

(1982.1.20)

データフロー制御データストリーム処理方式 データベースマシン (基本構想)

田 中 謙
北海道大学・工学部

1. はじめに

1970年代には種々のアーキテクチャのデータベースマシンが盛んに研究されたが、80年代に入って、この分野は反省期を迎えているように思われる。主な反省点は、

1. 汎用データベースマシンは存在し得るか?
2. 大容量データベースをどう扱うか?
3. RASISをどう支援するのか?

の3点である。

第1点、ニーズの分析の結果として生じた疑問というよりも、処理対象データのデータ長の柔軟性と、ジョイン等の複雑な演算の高速処理とを同時に満たすようなマシンを開発することの難しさに根差した疑問である。利用者のニーズとしては、IRやバンキングシステムにもDBMSを用いたいとの動きも出始めている。マシンの分化を考えると重要であるが、一方では、

1. 可変長データが扱え、
2. タプル数に 10^6 程度にもなる大規模な2つの関係のジョインを $O(n)$ で処理でき
3. ジョイン属性の値として、100文字以上の文字列が現われること許す

というようなマシン・アーキテクチャの研究が必要である。本論文ではそのようなマシンを提案する。

データベースマシンの大容量化に関する反省は、最下層メモリとしてCCDや磁気バブルメモリを用いるマシン・アーキテクチャに向けられたものである。本研究でも大前提として、

“最下層メモリには移動ヘッドディスクを用いる”

という条件を課す。これに伴い、関係のセグメンテーションも必要となる。

第2点、RASISの支援について考慮したマシン・アーキテクチャが皆無であったことについての反省である。これは、従来の研究の重点が、データベース処理全体の内の一部である関係代数演算の処理に置かれ過ぎたためである。本研究では、セグメント化された関係をオブジェクトと見做し、オブジェクト指向アーキテクチャで、他のデータベース処理モジュールを制御することを考えている。

本研究では、関係の各属性値をデータ・タイプ毎にコード化したコード化関係を用いて関係代数演算を行うことを提案し、この演算の高速処理のためのサブシステムと、符号化と復号化のそれぞれのためのサブシステムのアーキテクチャの概要を示す。

2. アーキテクチャの概要

2.1. 方式に関する基本的な考察

本研究の指針となつた基本的な考察を以下に列挙する。

a. 普通の移動ヘッドディスクを最下層メモリとして用いるという仮定のもとでは、セレクションの実行時間を、インデックス・ファイルを持ったソフトウェア後置型データベースマシンと、これを持たないハードウェア後置型データベースマシンとで比較し、いかに大差ないと予想される。

b. 更新処理に於て過大な負担とならないう限りは、データの冗長度を上げて

べり、処理の局所性や、処理系の単純化を図ることによって、総合的なコスト・パフォーマンスを上げるべきである。

c. 関係モデルでは、どの属性も一様に扱われるが、現実には、検索条件式中に現われることかなく、出力の際にのみ必要となる属性がある。このような read-only の属性を、関係代数演算の処理系に持ち込むことは、効率を下げるだけで、何等利点はない。

d. 性別のように、男と女の2値しか現れないような属性を持つ関係 $R(X, sex)$ は、 $R-male(X)$ と $R-female(X)$ のように2つの関係に分けて管理処理される。

e. データベースマシンとホストコンピュータの間、ホストコンピュータと入出力機器、およびネットワークとの間、入出力機器と利用者との間のそれぞれにチャネル容量の上限があり、これらも無視して、データベースマシンの高速化を図っても無駄である。データベースマシンの入出力のチャネル容量は、せいぜい1~10 MB/sec と見做してよいであろう。これは、検索結果の出力が5,000~10,000 タプル/sec 得られるスループットに対応する。

f. 検索結果の統計解析結果のみを出力する場合には、そのスループットの値はあてはまらない。このような統計解析用の検索処理のスループットは大きければ大きい程良い。

g. 属性値のデータ型を、数値と文字列、順序の定義の有無、出力の際にソートの指定を受けるか否か、不等号によるセリクシオン演算を受けるか否か、不等号によるジョインやリストリクシオンの対象となるか否か、四

