

TM-0178

関係型知識ベースにおける
単一化結合について

横田治夫(富士通)
森田幸伯, 西田健次, 伊藤英則

June, 1986

©1986, ICOT

ICOT

Mita Kokusai Bldg. 21F
4-28 Mita 1-Chome
Minato-ku Tokyo 108 Japan

(03) 456-3191-5
Telex ICOT J32964

Institute for New Generation Computer Technology

関係型知識ベースにおける単一化結合について

森田 幸伯, 横田 治夫*, 西田 健次, 伊藤 英則

(財) 新世代コンピュータ技術開発機構

* (株) 富士通

要約

関係型知識ベース (Relational Knowledge Base) に基づく知識ベースマシンにおける RBU 演算、特に単一化結合に対する処理方式およびその性質について述べる。関係型知識ベースは、知識ベースに対するひとつのモデルであり、知識は項関係 (term-relation) の形で格納・管理される。項関係のなかの項に対しては、単一化に基づく検索、RBU (Retrieval-By-Unification) 演算により検索を行う。関係型知識ベースは、知識ベースのモデルとして柔軟なモデルであるが、最も単純には単一化結合は、全ての可能なタプルのペアを生成し、そのタプルの指定された属性の値が単一化可能かどうかを調べることになり、その処理は、マシンにとって非常に重い負荷となる。本稿では、単一化結合について効率的な処理を行うための項の順序付けと、そのいくつかの性質について述べる。

U n i f i c a t i o n - J o i n O p e r a t i o n

O n a R e l a t i o n a l K n o w l e d g e B a s e

Yukihiro MORITA, Haruo YOKOTA*, Kenji NISHIDA, Hidenori ITOH

Institute for New Generation Computer Technology

*Fujitsu Laboratories Ltd.

This paper describes some properties of retrieval-by-unification (RBU) operations, especially unification-join, on a relational knowledge base, and the ordering of terms. The relational knowledge model is a conceptual model for a knowledge base. Knowledge is represented by term relations in this model. Terms in the relations are retrieved with operations called RBUs (i.e., unification-join and unification-restriction). To perform unification-join in the simplest manner, all possible pairs of tuples in term relations should be checked to see if each pair of terms in the tuple is unifiable or not. This would result in an extremely heavy processing load. The ordering of terms described here improves efficiency considerably.

1. はじめに

第5世代プロジェクトの中期の目的の一つとして知識ベースマシンの開発があげられる。知識情報処理を大規模なシステムで考えた場合、データ処理におけるデータベースのように、知識を効率的に管理・共有できるサブシステムが必要である。本稿では、そのようなサブシステムを効率よく実現するマシンを知識ベースマシンと呼ぶ。

知識ベースマシンは、さまざまなユーザ/ホストが利用するため、その概念構造は基本的に柔軟性に富んでいることが望ましい。[Yokota 86]で提案されている関係型知識ベースは、非常に柔軟な知識ベースマシンの概念モデルとなっている。知識を表現する基本要素としては、変数を含むある種の構造体である項(term)を用い、それらの管理の基本要素として関係(relation)を用いる。これらに対して関係代数を拡張したRBU: Retrieval-By-Unification演算を用意する。項関係で表現された知識は、ホーン節の集合や意味ネットワークモデルによる知識の集合などとみなすことができる。ある形式で格納した項関係は、ホーン節の集合とみなすことができ、その項関係に対して、RBU演算を行うことによりinput resolutionを行うことができることを[Yokota 86]は示した。[村上 86]では、意味ネットワーク的表現DCKR[小山 85]の一部を関係型知識ベースモデルの上で実現できることを示した。

しかし、RBU演算のうち単一化結合を単純に行うと2つの項関係のすべての可能な組合せに対して指定された属性が単一化可能かどうか調べなければならずマシン側の処理量が非常に大きくなる。

