

統計データベースにおける 一連検索用ファイル編成法

近藤 貴士

打浪 清一

手塚 慶一

Takashi Kondo

Seiichi Uchinami

Yoshikazu Tezuka

(大阪大学 工学部)

1. まえがき

近年、多目的有効利用可能な統計データベースの必要性が認識され、その実現が望まれている。データベース構築の際、ファイル編成法はその応答性に直接影響を与えるため、データの特性またその利用形態に適した手法を選ぶ必要がある。統計データは多次元性を有しているためその格納には多次元ファイル編成法が適している。さらに統計解析においては多数の項目のうち少数の項目のみに着目した一連検索が用いられる場合が多い。一連検索とはある項目に対し同じ項目値を持つ多くのレコードを一括して処理する必要がある検索をいう。しかし従来の多次元ファイル編成法はこのような検索に対し良好な特性を示さない。そこで本稿では統計データの一連検索に対し有効な多次元ファイル編成法を提案する。

2. 統計データの特性と要求分析

2. 1. 要求分析

統計データに見られる種々の特徴のうちファイル編成法を考えるうえで必要となる事項を挙げる。

(i) 多数のレコード

通常数千～数億程度である。

(ii) 多数の項目

一般に調査項目が多い。しかし検索にはごく少数の項目しか利用しない。

(iii) 多次元性

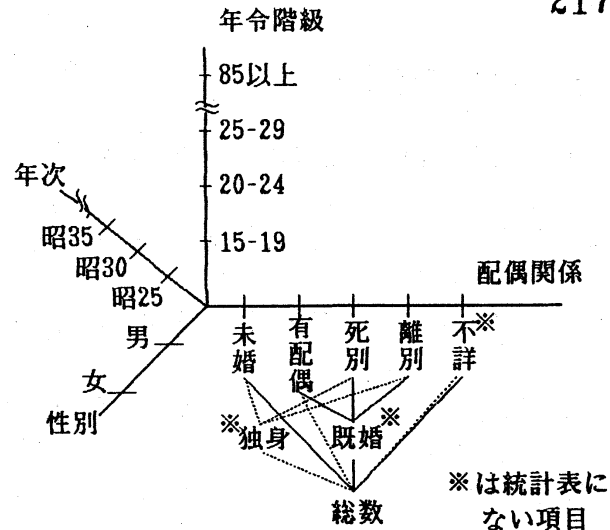


図1. 統計データの多次元性,階層性と特殊値

統計データの記述モデルとしては調査項目を軸とする多次元空間を用いると直感的に理解しやすい。利用者はこの多次元性に基づき様々なビューから検索を行う。

(iv) 項目値の階層性

各項目の項目値には階層性が存在する。階層性により上位のカテゴリを指定した検索は、下位のカテゴリに属するレコード全てが対象になる。

以上の特徴を持つデータに対する検索を必要レコード数によって次の2つに大別できる。

・ トランザクション検索(transaction retrieval)

ごく少数のレコードを処理するだけで要求を満足できるもの。

・ 一連検索(consecutive retrieval)

ある項目に対し同じ項目値を持つ多くのレコードを一括して処理してはじめて要求を満足できるもの。

統計処理では多数の多項目のレコードからなるファイル(データベース)を用いて少数項目のレコード少数からなる表を抽出する操作(要約)が多く、こ

れは一般に比較的多数のレコードを対象に行う処理で一連検索を要する。つまり一連検索は統計データを利用するうえで頻繁に行なわれる。

以上をまとめると、統計データベースには

- (1) 多数の多項目レコードの格納、管理が可能
 - (2) 少数項目に着目した、比較的多数のレコードを対象とする検索を高速に行える。
 - (3) 検索条件の選びかたによるパフォーマンスの低下がほとんどない。
- の条件を満たすファイル編成法が適している。

2. 2. 目的関数

統計データは一括入力されることが多く、更新あるいは削除されることが稀であるため目的関数 F は、 $F = Tr(R)$ とする。ここで Tr は検索に対する応答時間で、 R は検索条件である。CPU の演算コスト、仮想記憶域のトポロジ等を見捨てると、 $F \propto B_n(R)$ とおける。ここで B_n は検索条件 R に対する答えを返すのに必要なデータを含む仮想記憶域の単位データ(バケット)の総数である。以後 B_n をアクセスコストと呼び、これを小さくするファイル編成法を考える。ただしこれは検索条件・データの持つ性質などによって様々であるため、評価時はこれらに仮定を行う必要がある。