則演算の定義の有無を分類すると表2-1 のようになる。id# というのは

		order	sort	0-selection	0-join	0-restriction	小 リ ス ト リ ク シ ョ ン
Number	number	✓	✓	✓	✓	✓	
	id #	✓	✓	✓			
String	date	✓	✓	✓	✓	✓	✓*
	word	✓	✓	✓			
	text						

* : "1982-01-31" + 1 = "1982-02-1"
"1982-Jan-31" + 1 = "1982-Feb-1"

表2-1. データ型の分類

識別番号のようなデータ型で、date は年月日や、時分秒等のデータ型を指し、text型は順序が定義されない。word は date 型と text 型以外の文字列を指す。date 型以外の文字列型は、不等号によるジョインやリストリクシオンの対象になることはない。date 型は、number 型へ変換することが出来る。表2-1より、number 型と date 型、id# 型と word 型が同じ性質を持ち、これらと text 型の3つの型に分類することも出来ることかわかる。そこで、

A型 : number, date

B型 : id#, word

C型 : text

と分類する。

h. データベースマシンのアーキテクチャを考慮する上での大きな問題となるのが可変長データの取り扱いである。可変長データというのは、表2-1の word 型と text 型に現われる。word 型と text 型を、これらに対して定義されている順序を無視して、固定長のデータへと1対1に符号化できたと

すると、ジョインとリストリクションの処理は、符号化された世界で処理できることが表2-1よりわかる。プロジェクトの処理における重複タブルの除去も、符号化された世界で処理できる。word型とtext型に対しては、四則演算も定義されているから、このような型の属性について平均と総和の計算も要求されることはない。この型に対する統計演算は教之上げのみであり、これも符号化された世界で処理できる。復号化は、出力に対してのみ行えばよい。符号化された世界の関係代数演算処理は、固定長データの処理になり、このサブシステムのマシン・アーキテクチャを簡略化することができる。

2.2 符号化 / 復号化アーキテクチャ
 図2-1に、2.1節のよう述べた符号化関係の処理による要求処理の概念図を示す。図2-2には、処理の具体例を示す。

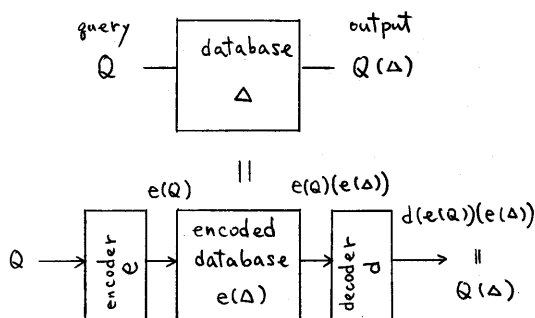


図2.1. 符号化データベースを用いた関係代数処理

示す。データベースの定義時に、各属性には固有のデータタイプが与えられるものとする。ジョインとリストリクションは、同じデータタイプを持つ属性の間の比較のみが許される。符号化は、各データタイプ毎に行われる。

(a) type

Name = word;
 Place = word;

attribute

Student: Name;
 Age: integer;
 Home-Town: Place;
 City: Place;
 Prefecture: Place;

(b) $Q = (R[Home-Town = City] S)[Age, Prefecture]$

R	Student	Age	Home-Town
	Nozaka	26	Kyoto
	Habata	24	Sapporo
	Masuyama	24	Shizunai

S	City	Prefecture
	Kyoto	Kyoto
	Sapporo	Hokkaido
	Shizunai	Hokkaido
	Tokyo	Tokyo

(c) encoder

Name	Use	Code
Habata	{R}	0
Masuyama	{R}	2
Nozaka	{R}	1

Place	Use	Code
Hokkaido	{S}	3
Kyoto	{R,S}	2
Sapporo	{R,S}	0
Shizunai	{R,S}	4
Tokyo	{S}	1

(d) encoded database

e(R)	Student	Age	Home-Town
	1	26	2
	0	24	0
	2	24	4

e(S)	City	Prefecture
	2	2
	0	3
	4	3
	1	1

(e)