[森田 86]では、項の順序付けを用いた組合せの省略による単一化結合の効率化について提案している。

[大森 86]では、ハッシュとソートによる単一化結合の効率化を提案している。その方式ではあらかじめ木構造をきめその構造(ハッシュ木)内の定数の情報は可能な限り利用するのでより強く絞り込みがおこなえる。しかし、きめられた構造を用いてハッシングを行うため扱える構造に制限ができる。実際にどちらの方式が優れているかは、絞り込みの後の実際の単一化を処理する方式や全体の実現方式などに依存していると思われる。

本稿では、単一化結合の効率化のための項の順序付けと単一化結合のいくつかの性質について述べる。以下、2章で関係型知識ベースとRBU演算について数学的な定義を与える。3章では、RBU演算に関するいくつかの性質について述べる。

2. 関係型知識ベースモデル

データベースでは、ユーザ/アプリケーションの多様な要求に対して、関係モデル(Relational model)という1つのデータモデルのもつ単純な構造とな操作で柔軟に対応することを可能にした。

知識ベースにおいてもさまざまなユーザ/アプリケーションに対応できるひとつの共通なモデルが必要とおもわれる。関係型知識ベースは、知識ベースマシンに対する1つの概念モデルであり、ユーザのさまざまな知識表現に対応できる単純な構造と操作からなる。

この章では、関係型知識ベースモデルについて、以後の議論に必要な事項について述べる。

2.1. 基本概念

関係データモデルでは、領域(domain) D_i 上の関係(Relation)を扱う。つまり、関係 R は D_i ($i=1, n$) の直積空間の部分集合と見ることができる。しかし、第一正規型で D_i の要素をアトミック(atomic)なものに制限している。

$$R \subseteq D_1 \times D_2 \times \cdots \times D_n$$

それに対して、関係型知識ベースモデルでは、領域 K_i を項(term)の集合に拡張する。領域 K_i の直積空間の部分集合を項関係(term relation) T と呼ぶ。

$$T \subseteq K_1 \times K_2 \times \cdots \times K_n$$

つまり、我々は取り扱う対象をデータから変数を含む一種の構造体である項に拡張する。

変数の集合を Var 、関数記号の集合を Fun とし、 $Fun \cap Var = \emptyset$ とする。 $Fun \cup Var$ 上の項は、以下のように再帰的に定義される。

1. 定数 $a \in D$ は項である。
2. 変数 $x \in V$ は項である。
3. もし n 引数の関数記号 $f \in F$ があり、
 $t_1, \dots, t_n$ が項なら、
 $f(t_1, \dots, t_n)$ は項である。
4. 上の規則を用いて生成されるものだけが項である。

Term を $\text{Fun} \cup \text{Var}$ 上の項の集合とする。
 置換 $\theta: \text{Var} \rightarrow \text{Term}$ は以下のような変数と項の組の順序集合で表現される。

$$\{ \{ t_i / x_i \} \mid t_i \in \text{Term}, \\ x_i \in \text{Var} \text{ and} \\ \text{if } i \neq j \text{ then } x_i \neq x_j \}$$

置換 θ を項 t に適用した結果を $t\theta$ と書く。
 $t\theta$ を t のインスタンス(instance)という。
 項 t_1, t_2 に対し、 $t_1\theta = t_2\theta$ を満足する置換 θ が存在する時 t_1, t_2 は単一化可能であるといい、置換 θ を単一化作用素(unifier)という。

項 t_1, t_2 の単一化作用素 θ に対し、項 t_1, t_2 の任意の単一化作用素 θ' に対し、
 c 置換 θ'' が存在し、 $\theta' = \theta \cdot \theta''$ であるとき θ は項 t_1, t_2 の最汎単一化作用素(most general unifier)という。

項 t_1, t_2 の最汎単一化作用素を $\text{mgu}(t_1, t_2)$ と書く。

項の組の集合 $\{ (t_i, u_i) \mid i = 1, \dots, n \}$ に対し、すべての i に対し、 $t_i\theta = u_i\theta$ となる置換 θ をその項の組の同時単一化作用素(simultaneous unifier)という。

