

TM-0151

情報処理学会第32回(昭和61年前期)

全国大会発表論文集

35 件

January, 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

データフロー方式並列推論マシン
—実験機の動特性—

31.

^{*} 六沢 一昭 大原 輝彦 来住 晶介 ^{**} 伊藤 徳義
(沖 電 気) (I C O T)

1. はじめに

我々は、データフロー方式並列推論マシン実験機 (PIM-D) を試作し、現在評価を行っている。本稿では、PIM-D の動特性について述べる。

2. PIM-D の構成

PIM-D は、複数台の演算エレメント (PE)、構造メモリモジュール (SM)、及び1台のサービスプロセッサ (SVP)、及びこれらを接続するネットワークからなる (図1)。

PEは、データフローグラフを解釈／実行する。PE は、トークンのバッファリングを行う PQU (Packet Queue Unit)、オペランドの待ち合わせを行う ICU (Instruction Control Unit)、発火した命令を実行する APU (Atomic Processing Unit)、及びローカルメモリ LM からなる(図2)。

SM は、構造データを格納する共有メモリである。SVP は、PIN-D の初期化、保守、及び計測データの収集を行う。

3. PIM-D の動特性

評価プログラムとして 7-queens を用いて測定した動特性を以下に示す。

3-1. 発火制御

ICU では、オペランドメモリの検索／書込みにハッシングを用いている。また、ハッシュ衝突の起きた場合にそなえ、各エントリにチェインポインタを設け、再ハッシュ(rehash)によりエントリを連鎖する方式をとっている。例を図3に示す。図中、searchとあるのはチェイン探索のことである。また、↓の数が「チェインの長さ」である。

測定データを表1に示す。表の値からハッシュ衝突の確立を求めると

$$\text{再ハッシュ回数} / (2 \text{ オペランド命令トークン数} / 2) = 0, 1, 8$$

となり、およそ20%（以下）である。オペランドメモリ使用率が最大で12%であるので、まだハッシュ関数に改良の余地がある。

確立を20%として、衝突を起こしたトークン1つあたりのチェーン探索回数を計算すると

$$\text{チェーン探索回数} / (2 \text{ オペランド命令トークン数} / 2 \times 100\%) = 0.74$$

となり、1回以下である。これより、チェインの長さ l は $0 \sim 1$ であるといえる。

3-2. 榨働率

ICU 及び APU の稼働率 (idle 率) を表 2、図 4 に示す。表中、待時間の 2 行目、3 行目はそれぞれ総クロック数に対する割合、1 回のバケット入出力あるいは T-BUS アクセスあたりの時間数である。図 4 では稼働率 (idle 率) を 1 バケットあたりのマイクロステップ数 (= クロック数: 1 ステップは 1 クロックで実行) で示した。

PE 1～3 台の範囲ではidle率（待時間）は一定である

稼働率は低い (APU 51%, iCU 35%)。これは、APU がバケット出力あるいは LM アクセスを行う際、T-BUS のバス権獲得に非常に長い時間がかかっているためである (6 クロック / 1 アクセス、総クロックの 35%)。図 2 からわかるように、T-BUS には多くのユニットが接続しており、これがネックになっている。

表1. 発火制御に関する計測データ

受信データ	714Kト-ク	再ハッシュ回数	42K回
2オペランド命令	474Kト-ク	チェーン探索回数	35K回
1オペランド命令	203Kト-ク	トークンキユ	(PQU:容量16Kト-ク)使用率
特殊パッケージ	37Kト-ク	最大長	1.6Kト-ク(10%)
送出トークン数	477Kト-ク	最大メモリ数	(容量32Kト-ク)使用率
			3.7Kト-ク(12%)