Age	Prefecture
26	2
24	3

(f) decoder

Code	Name	Use-Count
0	Habata	1
1	Nozaka	1
2	Masuyama	1

Code	Place	Use-Count
0	Sapporo	2
1	Tokyo	2
2	Kyoto	3
3	Hokkaido	2
4	Shizunai	2

(g)

Age	Prefecture
26	Kyoto
24	Hokkaido

図2.2. 符号化データベースの処理例

符号化の行われるのは、桁数の大きい id 型、date 型、word 型、text 型である。date 型は、適当な関数を用いて 1 対 1 に number 型へと変換できる。桁数の大きい id 型は、word 型として扱う。word 型と text 型は変換テーブルを用いるが、word 型は、変換テーブル中の word 型データは辞書式順序を保って管理される必要がある。したがって word 型は B-tree や、text 型は hash テーブルを管理する。図 2.2 (c) に見るように、word 型であっても符号化データは、符号化の前のデータ間の順序を保つ必要はない。0-ジョインや 0-リストリクションを行うことがないからである。R[Name > 'Ikeda'] のような 0-セレクションは、(c) の変換表より、Masuyama, Nozaka が条件に該当することを知り $(e(R)[Name = 2]) \vee (e(R)[Name = 1])$ を計算すればよい。

符号化モジュールと復号化モジュールでの処理は、基本的には、変換テーブルを関係と見做してセレクションの処理を行うことである。2.1 節の A の考察から、これら 2 つのモジュールは、ソフトウェア復置型データベースマシンを多重化することによって実現するのがよいと考える。多重化は符号化テーブル(図 2.2 (c))と復号化テーブル(図 2.2 (f))のセグメント化を意味する。符号器はデータタイプに定義されている順序やハッシュ値の順序でセグメント化するのが望ましい。これに対し、復号器は、符号の値の順序でセグメント化されることが望ましい。図 2.1 で更新処理を考慮すると、図 2.3 のよう

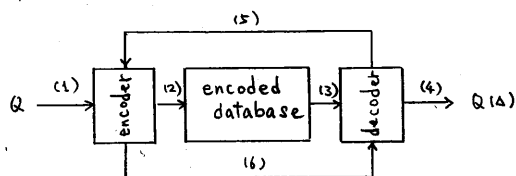
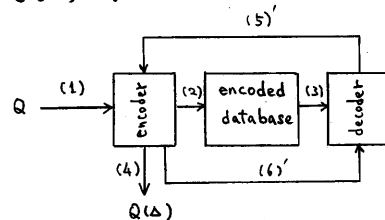


図 2.3 変換テーブルの更新

にデータバスの追加が必要である。(5)は、検索条件を満足するタプルを削除するような場合に必要である。そのようなタプルが符号化され (3) より得られると、各属性の符号値が復号器で調べられ、その符号値をデータバス中の何箇所かで用いているかを記憶している use-count (図 2.2 (f)) が 1 だけ減らされる。その結果、どこでも用いられなくなった符号値は、そのデータタイプ未使用の符号値を管理している場所へと戻され、符号器からも消去される。(6) のデータバスは、新しく関係を登録するような場合に必要である。

図 2.3 の構成では、(5) と (6) には生データが転送され、しかも符号器と復号器のセグメンテーションの仕方が異なっているために (5)、(6) のネットワークのスループットを上げることが難しい。そこでこれを図 2.4 のように改善する。復号器には、生データを持



encoder

P.A. : physical address

P.A.	Name	Use	Code
A1	Habata	{R}	0
A2	Masuyama	{R}	2
A3	Nozaka	{R}	1

P.A.	Place	Use	Code
A'1	Hokkaido	{S}	3
A'2	Kyoto	{R,S}	2
A'3	Sapporo	{R,S}	0
A'4	Sizunai	{R,S}	4
A'5	Tokyo	{S}	1

decoder

Code	Name	P.A.	U.C
0	A1	1	
1	A3	1	
2	A2	1	

Code	Place	P.A.	U.C
0	A'3	2	
1	A'5	2	
2	A'2	3	
3	A'1	2	
4	A'4	2	

図 2.4 符号化 / 復号化方式

つのだけでなく、生データが符号器や管理されている物理アドレス（物理アドレス番号）を持つ。これにより、(5) (6) のネットワークを流れるデータも固定長になり、変換テーブルの記憶に要する記憶領域を小さくもできる。この場合も、符号器と復号器へのセグメンテーション方式は異なるので、(5) (6) には、後述する動的クラスタリングを行う必要がある。

