

IMS・DBの特質と設計上のポイント

中田 康雄 (三菱レイヨン)

1. テーラベース利用上の利点

- (1) 当社では'75年4月よりIBMのIMS/360を全面的に活用したオンラインシステムの運用を開始した。IMS/360の利点は大きくDC利点とDB利点よりなるが、小論ではこのうちのDB利点を中心にその特質、構造、パフォーマンスおよび設計上のポイントについて述べる。展開の前提としてまずここには、DB利点を利用することによって達成された効果の面から考えたDBの利点を、経験に基づいて次の2点にあるとしてその詳細に触れておこう。

DB利用上の利点

- ・ダイレクト・アクセス
- ・重複データ保持の回避

(2) ダイレクト・アクセス

論理レコードの集合を階層木構造のDBとして構成することによって、対象とするレコードのダイレクト・アクセスが効果的に行なえることになる。ダイレクト・アクセスの利点としては、ある1件のtransactionを対象とするDP処理を迅速に行なうことが可能になるという点で、On line real timeの不可欠の前提となることの他に、DB処理上のLogicの簡易化をもたらすと謂ふ利点を挙げることもできる。これは例えば、在庫マスターが在庫場所別にソートされていて、在庫場所〈A〉から〈B〉への転送を1件のInvoiceとしてのtransactionによって認識しなければならぬとするアプリケーションを考えれば、通常のsequential処理では、例えばこのtransactionを2件に分解し他のtransactionと共に「マスター」と同一のKeyにてSortingしたのち「在庫マスター」とMatchingすると謂ふ如き工夫をしなければならぬ。これに対し在庫マスターをDBとして保持して置けばこうした類の工夫は不要となる。即ちダイレクト・アクセスはrandomに発生するtransactionを予めsortingすることなしに直接master fileとmatchingするうえで、transaction1件あたりの処理スピードを上げるだけでなく、matchingのLogicを簡易化する意味に於てDP設計上の効率化をもたらす。場合によっては総合的な効率化を結果せしめる武器ともなりうる。この意味で、DBの利用は従来のmaster file中心のDP処理からの解放という「自由を手にする」ことになるであろう。

(3) 重複データ保持の回避

ここでの重複データは2つの意味を持つ。1つはある論理レコードの集合を考えたときその集合内での重複である。例えば販売取引という論理レコードを考えると、これは契約情報と出荷情報とのサブ集合が存在し、そこで1件の契約に対して出荷が複数に分割される時、1件の出荷レコードに対して常に同一内容の契約情報を結合してレコード化すると、出荷の件数だけ重複して契約情報が保持されることになる。これに対してDBを利用すれば、契

納と出荷とを階層木構造の中で位置づけることによりデータの重複を回避できる。(図-1参照)

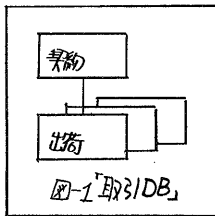


図-1「取引DB」

他方通常の Sequential 処理による DP 処理はアプリケーションの目的に応じて File を作成することによるデータの重複が避け難いことがある。例えば品名別の在庫とその受払を構成する取引を File 化するばあに、図-1の取引 DB に属する情報をこの「品名別受払 File」に取り込むことが必要になる。しかし FAS DB を利用することにより、「取引 DB」と「在庫 DB」なる2つの DB を構成し、その両方を向

かの手段を結合すれば、データの重複なしに新たに「品名別受払情報」を利用することが可能になる。(図-2参照)

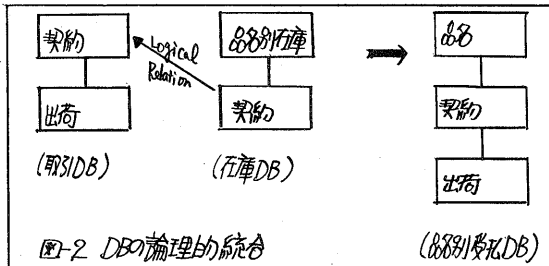


図-2 DBの論理的統合

(品名別受払DB)

この DB の論理的関係付けの手段は大きく IMS が用意している論理 DB の機能を利用することと、更に USER が独自にこの論理関係をアプリケーションによって maintenance することの2つである。前者は論理関係を維持する Direct Pointer 又は Symbolic

pointer を利用することにより、后者は例えば図-2の「在庫DB」の「契約セグメント」に取引 DB の「契約セグメント」の Key 部分だけを持ってこの関係付けをユーザーが管理することにより可能になる。当然のこと FAS のこともまた DB のダイレクト・アクセスの利点を利用することによりはじめが可能になる。