3. 多次元ファイル編成

従来提案されている種々の多次元ファイル編成のうち代表的な手法を取り上げ、これらを統計データに適用すると以下のようなになる。

(i) Inverted file

多数、多項目のレコードに対する転置リストが大きくなりすぎる。

(ii) k-d tree , Bucket Resolved Clustered Index File(BRCIF)

検索に関係する項目は全項目のうちのごくわずかであるため、レコード単位の格納形式であるこれらの手法は不要な項目の項目値へのアクセス分が本質的に多い。

(iii) Transposed file

検索はスキャンに基づいて行なわれるため不要なレコードの項目値に対するアクセスが多い。

(iv) Double Count Header Compression Scheme (DCHCS)

トランザクション処理に有効とされているが、アクセス対象データが散在するため一連検索に適さない。

検索レコード数と参照項目数によりこれらの手法が適している領域を図示すると図2となる。本稿で提案する統計データベースに適したファイル編成は*の部分に該当する。

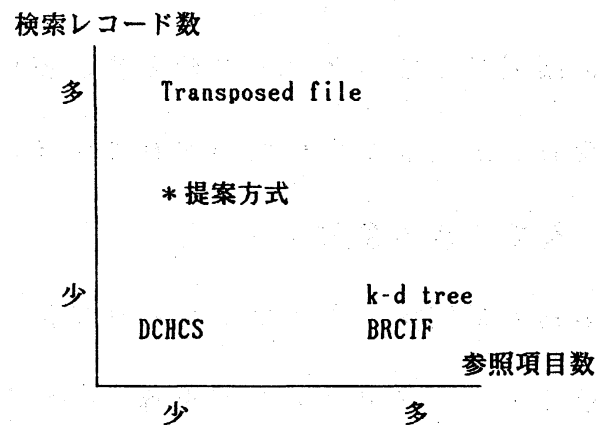


図2. ファイル編成法の最適領域

4. 本方式の設計概念

アクセスが必要なバケット数を低減させるための基礎概念は次の2つに大別できる。

(i) データのクラスタリング

クラスタリングとは、同時に必要となる確率の高いデータ(類似データ)を物理的に近い位置に配置させることをいう。

・局所的クラスタリング

類似したデータは同じバケットにまとめて格納する。これによりアクセス対象となるバケット数を抑えられる。

・大域的クラスタリング

質問に対しアクセスが必要となるバケットが全記憶域内で集中するようにバケットの位置付けを行う。本稿ではこれについては言及しない。

類似データを物理的近傍に配置することにより、アクセス時の物理レベルでの選択・射影の効果が期待できる。

(ii) データ圧縮

計算量に多少のオーバーヘッドが生じても適切な圧縮法により記憶域が大幅に削減できるとアクセス時間や記憶域の大きさの点で有効である。

なおデータ量が多いため冗長性が高いファイル編成法は不適當である。

5. 提案手法の概要

本稿で提案する方式はレコード単位のクラスタリングを行うファイル編成に項目別データ格納の考えを導入したもので、つまり多次元空間をいくつかの部分空間に分割しそれぞれについて項目別ファイルを作成する方法をとる。

このように本手法は多次元空間 $D_1 \times D_2 \times \dots \times D_n$ (n は項目数) を分割した部分空間が基本となっている。部分空間が全て直積ファイルとなるとときこれら1つ1つをブロックと呼ぶことにする。ブロックは類似したレコードを管理する単位となっており、項目数個の順次ファイルとそれらへのインデックスから構成される。レコードはそれが属するブロックに、項目ごとに分割され格納される。順次ファイルにはそれぞれある項目が割り当てられており項

目値の格納に用いる。

本方式は必要なブロック内レコードのみをアクセス対象とすることにより選択が行え、かつ選ばれたレコードのうち必要項目のみをアクセス対象とすることにより射影が行えるためアクセスコストを低減するのに有効である。