2.3 動的クラスタリングとソート

2.2 節の符号器/復号器を例にとり、説明する。両者は物理アドレス (P.A.) と符号 (C) の関係も情報として持っているが、符号器では、P.A. の順に並べられており、復号器では C の順に並べられている。符号器では P.A. の順に処理することが望ましく、符号器より復号器へ送られるデータは P.A. の順になる。一方、復号器では C の順に処理することが望ましい。したがって、符号器から復号器への転送では、C についてソートする必要がある。符号器と復号器がそれぞれ単一モジュールであればこれによれば、各々が複製モジュールより成る場合には、符号器の各モジュールからの出力を一つにまとめソートし、これを C について各復号器に分配する方法は多数のモジュールに対して転送路が一本となり、隘路を生じる (図 2.5(a))。逆に、C の値を先に分配し、各復号モジュールに転送されてからソートする場合には、分配の際にデータ転送の衝突を生じる (図 2.5(b))。

これに対し、ソートと分配をネットワークの機能として統合した SD ネットワークを提案する。これはソータや、着信等の開発したソートエンジン (SOE) を基本モジュールとする。分配は D-bus を用いて各符号モジュール毎に行われ、M-bus によって、各復号モジュール毎

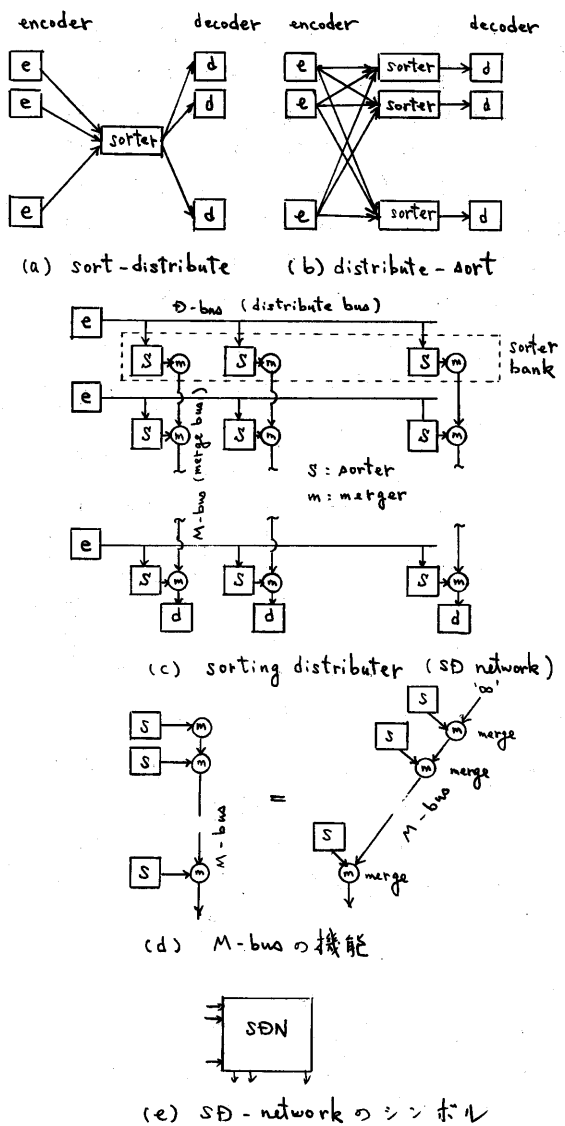


図 2.5 ソートと分配の方法

に、各列のソータの出力をマージして復号モジュールへと転送する。SD ネットワークはバッファメモリあり、大容量を必要とするが、後述のように大容量化は通常の RAM を用いて行う。各符号モジュールより一度に送出されるデータ量は、セグメントの大きさや上限が決まる。よって、図の (c) に示すソータバンク毎に要する RAM の大きさは一定である。よって、ソータの容量を増