2. IMS・DBの特徴と構造

(1) 序

DB 利用上の実感的利点は特にそのダイレクト・アクセス性能に負うところが大きい。IMS・DB はそのアクセスメソッドとして HFSAM, HISSAM, HIODAM, HDAM を準備しているが、ダイレクト・アクセスについては HDAM のアクセスメソッドが優れている。ここには、HDAM についてその特徴と構造を述べることでダイレクト・アクセス性能がどのようなものであるかを見る。

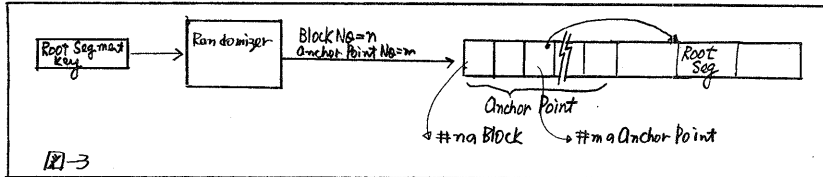
(2) HDAM DB の特徴

- DB 上の論理データベース・レコードはセグメントを上から下に、左から右にポインターを用いて順次展開することと、その階層的関連づけがなされる。このポインターは次の種類がある。

Pointerの種類	Application	Space	Retrieve Speed
Hierarchical Pointer	Sequential 処理に向く	節約できる	大
Direct Pointer └ Physical Child Physical Twin	Direct 処理に向く	節約が難しい	小

Direct Pointer を利用すると、Dependent Segment の検索が Direct に行なえるようになり、この意味で全体的に検索スピードが小さくて済むことになる。

- D. Root Segment の Addressing は ユーザーが準備するサブルーチン (Randomizer) により行われる。Address は BLOCK NO と ANCHOR POINT NO とにより決められる。Anchor Point は各 Block 内に保持され Block NO の Sub No 的な意味を持ち、この Anchor Point の中に、そこに Addressing された Root Segment が Store される Address が格納される。従って、Root Segment の Addressable Point は Block NO X Anchor Point NO の範囲内にある。以上の関係を図示すると図-3 のようになる。



1. Data Set は大きく Root Segment Addressable Area (R.A.A.) と Overflow Area (OVFL) とに分けられる。R.A.A. の大きさは上に述べた Root Segment の Addressing される Block NO の最大値によって規定される。
2. HDAM を利用する場合は USER は DBD に 'RMNAME = ' の後ろに - 9 によって次の情報を与えなければならない。
 - (1) Randomizer の Module Name
 - (2) R.A.A. の Total Size (Block 数)
 - (3) Anchor Point の数
 - (4) Data Base Record のうち R.A.A. に連続して何 bytes を収容するか。

(3) HDAM の構造

1. Bit Map

DB の各 1 Block は常に Bit Map として利用される。bit map の各 bit は、それに対応する Block の利用状況を示している。bit map 内の bit の数を N とすると 1 bit map は N block の利用状況を示すことになる。data Set が N block を越えて block size を持つ場合は $N+1$ 番目の block に別の bit map が作られる。bit map の format は図-4 に見るとおりである。各 1 word はこの block が bit map であることを示す flag を収容している。この次に位置するのは Anchor Point であり、Anchor Point に続く Area が bit map である。bit map の各 bit が 1 であることは、その block にその DB の最大セグメントを収容する余地があることを示している。従って bit map の各 1 bit は bit map の block に対応する 1 常に 0 である。

2. Space Address / Available Length

Data Block の初めには 2 bytes の space address が与えられる。これによってこの block 内の初めの Space Area の始まる位置が示される。この Space Address に続く 2 bytes はその Space Area の大きさを示している。

3. Anchor Point

Available Length を示す Field に続く Area は $n \times 4$ bytes の Anchor Point Field である。各 block の anchor point はその block に randomizer を通じて Addressing される root segment のうちの、最小の Key を持つ root segment が収容される。