同時単一化作用素 θ に対し、任意の同時単一化作用素 θ' に対し、置換 θ'' が存在し、 $\theta' = \theta \cdot \theta''$ であるとき θ は最汎同時単一化作用素という。項の組の集合 $\{ (t_i, u_i) \mid i = 1, \dots, n \}$ の最汎同時単一化作用素を $\text{mgsu}(\{ (t_1, u_1), \dots, (t_n, u_n) \})$ と書く。

2. 2. RBU

関係データベースのデータ操作言語としては、関係論理(Relational calculus)と関係代数(relational algebra)があった。関係代数が手続き型のシステムであるのに対し、関係論理は非手続き型である。しかし、関係代数のほうがマシンの演算を考えるうえでは考え

やすい。そこで、我々は、関係代数を拡張した知識操作言語を考える。

関係データベースから関係型知識ベースへの拡張に伴い、従来の関係代数の結合(join)や制約(restriction)といった演算は、単一化(unification)に基づくものに拡張される。すなわち、結合・制約のなかの等号条件を単一化可能条件に拡張した単一化結合(unification-join)、単一化制約(unification-restriction)を考える。これらを、RBU(Retrieval-by-unification)演算と呼ぶ。射影(projection)については関係代数のそれと同じである。

タブル μ の属性名の集合 X の各要素の属性の値からなるタブルを $\mu[X]$ と書く。

単一化制約は、項関係のなかから与えられた条件式を満足するタブルの集合をインスタンスの形で求めるものである。条件式は、以下のようなものとする。

$$F = \bigvee_{k=1, \dots, m} (\bigwedge_{j=1, \dots, m_k} (A_{i,j} \circ w_{j,k}))$$

ただし、各 A は属性名を w は属性名または項を示す。 w に対して w が属性名ならタブル μ に対するその属性の値を、 w が項なら w 自身を返す関数 take_term を考える。つまり、

$$\text{take_term}(w, \mu) = \begin{cases} \mu(w) & w \text{ は属性名} \\ w & w \text{ は項} \end{cases}$$

$\Lambda \Diamond w$ でタブル μ に対して $\mu[\Lambda]$ と $\text{take_term}(w, \mu)$ が単一化可能であるという条件を示す。 $\wedge, \vee$ はそれぞれ接合(conjunction)選言(disjunction)を意味する。

条件式 F をもつ項関係 T に対する単一化制約の結果関係 ($\sigma_F T$ と書く) は以下になる。

$$\sigma_F T = \{ \mu\theta \mid \exists \mu \in T, \exists k, \\ \theta = \text{mgsu}(\{ \{ \mu[A_{i,j}], \text{take_term}(w_{j,k}, \mu) \} \mid j = 1, \dots, m_k \}) \}$$

X, Y, Z, W をそれぞれ属性名の集合とする。単一化結合は、2つの項関係のタブルの相定された属性の値が単一化可能なタブルのインスタンスの対を合成したタブル集合を求

$\text{ancestor}(X, Y) \vee \neg \text{parent}(X, Y)$
 $\text{ancestor}(X, Y) \vee \neg \text{parent}(X, Z) \vee \neg \text{ancestor}(Z, Y)$
 $\text{parent}(\text{smith}, \text{clark})$
 $\text{parent}(\text{clark}, \text{turner})$
 $\vdots$

KB1	$\text{ancestor}(X, Y) S$	$\text{parent}(X, Y) S$
	$\text{ancestor}(X, Y) S$	$\text{parent}(X, Z), \text{ancestor}(Z, Y) S$
	$\text{parent}(\text{smith}, \text{clark}) S$	S
	$\text{parent}(\text{clark}, \text{turner}) S$	S
	$\vdots$	$\vdots$

KB2	$\text{ancestor}(\text{smith}, Y)$	$\text{parent}(\text{smith}, Y)$
	$\text{ancestor}(\text{smith}, Y)$	$\text{parent}(\text{smith}, Z), \text{ancestor}(Z, Y)$