(1) 入力アルゴリズム

(i) 投入されたレコードについてその属する部分空間に対応したブロックを求める。

(ii) 求めたブロックを構成する順次ファイルに投入されたレコードの各項目の値を別々に格納する。

順次ファイルについて次のことがいえる。

・一般に項目により必要な単位データ長(項目長)が異なるため順次ファイルの大きさは異なる。

・ブロックへの分割により各項目の項目値はブロック内では限定された値しかとり得ないことを利用してデータ

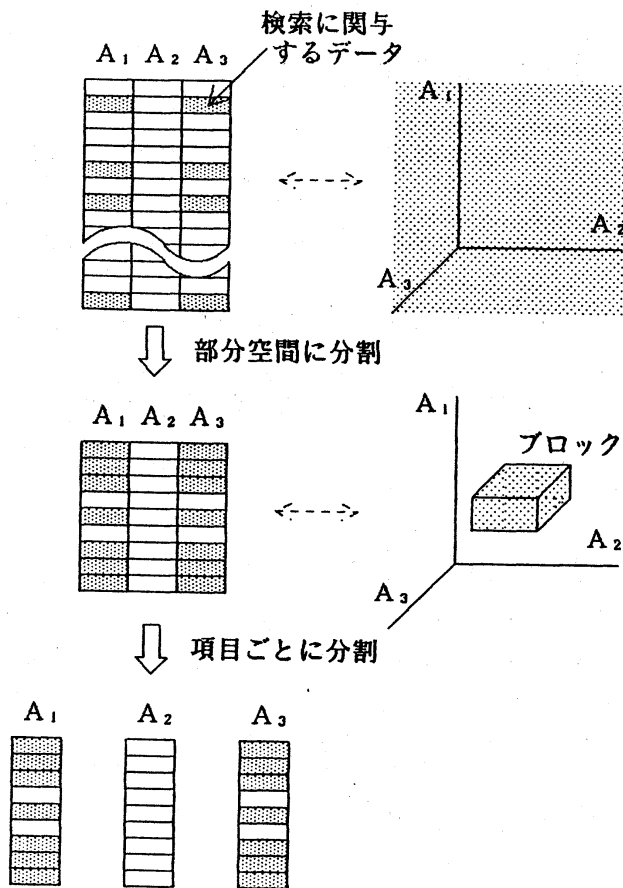


図3. 提案手法の概要(多次元空間との対比)

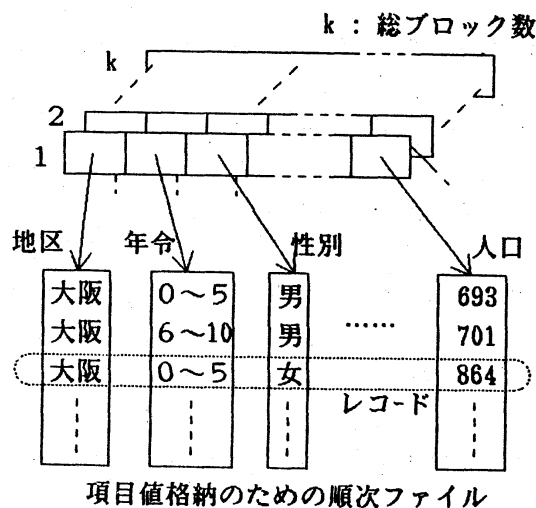


図4. ブロックの物理的構成図

圧縮が可能である。

(2) 検索アルゴリズム

ブロックは全多次元空間内で条件を与えている。ブロック i が示す条件を C_i で表すことにすると、与えられた検索条件 Q に対し、 $C_i \wedge Q$ が偽でないブロックのみに検索対象となるレコードが含まれている可能性がある。このようなブロックを検索対象ブロックという。検索対象ブロック全てに対し条件を満足するレコードをスキヤニングにより特定し必要な項目値を取り出す。検索条件を与えている項目を条件項目、出力要求されている項目を参照項目と呼ぶことにすると、アルゴリズムは次のようになる。