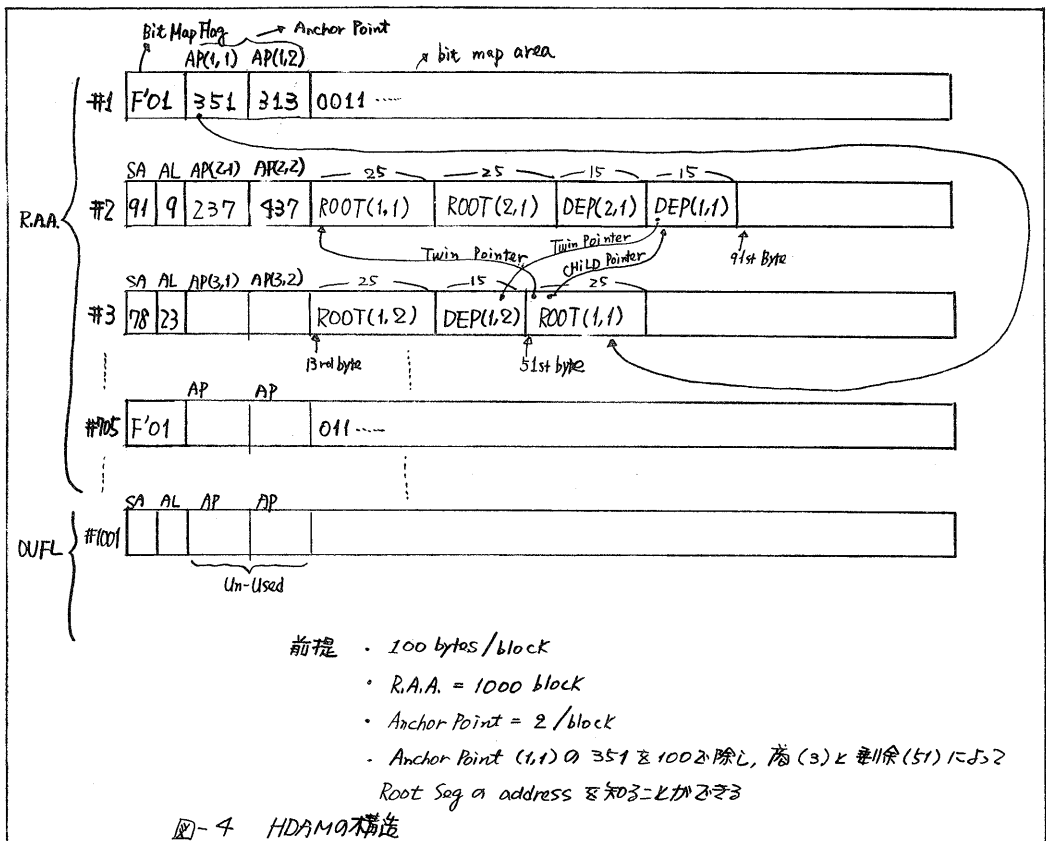
1) は address が保持される。即ち、複数の root segment が同一の block, 同一の anchor point の addressing されたとき (Synonym が発生), そのうちの最小の key を持つ root segment の address が anchor point に収容される。Synonym を有する複数の root segment は、この anchor point を親と見做し、physical twin pointer によって、ascending に chain される。

二. Free Space

ユーザは、Free Space としてある block の何%かを、又はいくつかの block 毎に1つの block の全体を確保することができる。これによって、DB の Load 時又は再編成時に、規則的に data Area の一部分を space Area として確保し、その後のセグメントの追加に於て、同一ADBLCコードが同一の Block 内に収容される確率を高めることが可能になる。この Free Space を規定する factor は次のように呼ばれる。

fbff : free block frequency factor

fspf : free space percentage factor



(4) Segment 追加 (Insertion) のルール

新しい Segment を Insert するにあたり、その Segment に関連する既に存在する Segment に最も近い空き Area に収容されるが、その場所を選択するルールは次のようになる。

1. 関連 Segment の収容できる block に Space があり, しかる New Segment を Insert してもその block の bit map が変わることはない (即ち Insert してのちにその DB の最大 Segment を収容する余地のある場合は), 次の優先順に条件を満たす Space に収容する。

- (1) New Segment に等しい大きさの Space
- (2) New Segment の 2 倍に等しい大きさの space
- (3) New Segment + 最小 Segment の大きさに等しい space
- (4) New Segment より大きい space

D. 関連 Segment の収容できる block に Space はあるが, New Segment を収容するとその bit map が変わる場合は, 次の優先順に条件を満たす Space を選択する

- (1) New Segment の 2 倍に等しい大きさの Space
- (2) New Segment + 最小 segment の大きさに等しい space
- (3) New Segment より大きい space

1. 関連 Segment の収容できる block と同一の track 内のいずれかの block で, 現在 buffer pool 中に存在する block

2. 関連 segment の収容できる block と同一の cylinder 内のいずれかの block で, 現在 buffer pool 中に存在する block

ホ. 最大の Segment Length を収容して 11 個の Track 内のいずれかの block

リ. 最大の Segment Length を収容して 11 個の cylinder 内のいずれかの block

ト. 関連 Segment を収容して 11 個の Cylinder を中心に $\pm n$ cylinder にあるいずれかの block で現在 buffer pool 中に存在する block

チ. 最大の segment length を収容して 11 個の cylinder を中心に $\pm n$ cylinder にあるいずれかの block

リ. Data Set の最後の block で buffer pool 中に存在する block.

