	112

și se observă că am obținut un număr total de biți mai mic decât în prima variantă, iar diferența este cu atât mai mare, cu cât avem texte de lungime mai mari.

Observăm, că pe lângă faptul că codul Huffman folosește în medie, un număr de biți mai mic decât codul ASCII pentru reprezentarea unui caracter, mai are proprietatea că secvența binară asociată fiecărui simbol, nu este prefixul unui alt simbol și de aceea, secvențele de biți sunt concatenate fără să se folosească nici o punctuație de delimitare a acestora.