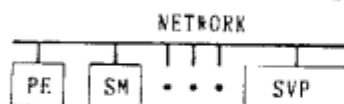


図1. PIM-D の構成

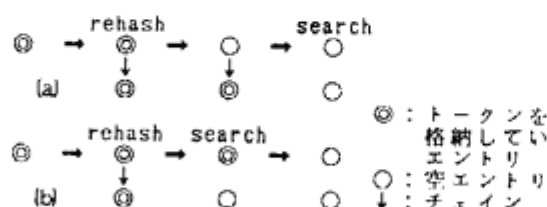


図3. チェイン操作 - チェインの長さ0~1
(a), (b)ともrehash1回、search1回)

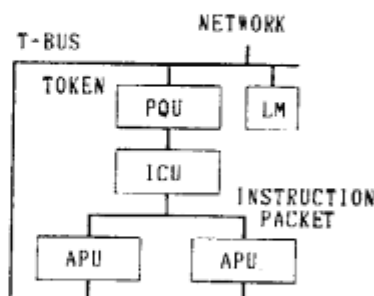


図 2. PE の構成

3-3. 実行命令解析

OR 並列及び GHC のプログラムを実行させた時の主な命令の実行回数などを表3に示す（実行された命令の総数は、OR並列が21種類440K回、GHC が20種類661K回）。表中、D換算マイクロステップは、1回の T-BUSアクセスに6ステップ（クロック）かかるとして換算した値（ $B \times C \times 5$ ）。Eの（）内は、総マイクロ実行時間（T-BUS アクセス待時間を含む）に対する割合。Gは、マイクロ実行における T-BUSアクセスに費やす時間の割合（ $C \times 6 / D$ ）。

OR並列、GHC とも実行ステップ数の上位（Eの上位）3命令で総マイクロ実行時間の40～50%を、また表に示した命令（総種類数の1/4）だけでは2/3をしめている。これらは link を除いて T-BUSアクセスの割合が大きい（平均41% $=17/(24+17)$ ：図4より）。このこともネックが T-BUSにあることを示している。

4. まとめ

PIM-D の動特性について述べた。結果は T-BUSがネックになっていることを示している。IM は APUによってのみアクセスされ、そのアクセスが全 T-BUSアクセスの大きな比重をしめていることから、IM を T-BUSから切離し、別のバスを APUとの間に設けることで大きな改善が期待できる。

謝辞 ICOT 関係各位に感謝する。

【参考文献】

- 1)伊藤, "並列型 Prolog のデータフローマシンによる実行方式," 日経エレクトロニクス, 1984, 11, 5号.
- 2)来住 他, "データフロー方式並列推論マシン実験機の構成," 情報処理学会第30回全国大会 6C-2.
- 3)久野、大原 他, "データフロー方式並列推論マシン 一実験機の性能評価一, 一実験機のプロセッサ割り付け方式一," 本大会 2R-3, 4.

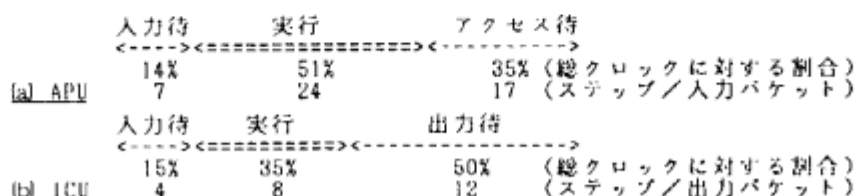


図4. APU, ICU の稼働率

表2. ICU, APU の idle率

PE 台数	総 クロック 数 (クロック)	A P U				I C U			
		バス 入力 待時間 (クロック)	入力 バス 回数	T-BUS バス 待時間 (クロック)	T-BUS バス 回数	バス 入力 待時間 (クロック)	入力 バス 回数	バス 出力 待時間 (クロック)	出力 バス 回数
1 台	11,409K	1,580K 14% 6.6	239K	4,160K 36% 6.2	668K	1,671K 15% 2.3	714K	5,956K 52% 12.5	477K
2 台	6,012K	806K 13% 6.5	124K	2,082K 35% 6.1	334K	878K 15% 2.5	356K	2,957K 49% 11.9	248K
3 台	4,044K	552K 14% 6.5	85K	1,341K 33% 6.0	223K	615K 15% 2.5	248K	1,975K 49% 11.7	169K