KB3	$\text{ancestor}(\text{smith}, \text{clark})$	$\text{parent}(\text{smith}, \text{clark})$	
	$\text{ancestor}(\text{smith}, Y)$	$\text{parent}(\text{smith}, \text{clark}), \text{ancestor}(\text{clark}, Y)$	$\text{ancestor}(\text{clark}, Y)$

図1. 項関係とRBU演算の例

めるものである。適当な関数記号を導入することにより $\mu[Y]$ は Y が属性名の集合の場合でも一つの項とみなせることに注意しよう。

項関係 $T_1(X, Y)$ と $T_2(Z, W)$ の $Y \diamond Z$ の条件での単一化結合の結果関係は $X \cup Y \cup W$ 上の関係で以下になる。

$$\begin{aligned}
 T_1(X, Y) \underset{Y \diamond Z}{\bowtie} T_2(Z, W) \\
 = \{ \mu \mid \exists \mu_1 \in T_1, \exists \mu_2 \in T_2, \\
 \theta = \text{mgu}(\mu_1[Y], \mu_2[Z]), \\
 \mu[X] = \mu_1[X]\theta, \\
 \mu[Y] = \mu_1[Y]\theta = \mu_2[Z]\theta, \\
 \mu[W] = \mu_2[W]\theta \}.
 \end{aligned}$$

但し、 $X \cap W = \emptyset$ とし、必要なら属性名は適当に付け替えるものとする。

図1は、項関係の例を示している。

図の上の方に示されたホーン節を表わす項関係がKB1 (head, body), である。KB1に $\text{head} \diamond \{\text{ancestor}(\text{smith}, W)\}$ という条件で単一化制約を行った結果がKB2 (head, body) である。さらにKB2とKB1を $\text{kb2.body} \diamond \text{kb1.head}$ で単一化結合した結果がKB3である。

3. RBU演算と項の順序付け

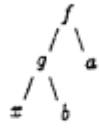
関係モデルは、ユーザに対して柔軟なデータモデルを与えるが、マシンの処理量が大きくなる。特に、大きな関係どうしの結合演算の処理は、非常に大きな計算量を必要とする。そのためDeitman (Kakuta 85)では、関係演算専用のハードウェア・関係代数エンジン (RE) を用いて効率化を計っている。REではハードウェアソータを用いて結合(join)などの処理時間の低減をはかっている。

関係型知識ベースモデルにおいても、単一化結合の処理はマシンに対して大きな負担となる可能性がある。単一化結合の効率化が必要である。そのためのRBU演算と項の順序付けに関するいくつかの性質について述べる。

3.1. 項の順序付け

我々は、非常に大抵な知識が知識ベースマシンに格納されることを仮定している。その格納場所としては2次記憶を想定している。最も単純に考えれば、単一化結合は2つの項関係 (term-relation) のタプル (tuple) のすべての組合せに対して条件を満足するかどうかを調べればよい。しかし、すべての組合せを生成することは、非常に大きな処理負荷がかかる。

そこで関係のタブルを何等かの順序で順序付けることにより組合せを省略したい。



Family order: (f2) (g2) (x) (b0) (a0)
Level order: (f2) (g2) (a0) (x) (b0)

図2. 木の線形化表現

{Yokota 86}では項を順序づけるためにジェネラリティ (generality) の概念を定義した。 t_1, t_2 を項とすると、 t_1 が t_2 のインスタンスであるとき t_1 は t_2 よりジェネラルであるという。すなわち、

$t_1 \sqsupset t_2$ iff $\exists \theta \quad t_2 \theta = t_1$
(但し、 θ は置換 (substitution))

ジェネラリティによる順序付けは全順序ではない。そこで、ジェネラリティを保存しながら項に対して全順序となる順序付けを考えたい。そのためには、項が木で表現されることに注意し (図2)、その木を線形化して、文字列として表わし、その文字列上の順序を用いれば良い。木の線形化としては、例えば {Knuth 85} によれば家族順 (Family order) のものと、レベル順 (Level order) のものがある。本稿では、家族順のものを leftmost 方式の順序と呼び、レベル順のものを outermost 方式と呼ぶ。m 方式による項 t の線形化 (文字列) 表現を $rep_m(t)$ と書く。