ス. Data Set の最後の block のいずれかを bit map を参照して選択

3. IMS・DB のパフォーマンス

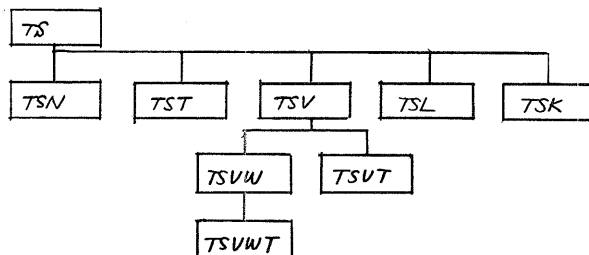
(1) IMS・DB を利用したアプリケーションを設計するうえで, DB 利用上の各機能に關するパフォーマンスを知ることは不可欠の前提であり, 当社はこれを次の条件の下に測定した。

1. Machine : IBM/370 M145 (700K bytes)

2. DISK UNIT: 3333-E

3. DB Buffer Size : 10KB

4. DB STRUCTURE



ホ. Segment Length / Occurrence / Key Length

Seg Name	S.L.	K.L.	Occurrence
TS	190	6	3000
TSN	30	6	1
TST	50	6	1
TSV	130	6	2
TSUV	170	6	2
TSUVWT	50	6	1
TSVT	50	6	1
TSL	15	6	1
TSK	15	6	1

(註)

TS の Key Field は

Initial Value = 000005

Increment = 000005

1. OS : VS 2.6

(2) 測定結果は次の通り.

表-1. CALL Function別測定値

CASE		1		2		3		4		5		6		7		8		
CALL Function		GU TS		GU TS/TSV		GN TSVW		GNPTSUV		GHNTSVW		REPL TSVW		ISRT TS		ISRT TS		
Factor		Total	AVE	Total	AVE	TOT	AVE	TOT	AVE	TOT	AVE	TOT	AVE	TOT	AVE	TOT	AVE	
HIDAM	CPU TIME	ms	1624	16	1628	16	6	3	7	3	6	3	17	9	6960	69	7271	72
	OS	ms	1239	12	1442	14	16	8	16	8	16	8	18	9	6028	60	10010	100
	IMS	ms	34	29	3094	31	24	12	23	12	24	12	37	18	13022	130	17316	173
	TOTAL	ms	4470		4690		30		30		20		40		34670		37280	
	ELAPSE	ms																
	CPU UTILIZATION	%	64		66		80						93		36		46	
	INDEX DB SEEK	B	60		60										326		601	
	DATA DB SEEK	B	100		100										663		448	
	CHANNEL	ms	813		901										4870		7211	
	HIDAM	CPU TIME	ms	1553	15	1536	15	7	3	8	4	6	3	16	8	2892	29	2815
OS		ms	1187	11	1397	14	16	8	16	8	16	8	18	9	3086	30	2958	29
IMS		ms	2775	27	35	29	23	12	25	12	24	12	35	18	6012	60	5808	58
TOTAL		ms	4420		4490		30		30		30		40		12180		11870	
ELAPSE		ms																
CPU UTILIZATION		%	56		59				84		80		89		49		49	
DATA DB SEEK		B	148		148										335		336	
CHANNEL		ms	519		519										1455		1213	
CALL回数		100		100		2		2		2		2		100		100		

(註) CASE1 : TSセグメントを initial value = 000005, increment = 000100 で 100回 Get.

CASE2 : TS,TSVをD'commandを利用し, CASE1と同様に100回 Get

CASE3~5: Get Uniqueで TS,TSVを RetrieveしてTS TSVを Unqualified SSAで Get.

CASE6 : GHNTSVWの時間を含む. 従って replace のためのTimeは Case6より Case5を3112

HIDAM CPU Total = 5ms

HDAM CPU Total = 6ms

CASE7 : TSの Key Fieldの initial value = 000006, increment value = 000005
に2 ISRT.