するために使用するRAMバンクは、ソー
タバンク毎に分割することができる。

データ長が固定の場合、2つの処理
システム間のデータ転送において、間
にリータを介在させることが望ましい
場合、各システムが多重化された場合
にはSDネットワークを用いることが
できる(図2.6(a))。ただし、各デー

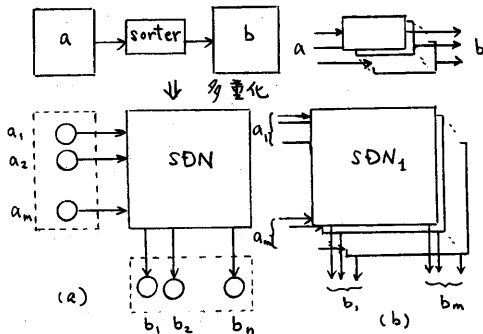


図2.6 リータを介した転送の多重化

の転送先は、モジュラスを計算する等
によって、データより計算可能とする。
さらに、転送ジョブの多重化は、同図
(b)のように多重度だけSDネットワ
ークを用いることにより実現できる。

2.4 大容量リータ

本節では、読者がソートエンジン
(SOE)について既に知っていること
と仮定する。(1)

SOEはヒープ構造を用い、図2.7(a)
のような状態やキーを入力すると、
同図の次経路上のデータ列にこの経路
上へ順序正しく新しいキーが挿入さ
れる。SOEが図(c)のような時、最下位
レベルの(i+1)番目のノードと根ノ
ードを結成経路に次経路を固定し、ヒ
ープの順序と同順序(図では昇順)にソ
ートされたキー列を入力した後、あ
ふれたキーをあふれた順に上から下
へ(i+1)番目のノードの下に並べか
けると、このアンバランスな木におい
て根から各葉へ至るキー列はソート

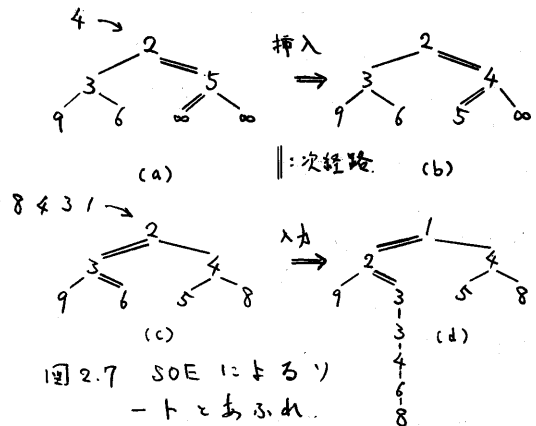
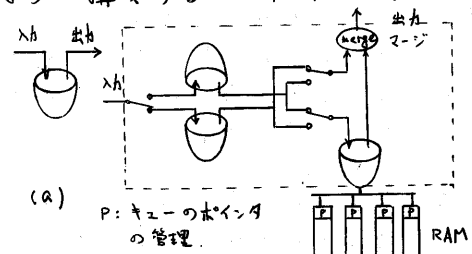


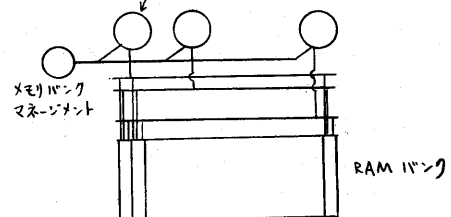
図2.7 SOEによるソ
ートとあふれ

されている。このあふれの部分をRAM
に記憶し、SOEの出力時に(i+1)番目
のノードに空孔が生じたとき、あふれ
部の上から順に、SOEへ送り出すと、
あふれ部と一緒にリータ出力する
ことができる。昇順にソートされたデ
ータ毎にiを变化させれば、ヒープ
の最下端のすべてのノードの下にあふ
れ部を設けることができる。今SOEを
図2.8(a)のように表わすと、以上の
ことを用いて、大容量リータを(b)の
ように構成することができる。RAMは



(b)大容量リータ

(b)図の破線で囲まれた部分



(c) RAMバンクの共用

図2.8 大容量リータとリータバンク

(1) Y. Tanaka et al. IFIP Congress 80 pp. 629-632.