(i) 条件を満足しうるブロックを全て求める。

(ii) 求めたブロックに対し、条件項目について項目値ファイルを先頭からスキャンしていき、条件が満たされたレコードの参照項目の値を出力する。

アクセスコストは、項目値順次ファイルを構成するバケット数を、対象ブロック全てと条件項目、参照項目全てにわたって総和をとった値となる。実際にはデータ圧縮が行えるためこれより少し小さい値となる。このコストは分割法によって様々になる。より分割数を大きくするほど不要なデータへのアクセスを減少できるが、逆にバケット内のメモリ利用率低下により結果としてアクセスが増大する。このように分割にはトレードオフが存在するため適切な分割を定める必要がある。

レコードが存在しうる全空間をブロックへ分割する方法としては種々考えられるが、ここではまずファイル作成時に分割を決定する方法について述べ

る。これを静的ファイル編成法とよぶ。説明のために以下の記号を定義する。

n : 全項目数

A_i ($1 \leq i \leq n$) : 項目

N : 総レコード数

D_i : 項目 i の領域(項目値の全体集合)

6. 静的分割法

6. 1 静的分割法の概要

ブロックの定義をレコード入力の前に行う。

ブロック管理が最も簡単になる方法は全定義域を塞の目状に分割し1つ1つをブロックとするものである。この例を図5に示す。

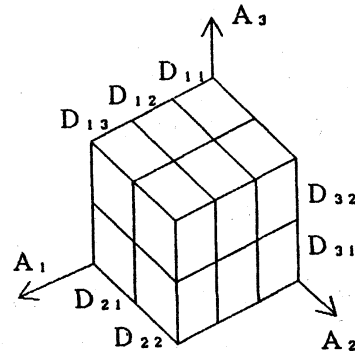


図5. 静的分割法の例

静的分割法の評価のための準備として分割

に対する評価法について述べる。

6. 2 静的分割法の評価法

ここでは確率を用いた評価法を採用した。

レコードはあるブロックへ確率 p で落ち込むとする。前述したとおり、評価対象はアクセスされるべきバケットの総数であるが、これは検索対象ブロックについて条件項目と参照項目の項目値を格納するバケット数 $X = \lceil r \cdot \ell / B \rceil$ の総和をとったものである。ここで r はブロック内レコード数(確率変数), ℓ は項目長, B はバケット容量である。以下では1つの項目を代表させ、そのバケット数を考える。

レコードが全部で N 個投入されるとすると、あるブロックのレコード数が r となる確率 $P(r)$ は2項分布 $B(N, p)$ に従う(図6)。ここで $N \cdot p$ が十分大きいと仮定できるとすると2項分布 $B(N, p)$ を正規分布 $N(Np, Npq)$ に近似可能

となる。そこで正規化を行い、 $P(r)$

誤差関数 $\text{erf}(x)$ を用いると、 X

の期待値 $E(X)$ は簡単に計算で

きる。

さて、確率変数の平均につい

ては公式、

$$E(X_1 + X_2 + \dots + X_i)$$

$$= E(X_1) + E(X_2) + \dots + E(X_i)$$

が成立するため、ある項目に絞っ

て見たときある検索条件に対し

アクセスが必要となるバケット

数は、平均値を考える限りバケットごとに算出される平均値の和をとるだけで計算可能である。これを検索に関与する項目全てについて和をとると求めるアクセスコストが得られる。

静的分割法について適切な分割法を決定する一般的規則は未知であるが、与えられた分割法に対する評価を以上の方法により行うことができる。

6. 3 静的分割法の評価と考察

簡単のためにつぎのように仮定したデータについて、前述した評価法に基づきアクセスコストの期待値を算出する。

- ・レコード数 : 100万
- ・項目数 : 5
- ・項目値の発生確率は一様で各項目で独立している。
- ・項目長 : すべて8 (bit)
- ・バケット容量 : 16384 (bit)