表3. 実行命令

	命 令	A 実行 回数	B マイ クロ ス テ ッ プ 数	C T-BUS バス 回数	D 換 算 マイ クロ ス テ ッ プ 数	E D * A	G T-BUS バス 割合
O R 並 列	call (成功)	10K	100	18	190	1,900K(19%)	57%
	copy	54K	15	3	30	1,620K(16%)	60%
	link	66K	20	0.5	22.5	1,485K(15%)	13%
	append-stream	7K	50	8.5	92.5	648K(7%)	55%
	return	5K	70	11.5	127.5	638K(6%)	54%
G H C	call (成功)	12K	100	18	190	2,280K(15%)	57%
	return	17K	70	11.5	127.5	2,168K(14%)	54%
	copy	58K	15	3	30	1,740K(11%)	60%
	switch-by-type	33K	20	4	40	1,320K(9%)	60%
	link	53K	20	0.5	22.5	1,193K(8%)	13%
	create-global	14K	35	7	70	980K(6%)	60%

データフロー方式並列推論マシン 実験機のプロセッサ割付け方式

2R-6

大塚 肇[◎] 久野 英治 桑佳 品介 伊藤 徳義

(神奈川総合システム研究所)

1. はじめに

PIM-Dは第5世代コンピュータ・プロジェクトの一環として研究、開発されたデータ・フロー方式の並列推論マシン実験機である。PIM-D実験機のアーキテクチャはデータ・フローグラフを解釈、実行する演算エレメント・モジュール(PE)と、共有メモリとして用いられ構造体データを保持する構造メモリ・モジュール(SM)、及びこれらを結合している階層バスからなっている。階層バスはこれらのモジュール間の通信用トークンの交換網となっている。また4PE、4SMを一つのクラスタとしてこのクラスタを同一の階層バスで多数結合できるようにしている。PIM-D実験機ではPEとSMとは機能が分割されていることにより両者の処理が並列に行なえ、かつプロセスの割付けと構造データの割付けが独立して容易に行なうことが可能となっており、それらの割付け戦略を検討することにより、リソースの有効利用、処理時間の高速化が期待できる。

筆者らは、現在PIM-D実験機の試作を終え評価を行っている。本稿では、OR並列型Prologで記述されたプログラムについて評価を行なったものを報告する。

2. OR並列型Prologの実行方式

Prologにおいては、1つのゴール・リテラルと、その述語に対する定義が与えられたとき、そのゴールと定義と定義に含まれるすべての節の間でユニフィケーションが試行される。通常の逐次型推論システムの処理系での節の呼出しは、プログラム・リスト上で現われた順に、ゴール・リテラルの呼出しは左から右に行う。ゴールが与えられたときその一番左のリテラルが選られ、その述語に対する定義中の最初の節との間でユニフィケーションをおこなう。これが失敗したとき、バックトラッキング機構による制御により次の節を呼び、ユニフィケーションの再試行をおこなう。これに対しOR並列Prolog方式ではこのユニフィケーションを並列に行なうことによって、推論の高速処理を可能にしようとするものである。すなわちゴールが与えられたとき対応する定義中の各節に関するユニフィケーション・プロセスを並列に生成することにより実現できる。ユニフィケーションが成功した場合、一般に出力される解の順序は一意的には定まらない。そこでPIM-D実験機にインプリメントされたOR並列型Prologにはストリームの概念が導入されており、OR並列で動作しているユニフィケーションプロセスと、その結果を使用するプロセスとの非同期通信手段をこれにより提供している。ストリームはnon_strictなデータ構造である。ここで言うno-