これらの文字列表現では、各文字は $Var \cup Fun$ の要素に対応している。文字 c の対応する $Var \cup Fun$ の元を $node(c)$ と書く。項を順序付けるためにその文字列表現 $rep_m(t)$ に対して図3に示す辞書的順序を用いる。ただし、図では以下のような記号を用いている。

文字列 s の i 文字目から j 文字目までを $s[i:j]$ で表わす。ただし、 $i > j$ のときは空の文字列を表わす。また、 $length(s)$ で文字列 s の長さを表わす。文字列 s の中に変数 x に対応する文字がはじめてでてくる位置を $posv(s, x)$ で表わす。文字列 s の中に変数 x が現れない場合は、

$s_1 > s_2$ if and only if $n = difpos(s_1, s_2)$, and

1. $node(s_1[n;n]), node(s_2[n;n]) \in Var$,

(a) $posv(s_1, node(s_1[n;n])) > posv(s_2, node(s_2[n;n]))$.

(b) $posv(s_1, node(s_1[n;n])) = posv(s_2, node(s_2[n;n])),$
 $s_1[n+1;length(s_1)] > s_2[n+1;length(s_2)]$.

2. $node(s_1[n;n]) \in Var, node(s_2[n;n]) \in Fun$.

3. $node(s_1[n;n]), node(s_2[n;n]) \in Fun$,

$node(s_1[n;n]) > node(s_2[n;n])$ in arbitrary order in Fun

図3. 項の文字列表現の辞書的順序

$posv(s, x) = length(s) + 1$ とする。

同様に、文字列 s の中に変数に対応する文字がはじめてでてくる位置を $posv(s)$ で表わし、現れなければ、 $posv(s) = length(s) + 1$ とする。さらに、 $s[1:posv(s)-1]$ を $prefv(s)$ と書く。さらに、文字列 s と s' がはじめて違う文字を持つ位置を示す $difpos(s, s')$ を以下のように定義する。

$$difpos(s_1, s_2) = \begin{cases} 1, & \text{if } s_1[1;1] \neq s_2[1;1] \text{ or,} \\ & s_1 \text{ or } s_2 \text{ is null-string;} \\ n, & \text{if } s_1[1;n-1] = s_2[1;n-1], \\ & s_1[n;n] \neq s_2[n;n] \ (n \geq 2). \end{cases}$$

この2つの方式 (left most, outer most) は、ジェネラリティによる順序を保存することが示せる。2つの方式では、あるノードに対してその子ノードを変化させてもその文字列表現におけるノードに相当する文字より前は変化しない。一般にそのような文字列化の方式 m の集合を M とし、以下 $m \in M$ である rep_m について考える。この場合、置換 θ は、項の変数のノードをある部分木で置き換えるので $k = posv(rep_m(t)) - 1$ とすると、 $rep_m(t)[1:k] = rep_m(t\theta)[1:k]$ である。

まず、 $m \in M$ のとき、 rep_m の辞書的順序を用いた項の順序付けは、ジェネラリティを保存することを示す。

定理1

$m \in M$, かつ $t \sqsupset t'$ ならば
 $rep_m(t) > rep_m(t')$

[証明]

$t \sqsupset t'$ よりある置換 θ が存在して、
 $t \theta = t'$

任意の置換 θ は、変数 1 個の置換の合成で表わされることから、変数 1 個に対する置換を $\theta_1 = \{t_1 / x_1\}$ とし、

$s = \text{rep}_m(t)$, $s' = \text{rep}_m(t \theta_1)$
 $k = \text{posv}(s, x_1)$