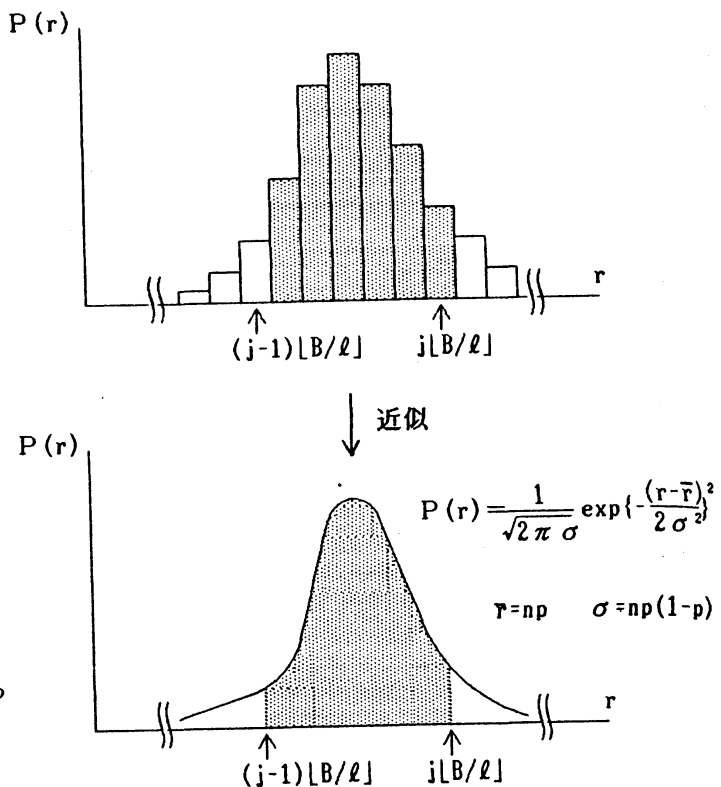


図6. バケット数が j になる確率

- ・ 分割は各項目同様に等分割を行う。
- ・ 検索条件は partial match queryとする。

以上の仮定を用い分割数 1 (分割しない), 2, ..., 5 について 1 項目あたりのコストの期待値を算出した結果を表 1 に示す。ここではデータ圧縮を無視している。

表 1. アクセスコストの期待値(小数部は四捨五入)

c	例. k	1	2	3	4	5
0	(*,*,*,*,*)	489	512	652	1024	3125
1	(0,*,*,*,*)	489	256	217	256	625
2	(0,0,*,*,*)	489	128	72	64	125
3	(0,0,0,*,*)	489	64	24	16	25
4	(0,0,0,0,*)	489	32	8	4	5
5	(0,0,0,0,0)	489	16	3	1	1
総ブロック数		1	2 ⁵	3 ⁵	4 ⁵	5 ⁵

c : 条件項目数 k : 分割数 * : don't care condition

計算結果の表で分割数 1 は transposed file と同じである。そこで表から提案方式を transposed file と比較すると次のことがわかる。

- ・ 条件項目が存在すると Transposed file よりアクセスコストが少なくてすむ分割法が存在する。

- ・ 条件項目が多いほど必要となる 1 項目あたりのバケット数が少ない。

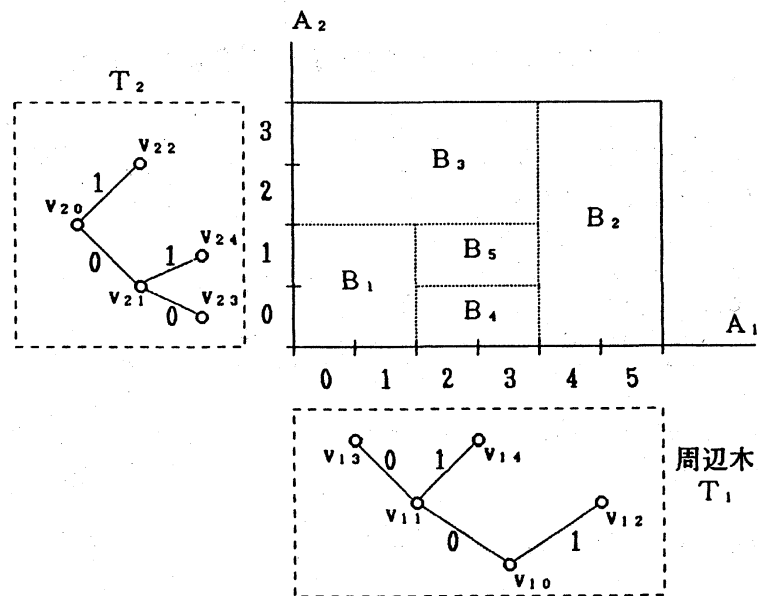
なお項目長などと密接な関係があり最適な分割法は case by case となる。