多バンク構成でもよいが、同時に複数バンクをアクセスすることはないので、1バンクでもよい。多バンク構成の場合には、各バンクの割り当ても管理することにより、多バンクのRAMモジュールを複数のソータイ共用し、ソータイバンクを構成することが出来る(図2.8(c))。

2.5. 処理方式

符号化データベースの処理は、

1. タプルセレクタ (TS)
2. タプルメカ (TM)
3. アグリゲートオペレータ (AO)

の3つのモジュールで行う。TS, TM, AOは図2.9のように結合されている。

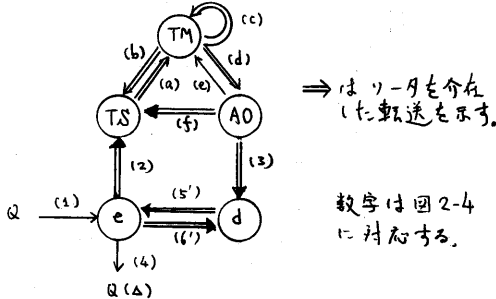


図2.9. 抽象アーキテクチャ

関係 $R(A_1, A_2, \dots, A_n)$ の各タプルにはタプル番号を定義することが出来るので、これを $@R$ と表し、 R があたかも $@R$ という属性をもつかのように見做し、 $R(@R, A_1, A_2, \dots, A_n)$ と表すことにする。 R は符号化されているとする。

TSは、各関係毎に、その関係の read-only な属性毎に、 $R(@R, A_1, \dots, A_n)$ の A_i と $@R$ へのプロジェクション $R[A_i, @R]$ を A_i についてソートしたファイル $R[A_i, @R]$ を持つ。TSが多重プロセッサ化される場合には、各 $R[A_i, @R]$ 型のファイルは、 A_i が number 型や date 型の場合には、 A_i の値をレンジ分けし、レンジ値のモジュラスの値でプロセッサに分配する。それ以外、 A_i の値の

モジュラスでプロセッサに分配する。プロセッサ内ではさらに $R[A_i, @R]$ を分割する場合は、 A_i の値の順序で分割する。各プロセッサは、図2.10に示すような構成で、プロセッサ間は同図(b)のように結合される。TSでは、セレ

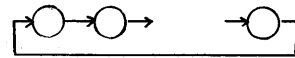
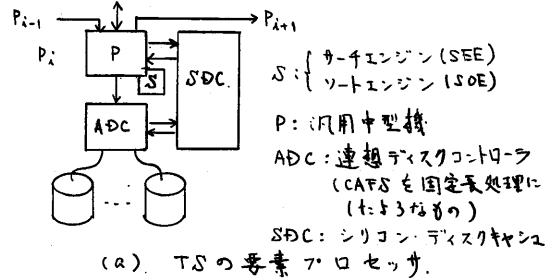


図2.10. タプルセレクタの要素プロセッサ

クションを行い、該当するタプルのタプル番号をTMへ送ったり、等号によるジョインを各プロセッサ毎に独立に処理し、ジョインされるタプル対のタプル番号対の集合をTMに送る。セグメンテーションと符号化の仕方により、等号ジョインはプロセッサ毎に独立に処理することが出来る。0-ジョインは、図2.10(b)の結合ネットワークを用いて処理する。TSからTMへの転送はソータイを介して行われる。多重プロセッサの場合はSDネットワークを介する。