とすると、 $s > s'$ を示せばよい。

$m \in M$ より、 s と s' は、 $k-1$ 文字目までは一致している。従って、

$s[1; k-1] = s'[1; k-1]$

ここで、 $s[k; k] \in \text{Var}$ であるから

$s'[k; k] \in \text{Fun}$ なら $s > s'$ 、

$s'[k; k] \in \text{Var}$ の場合は、 t_1 が変数の場合であり $\text{posv}(s', t_1) \leq k$ である。

$\text{posv}(s', t_1) < k$ なら

図 3. 1(b)より $s > s'$ 。

$\text{posv}(s', t_1) = k$ なら

(1) $\text{posv}(s, t_1) = \text{length}(s) + 1$ の時は

$s[k; \text{length}(s)]$ と $s'[k; \text{length}(s')]$

は、 x に相当する文字が t_1 に相当する文字に置き換わっただけで、 $\text{posv}(x)$

$= \text{posv}(t_1) = k$ なので $s = s'$

(順序の意味で) である。

(2) $\text{posv}(s, t_1) < \text{length}(s) + 1$ の時は

$\text{posv}(s, t_1) = k$ とすると、

$s[1; k-1] = s'[1; k-1]$ であったから $k \leq k'$ 、

$k = k'$ は $x = t_1$ の時だから $s = s'$

$k < k'$ のときは、 s , s' の k' 文字目がどちらも t_1 で $\text{posv}(s, t_1) = k < k'$

$= \text{posv}(s', t_1)$ だから $s > s'$

結局、 $s > s'$ 故に任意の置換 θ に対しても
 $\text{rep}_m(t) \geq \text{rep}_m(t \theta)$ 。 □

系 1

left most 方式, outer most 方式はジェネラリティを保存する。

定理 2

2 つの項が t, t' に対し $m \in M$ とし、
 $s = \text{rep}_m(t)$, $s' = \text{rep}_m(t')$ とするとき

t, t' が単一化可能であるなら

$s[1; k-1] = s'[1; k-1]$ (1)

但し、 $k = \min(\text{posv}(s), \text{posv}(s'))$

INPUT: Two sets of terms T_1 and T_2 .

OUTPUT: All the possible pairs of terms.

Step 1: Set $fs \leftarrow \text{null-string}$, $class_i(k) \leftarrow \phi$
for $i = 1$ or 2 and $k = 0, 1, 2, \dots$

Step 2: Take a term $t \in T_i$, ($i = 1$ or 2),
such that $t \geq t'$ for all $t' \in T_1 \cup T_2$.
And let $T_i \leftarrow T_i - \{t\}$.
If there are no terms in $T_1 \cup T_2$, then stop.

Step 3: Set $p \leftarrow \text{prefv}(\text{rep}_m(t))$, and $n \leftarrow \text{difpos}(p, fs)$.

Step 4: Let $class_1(k) \leftarrow \phi$, and $class_2(k) \leftarrow \phi$ for $k \geq n$,
Let $fs \leftarrow p$.

Step 5: Output all pairs (t, t') ,
where $t' \in class_j(k)$, $1 \leq k \leq n-1$, $j \neq i$, ($j = 1$ or 2).

Step 6: Let $class_i(\text{length}(p)) \leftarrow class_i(\text{length}(p)) \cup \{t\}$
and go to step 2.

図 4. ペア生成アルゴリズム

である。

[証明]

t と t' が単一化可能であるとする、 θ が存在して

$t \theta = t' \theta$

$t'' = t \theta = t' \theta$ と置いて、 $s'' = \text{rep}_m$

(t'') とすると

$m \in M$, $k \leq \text{posv}(s)$, $k \leq \text{posv}(s')$ だから

$s[1; k-1] = s''[1; k-1] = s'[1; k-1]$ □

3. 2. 単一化結合における順序付けの効果

項の集合 T_1, T_2 に対し、

単一化結合の処理を削減するため単一化の可能性のあるペア (t_1, t_2) ($t_1 \in T_1$, $t_2 \in T_2$) を選び出すことを考える。

定理 2 より $t_1 \in T_1$, $t_2 \in T_2$ が単一化可能ならば、 $\text{rep}_m(t_1)$, $\text{rep}_m(t_2)$ について変数に相当する文字が現れるまで両者の前方の文字列は一致している。単一化の可能性のあるペアは、その文字列表現 $\text{rep}_m(t_1)$, $\text{rep}_m(t_2)$ に対する比較操作だけで選びだすことができる。この選び出す操作は図 4 のアルゴリズムにより容易に行うことができる。