静的分割法は格納されるレコードの数や分布, 項目長などが既知であるときは最適分割を考えることができようが、実際には最適分割を定めるのは困難であることが予想され、データは逐次投入され分布も未知である。そこで投入されたデータにより自動的にブロック分割を行う動的分割法が実用的であると考えられる。

7. 動的分割法

7. 1 動的分割法の概要

本方式は最初全レコードを格納するブロックが1つだけ定義されており、レコード格納によりブロック内レコード数が大きくなり後述の基準を超えたとき、ブロック内レコードを後述の計算により定まる項目の項目値によって



(a) 周辺木とブロック

ノード	子ノード		ブロック				
	左	右	1	2	3	4	5
v10	v11	v12	0	0	0	0	0
v11	v13	v14	0	0	1	0	0
v12	—	—	0	1	0	0	0
v13	—	—	1	0	0	0	0
v14	—	—	0	0	0	1	1
v20	v21	v22	0	1	0	0	0
v21	v23	v24	1	0	0	0	0
v22	—	—	0	0	1	0	0
v23	—	—	0	0	0	1	0
v24	—	—	0	0	0	0	1

$v_{10}=\{0,1,2,3,4,5\}$

$v_{11}=\{0,1,2,3\}$

$v_{12}=\{4,5\}$

$v_{13}=\{0,1\}$

$v_{14}=\{2,3\}$

$v_{20}=\{0,1,2,3\}$

$v_{21}=\{0,1\}$

$v_{22}=\{2,3\}$

$v_{23}=\{0\}$

$v_{24}=\{1\}$

(b) ビットマップファイル

図7. 動的分割法のブロック管理法

2分割する方式である。

本方式はブロック管理に n 個の周辺木とビットマップを用いる。周辺木 T_i は2分木で、項目 i の項目値の集合に対応するノードから構成される。根は項目 i の全項目値に対応する。あるノード j がレベル h で、ある項目値集合 S_j に対応するとき、ノード j の左子ノードは S_j のうち第 h ビットが“0”である項目値部分集合に対応し、ノード j の右子ノードは S_j のうち第 h ビットが“1”である項目値部分集合に対応する。このようにノードは項目値集合をあらわす。またこれに従い、左子ノードへのアークを“0”,右子ノードへのアークを“1”とラベル付けする。

全てのブロックはこれらノードが示す項目値集合を一辺とする超直方体に対応させる。このような、ノードとブロックの対応づけを行うのがビットマップである。つまりあるノード v_{ij} のビットマップで“1”であるブロックは項目 i のとりうる項目値は v_{ij} である。

ブロックが特定されるのに用いた周辺木 T_i 上の経路は0,1系列で示せる。ブロック内のレコード全ては項目 i の項目値はこの系列を上位に持っているため本方式ではこの系列を省略して格納することによりデータ圧縮を行う。

(1) 入力アルゴリズム

- (i) 全ての項目について、投入されたレコードを構成する項目値のビット系列に従って各周辺木を根から順に葉までたどる。
- (ii) 各項目について(i)で求めた経路上のノードのビットマップの論理和をとる。これによりブロックの集合を得る。
- (iii) 求めたブロック集合の積集合をとるとただ1つのブロックが定まる。

(iv) 求めたブロックが分割基準に達したときは分割を行う。これについては後述する。

(v) 求めたブロックにレコードを項目ごとに分割格納する。

以上の方法によると、 $A_i = a$ であるレコード全ては、周辺木 T_i をその根から a のビット系列に従ってビットマップをその論理和をとりながら葉に到着するまでたどっていき、結果が“1”であるビットに対応するブロックに格納されていることになる。