CASE8 : TSの Key Fieldの initial value = 040000, increment value = 000005
に2 ISRT.

表-2 DATA BASE OPEN TIME

		HIDAM	HDAM
C P U T I M E	OS	7.91	287
	IMS	83	59
	MPP	0	0
	TOTAL	875	347
D I S K S E E K C O U N T	DATA BASE	14	5
	SYSRES	60	20

(3) Comment

- I. Root Segment の Retrieve については、HDAM, HIDAM の間には、時間的に有意差は認められない。但しこれは、HIDAM の場合は Index DB の OVFL chain が殆んど存在しない前提においてのみである。
- D. Dependent Segment の Retrieve において、GN, GHN, GHNPA などの Function を用いても時間的に差はない。
- II. REPLACE は極めて小時間で行われる。但し Replace の場合は、Synchronization point において Disk 1 の Physical I/O が生ずるので、この時間を考えると全体的にはそれほど小さくはない。
- III. CASE 7, 8 から窺えるところ、Root Segment の追加に於ては、HIDAM と HDAM の差は極めて小さくなる。しかも HIDAM は Index DB の OVFL chain の増加によるパフォーマンス低下が著しい。
- IV. DB Open の時間は極めて大であり、このオーガの時間は real time 処理に於て致命的となる。

4. 設計上のポイント

(1) R.A.A. の Size の決定

- I. HDAM における R.A.A. のサイズを決定するにあたっての一般的基準は次式によって表わされる。

$$\frac{\text{R.A.A. 内の DBLコードサイズ (Bytes)} \times \text{DBLコード発生件数}}{\text{BLOCKサイズ (Bytes)}} = \text{R.A.A. サイズ (Blocks)} \quad \dots (1)$$

- (1) 式による R.A.A. のサイズは Free Space を考慮すると、いわば minimum なサイズであり、この場合の Total Size は次の式によって表わされる。

$$\text{minimum size} \times \frac{fbff}{fbff-1} \times \frac{1}{1 - \frac{fsPf}{100}} = \text{Total Size} \quad \dots (2)$$

- D. (1) 式を構成する各 Factor について更に詳細にみると次のようになる。

(a) DBLコード発生件数

アプリケーション毎にこの数値は既知を見做すことができる。

(b) DBLコードサイズ

このサイズは論理データ・ベース・レコード長の最大値を最高とし、ルート・セグメントの長さを最小の値とする範囲内で種々の選択が可能となる。Disk Space の Utilization を最大にする意味からは、最小の値に近い値を選択すべきであり、Hierarchical Chain を短くし、Retrieve の効率をあげるためには最大の値に近い値を選び、同一データ・ベース・レコード・ブロック内に

収容する確率を高める方がよい。DBロードサイズの決定にあたっては、このように、DISK SPACEの物理的制限、Root Segment Keyの特質、Randomizerの特質、アプリケーションの特性等の要因の相互連関を把握することが必須であるので、何度かの試行を要するが、一般的にはルートセグメントの長さとしてActiveなDependentセグメントの長さを加えたものに落着くであろう。理想的なRandomizer(変換後のAddressの分布がポアソン分布に近似する)を前提として、Disk Spaceを有効に利用することだけを制約条件とするなら、Root SegmentをKey情報のみに限定し、その他の情報は全て従属セグメントに付随構造を考えたうえで(HIDAMにおけるINDEX DBのHIDAM化と考えてもよい)DBロード長をRootセグメントの長さで等しくすると、工夫も一考に値するだろう。

(1) Block Size.

DB Buffer Poolの有効利用の点からほぼ外生的に決定される。即ち、DB Buffer PoolはIMSのコントロールの下で、全てのDBが共用することになり、このBuffer Pool内のDB Blockの収容にあたって、各々のDBのBlock Sizeが異なるとBuffer内に連続してSpaceを形成するためDataのMoveによるSpace Areaの連結(Data Compaction)が生ずる。この時間的ロスを防ぐため全てのDBのBlock Sizeは共通にすることが望ましい。当社は2KBの値を採用している。