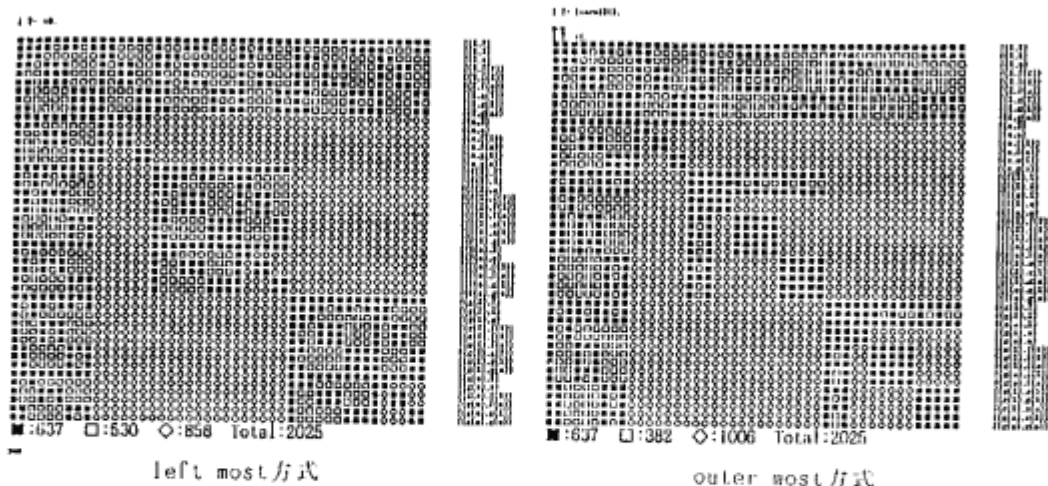


図5. 順序付けによる組合せの省略の例

二つの項の集合から順序付けにしたがって(辞書的順序で大きい順に)マージしながら項を一つずつ取出す。

項 t が入力されると「 s 」の値は $\text{pref}(\text{ref}_k(t))$ に変更され、どちらの項集合 (T_i) に属していたか及び文字列表現したときの最初に現れる変数の位置 k により $\text{class}_i(k)$ に入れられる。

同時に「 $s[1:k-1] \neq \text{ref}_k(t)$ 」($1:k-1$)となる $\text{class}_i(k)$ を空集合する。残りのすべての class_j の内容としては定理2の(1)を満足する。従ってStep5でこれらのうちもう一つの項集合に属していたものとのペアが出力される。

一方、 t が入力されたときそれより以前に入力された t' に対してで定理2の(1)を満足するものがこの時点でどの class にも入っていないということはない。

もし、どの class にも含まれていないとすると、 t' の属していた class を空集合にした項 t'' が t' 以降 t 以前に入力されていたはずである。順序付けに従って順に入力していることから $s' = \text{ref}_k(t')$, $s = \text{ref}_k(t)$, $s'' = \text{ref}_k(t'')$, $k = \text{posv}(t)$ とおくと

$$s < s'' < s'.$$

t'' が t' の属していた class を空集合にしたことから、ある $k' < k$ があって

$$\text{prefv}(\text{ref}_k(t''))[k'; k']$$

$$\neq \text{prefv}(s)[k'; k']. \quad (2)$$

ところが、 t と t' が定理2の(1)を満足し、

$s < s'$ であるから、

$$s[1:k-1] = s'[1:k-1]$$

$k = \text{posv}(t)$ だからこれらはすべて定数に相当する文字のはずだから(2)に反する。

よって、それより以前に入力された項のうち定理2の(1)を満足するものは、どれかの class の中にある。

但し、その場合の削減される量は、順序付けにより図5のようになる。

図5は、縦軸、横軸に単一化結合を行う2つの項関係のなかの項を順にならべたものである。最も単純にはこの全平面の相当する項のペアについて単一化可能であるかを調べなければならない。図4のアルゴリズムでは定理2の(1)の条件を調べることにより単一化可能かどうかの判定の回数を減らす。図5はトップレベルが2引数関数1種 (f)、その引数として1引数関数2種 (f または g) および定数1種 (a) からなる項の集合の場合である。