(2) 分割アルゴリズム

分割対象ブロック番号を p 、未使用のブロック番号を f とすると、選定された分割項目 d から分割は図 8 のように行えばよい。

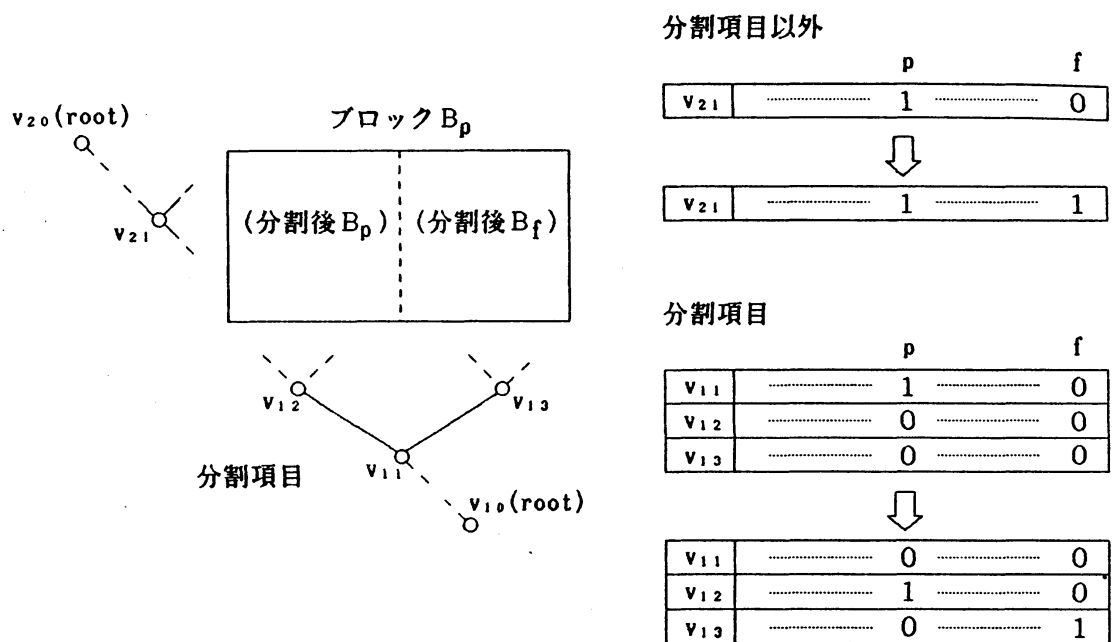


図 8. ブロック分割によるビットマップの変化

(3) 検索アルゴリズム

(i) 条件項目全てについて、その項目値に対応するノードを求め、ビットマップからブロックの集合を得る。

(ii) 求めたブロック集合の積集合をとる。

(iii) 得られたブロック全てに対し、条件を満たすレコードの参照項目を出力する。

本方式では分割は項目値のビット系列で行うが、分割を行う時期やどの項目の値に着目して分割するかは非常に重要な問題点となるので次に分割基準について述べる。

7.2 分割基準

分割はアクセスコストを減少させるのが目的であるがその方法を検討するため分割によるアクセスコストの変化を考える。このためにqueryを次のように仮定する。

(i) 検索条件は $Q(A_i = a)$ の形に限定する。ここで条件項目 i はランダムに等確率で選ぶとし、 a は D_i からランダムに等確率で選ぶ。

(ii) 参照項目 j は1つでランダムに等確率で選ぶ。

全ての i, j について検索コストを考え、これらの平均値を最も減少させる項目 d による分割を局所的最適分割とする。

いま分割を考えるブロック B_0 を分割対象ブロックと呼ぶことにし r 個のレコードがあるとする。分割は項目 d で行うこととし、分割を行うとレコードはブロック B_1 に r_1 個、ブロック B_2 に r_2 個と分かれるとする。

アクセスコストは検索対象となるバケット全てについて、条件項目に対す

るアクセスコストと参照項目に対するアクセスコストの和をとったものである。ある項目が条件項目になる確率、また参照項目になる確率は共に $1/n$ で、ある項目 i のある項目値が条件として選ばれる確率は $1/\{n \cdot \text{card}(D_i)\}$ であるから平均コストは、

$$\sum_i \sum_{a \in D_i} \sum_j \text{ACOST} / n^2 \text{card}(D_i)$$

である。ここでACOSTは条件 $Q(A_i = a)$ に対し参照項目 j の値が要求されたときの全ブロックに対するアクセスコストである。

分割により変化するのは分割対象ブロック分だけであり、これは

$$(1/n^2) \sum_i \sum_{a \in D_i} \sum_j \text{COST} / \text{card}(D_i) \dots \textcircled{1}$$

である。ここでCOSTは条件 $Q(A_i = a)$ に対し参照項目 j の項目値が要求されたときの、分割対象ブロックの項目 j の順次ファイルを構成するバケット数であるが、これは i, j, d の選びかたによって様々である。