(2) Anchor Point

Anchor Pointの数を増大することは、Randomizerによって得られるAddressing Pointの量を、R.A.A.の大きさ(即ちBlock数)を変えずに増やすことができる。従ってSynonymの発生を少なく抑えることが可能となる。しかし同一のSynonym chain内のRootを同一のBlockに収容してretrieveの効率をあげるには、1Block内のDBロード数との関連に於てAnchor Pointの数を選択すべきである。

$$1 \text{ Anchor Point ありのロード数の期待値 (B)} = \frac{1 \text{ DB ロード数 (G)}}{1 \text{ Block ありの Anchor Point 数 (E)} \times \text{R.A.A. の Block 数 (F)}} \quad (3)$$

(3) 再び,

$$1 \text{ Block ありのロード数の期待値 (A)} = \frac{(G)}{(F)} = (B) \times (E) \quad \dots (4)$$

一方,

$$(A)' = \frac{\text{Block Size (C)}}{\text{DB Record Size (D)}} \quad \dots (5)$$

よって(A)'を導くことができる。11は同一のSynonym chain内のRootを同一のBlock内に収容しようと考えれば(A)=(A)'となるはずである。更にRandomizerの特性によれば(B)はこれを平均値として11からかの偏差を生ずるので(A)=(A)' $\times\alpha$ ($\alpha>0$) となるはずである。こうして、(A)=(A)' $\times\alpha$ を前提とし、(C)、(D)、(G)を既知とすれば(F)が決まり、更に(E)を既知とすると(B)が決まることになる。

Anchor Pointの数はDisk Spaceの有効利用の面からは少なく抑えた方がよい。なぜなら、図-4によっても知られる如くAnchor Pointは利用できない余分なSpaceとしてOVL Areaにも設定される。従ってOVL Areaのサイズが太いときAnchor Pointの分だけ無駄が生ずる。

(2) Bit Map の考慮

Bit Map の利用が HDAM の複雑性を増大せしめてゐることは疑いえない。Randomizer が Bit Map Block の Relative Block No に Root Key を Addressing すると、Bit Map Block には DB レコードを収容することほゞきないから、そのレコードは別の Block に収容され、anchor point chain が Block を越えて延びることになる。従つて HDAM の Load 時又は Reorganization に際して、場合によつては全ての anchor point chain が Block を越えて延びる結果を招くことがある。これを防ぐには前頁の $(A) = (A') \times \alpha$ の α を大きめに設定して置き、Bit Map による anchor point chain の他 Block への延長がその Block の anchor point chain へ影響しないように考えておくことが必要となる。

(3) Formatting.

HDAM の Bit Map の設定及び Formatting は該当する領域の Data が Store される時に始めて行なわれる。Insert の効率を高めるには R.A.A. の全領域の Formatting を Load 時に行なう必要がある。このための Load 時に最高位のブロックへアドレッシングする Dummy Root を作り出す工夫をしておくべきである。

(4) Loading

HDAM の Load 時の効率をあげるには、DB レコードを Randomizer によつて予めの Address 変換し、この Address を Key とし Sort したのちに Load するとよい。

(5) DB OPEN

DB OPEN が頻発することによるパフォーマンスの低下は極めて著しい。これを防ぐには DMB Buffer を大きくすること、各 DB の DBD を小さくする（特に Field の定義を極小に抑える）こと、HDAM より HDAM を利用して DMB の数を減らすことが望ましい。

(6) DB Structure.

Performance Up の観点より DB Structure の望ましいあり方を考えるとき次のようになる。

1. Root Segment の Occurrence は大であらばあるほどよい
2. Dependent の Occurrence は小であらばあるほどよい。
3. 上記 2 を達成するためには Dependent の同一レベルのセグメント数を増大させることは厭うべきでない。
4. 上記 1, 2 を達成することができうるなら DB のセグメントレベルが大になることも厭うべきでない。

(7) Free Space

Free Space の Factor には fbf と fspta 2 つがあるが、このうち fbf は DB の Load および Reorganization の際に未使用の block を作り出す意味で上記 (2) の Bit Map と同様の理由によつて、anchor point chain が Block を越えて延びるという結果を招くばかりである。従つて、HDAM に於ては fbf の利用を控えるほうが賢明であろう。

(8) Twin Backward Pointer の利用

Deletion が頻繁に行なわれるセグメントには Twin Backward Pointer を付加すると効果的である。これは Deletion に際してその直前の兄弟セグメントの Twin Forward Pointer を付け替える時に、Backward Pointer が有効に機能するからである。

(9) Child Last Pointer の利用

Dependent Segment の Occurrence が大で、しかもこの Segment が Key Field を持たない場合は、この Dependent Segment の Parent に Child Last Pointer を持たせておくことで Dependent の追加にあたって極めて大きな効果をもたらす。