図中■は、単一化可能な項のペアを、◇は条件(1)を満足しないペアを、□は条件(1)は満足するが単一化不可能な項のペアを示す。

項を木で表現した場合、変数は、その葉にしか現れない。従って線形化表現ではレベル優先のものの方が深さ優先のものに比べて変数が後ろに現れやすい。本方式では $\text{posv}(\text{ref}_k(t))$ だけを調べるのでレベル優先(outer most)の方が絞り込みが強い。

3. 3. 単一化制約と単一化結合の関係

関係代数である種の制約が結合演算によって行う事ができたように、2. 2で述べた条件式を持つ単一化制約は単一化結合で行う事ができる。

例えば、項関係 $T(A_1, A_2, A_3, A_4)$ に対する単一化制約で、条件式として $\Gamma = (A_1 \diamond t_1 \wedge A_2 \diamond t_2) \vee (A_3 \diamond t_2)$ を持つ単一化制約の場合、条件を表現する関係として、

$$\begin{aligned} T'(\Lambda_1, \Lambda_2, \Lambda_3, \Lambda_4) = \\ \{(t_1, x_2, x_2, x_4), \\ (y_1, y_2, t_2, y_4)\}. \end{aligned}$$

を用意し、お互いの属性 $(\Lambda_1, \Lambda_2, \Lambda_3, \Lambda_4)$ が単一化可能という条件で結合を行えばよい。

つまり

$$\sigma_{\Gamma} T = T \bowtie T'$$

但し t_1, t_2 は項で、 x_2, x_4, y_1, y_2, y_4 は変数である。

4. おわりに

関係型知識ベースモデルにおけるRBU演算および項の順序付けについて、いくつかの性質を述べた。また、項の順序付けにより単一化結合を効率的に行うことができることを示した。

現在は、同アルゴリズムおよびそれを用いてRBU演算専用装置(エンジン)

[Morita 86]についてシミュレーション等により評価中である。

一方、関係型知識ベースモデルそれ自身の研究やそれを用いた知識ベース管理システムの構築についても現在研究が進められている。

[謝辞]

本検討を進めるにあたり、有益なご示唆を戴いた知識ベースマシンWGのメンバーの方々に感謝致します。

[参考文献]

- [Kakuta 85] Kakuta, T., et al., "The Design and Implementation of Relational Database Machine Delta", Proceedings of the International Workshop on Database Machines'85., March 1985.
- [Knuth 73] Knuth, D. E., The Art of Computer Programming, Vol. 1, Fundamental Algorithms, second edition, Addison-Wesley, 1973.
- [小山 85] 小山, 田中, "Definite Clause Knowledge Representation", Logic Programming Conference 85, July 1985.
- [村上 86] 村上, 他, 「単一化検索言語による知識ベースソフトウェアの記述」, 情報処理学会第32回全国大会, 1M-9, 1986.
- [Murakami 85] Murakami, M., Yokota, H., Itoh, H., "Formal Semantics of a Relational Knowledge Base", ICOT TR-149., Dec. 1985.
- [森田 86] 森田, 他, 「単一化結合の処理方式」, 情報処理学会第32回全国大会, 1M-7, 1986.
- [Morita 86] Morita, Y., et al., "Retrieval-By-Unification Operation on a Relational Knowledge Base", To appear in Proceedings of the 12th International Conference on VLDB, August 1986.
- [大森 86] 大森, 他, 「推論機能と関係データベースの融合-ハッシュとソートによる検索-」, 情報処理学会第32回全国大会, 1M-6, 1986.
- [Yokota 86] Yokota, H., and Itoh, H., "A Model and Architecture for a Relational Knowledge Base", Proceedings of the 13th International Symposium on Computer Architecture, pp. 2-9, June 1986.