$R[A_i = B_j] S[B_k = C_k] T$ のような演算や TS や $((R[A_i, @R]) [A_i = B_j] (S[B_j, @S]) [@S, @R])$ と、同様の $(@S, @T)$ の対が求められ、両方とも $@S$ につきソートするようにSDネットワークに転送することにより、この3関係のジョインが残る $(@S, @R, @T)$ の3項組を、TMの各プロセッサ中に求めることが出来る。

TMは、各関係 R をタプル番号のモジュラスで要素プロセッサに分割した

ファイルも、タプル番号の順序でセグメント化して構っている。各要素プロセッサは、図2.10の(a)と同様であるが同図(b)の結合ネットワークは持たない。

$R(A_i)$, $S(B_i)$ に対して、

$$((R(A_i = 'v')) [A_j = B_k] S) [A_l, B_m]$$

のような処理では、符号器EやVがコード化され、 $((R(A_i = 'v')) [A_j = B_k] S)$ がTMへ送られ、TMでは、この集合をXとして、各プロセッサで、 $Y = ((R(A_i = 'v')) [A_j = B_k] S)$ をサーチエンジンを用いて求め、これを A_j でソートしてTSに送る。TSの各プロセッサでは、 $Z = ((R(A_i = 'v')) [A_j = B_k] S)$ を求め、 OR をソートしてTMに送る。TMでは、各プロセッサ毎に独立に、 $((R(A_i = 'v')) [A_j = B_k] S)$ を求め、これをWとして、転送経路(c)を用いて、 OS でソートしてTMに送る。各プロセッサは独立に $((S(B_i = 'v')) [A_l = B_m] W)$ $[A_l, B_m]$ をサーチエンジンを用いて処理する。結果は、出力指定に従って、 A_l か B_m のいずれか、指定がない場合は任意に決める。この属性についてソートしてAOに送る。

AOはグループ化して統計演算や、重複タプルの除去などを行う。これは各プロセッサ毎に独立の多重プロセッシングで処理できる。AOの要素プロセッサは、サーチエンジンとソートエンジンを持った汎用コンピュータである。

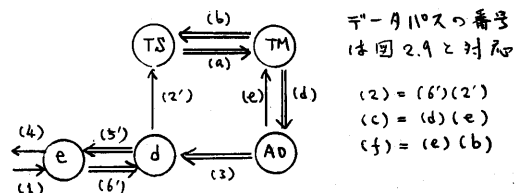
(e), (f) のデータバスは、結果として得られた関係の登録や、更新処理に用いる。

結果として得られた関係の各タプルには、AOの各プロセッサ毎に、プロセッサ番号とタプルの順番からなるタプルidがつけられ、各属性毎にバラバラにし、上の例では、 A_l とタプルid、 B_m とタプルidの対を A_l , B_m の値でソート

して、復号器dに送る。復号器では、各符号をこれに対応する生データの入っている物理アドレスに変換し、符号器に転送し、符号器では、各属性毎に、生データとタプルidの対が得られる。符号器から、これらの対をタプルidの順序で1個ずつ取り出し、タプルを構成することにより結果が得られる。

符号器と復号器も図2.10(a)のよう要素プロセッサで構成され、この2つと、TS, TMの各プロセッサでは、検索要求にしたがって、セグメント・ファイルの先読みがなされ、シリコンメモリで構成される大容量(100~500MB)のディスク・キャッシュにステージングされるものとする。先読みと、ステージングがレディーになることをうける処理のアクティヴェートは、データフロー制御で行う。処理の単位は、1セグメントであり、これはオブジェクトと見做して管理と処理がなされる。

図2.9の2重矢印の部分にはSDネットワークが用いられるが、SDネットワークを少なくするために、図2.9のかわりに図2.11のようを考える。



← は i 番目のプロセッサを j 番目に結ぶバス結合が、(4), (2) は Σ バスを表す。

図2.11. 改良した抽象アーキテクチャ

3. おわりに、

本マシンは、対象データベースの規模に加えて、各モジュール毎のプロセッサの多重度、SDネットワークの多重度も自由に設定できる。制御方式の詳細は別稿に譲る。

謝辞 通産省第5世代コンピュータDDG委員の方々の議論に感謝します