詳細は省略するが(分割後の①式) - (分割前の①式)つまり分割による①式の増加は、 l が大、 n が大なる仮定を置き近似すると、

$$(1/n^2) \left[\sum_{i \neq d} P_i \left\{ \sum_j I(j) + (n-1) I(i) \right\} - P_d \left\{ \sum_j L(j) + (n-1) L(d) \right\} \right] \dots \textcircled{2}$$

と表せる。ここで $L(i)$ は分割前の項目 i の項目値格納のためのバケット数、 $I(i)$ は分割による項目 i 格納のためのバケット数の増加である。②式の[]内の第1項は分割を行うことによるアクセスコストの増加すなわちデメリット、第2項はメリットをあらわしている。

そこで本方式では、レコードの格納のために新たなバケットが必要となるとき、

$$\begin{aligned} \text{Div}(d) = & \sum_{i \neq d} P_i \left\{ \sum_j I(j) + (n-1) I(i) \right\} \\ & - P_d \left\{ \sum_j L(j) + (n-1) L(d) \right\} \dots \textcircled{3} \end{aligned}$$

を計算し、負となる d が存在するとき ③ 式を最小にする d で分割するという分割法を採用する。これにより平均コストを最小にする準局所的最適分割を与える。

8. あとがき

統計データの特徴を分析し、その一連検索に適したファイル編成としてブロック分割による多次元ファイル編成法についてその基礎概念および概要を述べた。まず静的分割法についての評価法として確率を用いた方法を示し、Transposed file と比較して有効である分割法が存在することを示した。これにより分割の有効性が示せた。また動的分割法は対象分野のデータに応じた分割を動的に行える有効な手法であり、その分割法、分割基準とその基準の理論的根拠を示した。本ファイル編成法は大量データを対象とし、多数の項目のうち少数項目に着目した一連検索が多用される分野、例えば文献検索システム、実験値計測システムなどに有効であると考えられる。動的分割法の実用性評価のためのシミュレーションによる他の編成法との比較などが今後の課題として残されている。

参 考 文 献

- [1] D.S.Batory "On Searching Transposed Files" ACM Transaction on Database Systems, Vol.4, No.4, pp.531-544(1979).

- [2] Susan J.Eggers,Arie Shoshani "Efficient Access of Compressed Data" A LBL Perspective on Statistical Database Management, pp.179-189 (1982).
- [3] 行政管理庁行政管理局統計主幹 "統計調査総覧" 全国統計協会連合会 (1976).
- [4] 近藤貴士 "統計データベースにおける一連検索用ファイル編成法について" 信学AL研資, AL85-68 (1986).
- [5] W.C.Lin,R.C.T.Lee,H.C.Du "Common Properties of Some Multi-attribute File Systems" IEEE Transaction on Software Eng. SE-5,2 pp.160-174 (1979).
- [6] J.Martin著,国友他訳 "データベース" 日本コンピュータ協会(1983).
- [7] 小菊他 "リレーショナルデータベースを用いた統計検索機能と実現方式" 情処DB研資42-2(1984).
- [8] 大沢裕,坂内正夫 "良好な動特性をもつ多次元点データ管理構造の一提案" 信学論Vol.J66-D, No.10,pp.1193-1200(1983).
- [9] 坂内正夫,大沢裕 "画像データベースにおけるデータ表現・管理方式" 信学論(D), J68-D, 4, pp.434-441 (1985).
- [10] 総理府統計局 "第三十二回日本統計年鑑"(1982).
- [11] 田中議 "複数属性ファイルの適応型最適ファイル分割法" コンピュータソフトウェア, Vol.2, No.2, pp.15-28 (1985).
- [12] 特定研究「多目的統合統計データバンクの開発」総括班事務局
"STATISTICAL DATA BANK" Vol.1, No.1-4 (1983~1984).

[13] M.J.Turner,R.Hammond,P.Cotton "A DBMS for Large Statistical
Databases" VLDB '79, pp.319-327.

[14] 打浪,手塚 "統計用DBMSの構成に関する一考察" 信学総全大(1984).

[15] 上田尚一 "統計データの見方・使い方" 朝倉書店(1981).