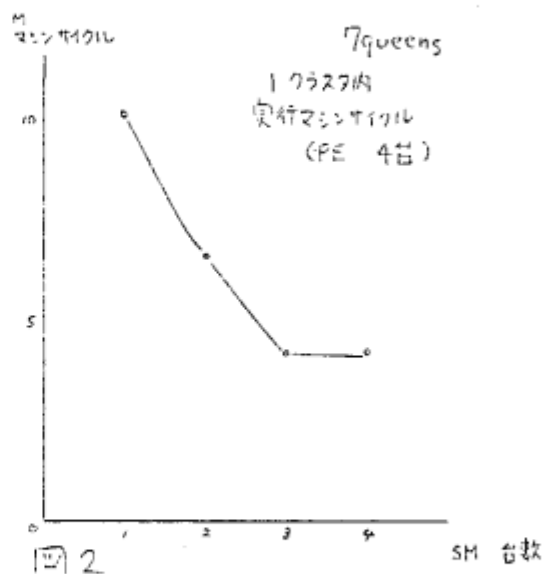
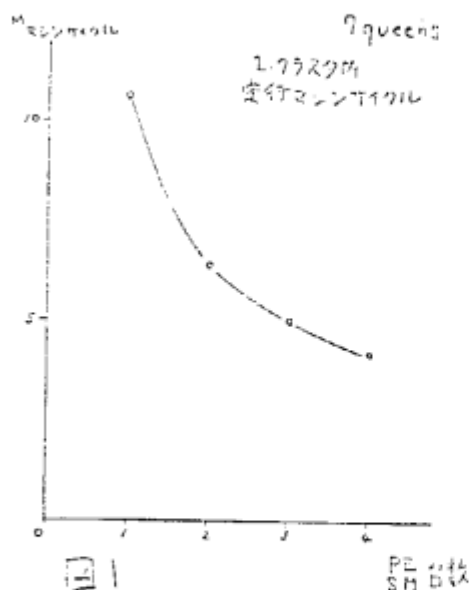
n_strictなデータ構造とは構造の要素が必ずしもすべて確定していない時点で生成されるような構造のことである。この構造は節の起動と同時にSM上に生成され、その要素は解が得られた順に書き込まれる。解を消費するゴール側はすべての解が生成されるのを待つことなくそれらを参照することが可能である。[1][2]

プロセッサの割付けはプログラムの実行前に静的に決定される。PE内のLMにはプロセスを実行するPEとそのプロセス管理テーブルのポイントが格納されている領域がある。PEがプロセスを起動する際にはこの領域をサーチして、実行モジュール、プロセス管理テーブル、及びプロセス識別子を決定している。(色付きトークン方式)したがってこのテーブル内のポイントを操作することによってプロセッサとプロセスの割付けを容易に制御できる。

3. 評価プログラムによる実験

評価用プログラムとして7 queensのプログラムを選んだ。これは探索型のプログラムとして典型的なものであり、並列度が高くプロセッサ割当効果がよく現われることが期待できると思われるからである[3]。現PIM-D実験機は2つのクラスタから構成されており評価実験は1クラスタの場合と2クラスタの場合について行った。

図1には1クラスタ内でPE、SMの台数をパラメータにとったときの実行マシンサイクル数が示されている。また図2にはPEを4台としSMの台数をパラメータにとったものが示されている。図1と図2を比較すると1クラスタ内ではPE、SMを共に同じ台数ずつ増加させたほうが実行マシンサイクルが減少することがわかる。SMが1台である場合は4PEからのアクセスが集中し1PE、1SMの場合と実行サイクルはあまり変わらないが、SMの台数を増加してゆくとそれが減少してゆくのがわかる。図3には2クラスタの場合の実行マシンサイクルを示した。ここではクラスタ間のプロセッサ割付け比を1:8から1:64まで変化させたものをプロットしてある。ここで言うプロセッサ割付け比とは、起動するプロセスをどのくらいの割合でクラスタ外のプロセッサに割り当てるかを表したものである。ここではプロセッサ割付け比を小さくすると実行サイクルが減少するがあまり小さくするとクラスタ間階層バスの通信量が大きくなり、このバスがデッドロックを起こした。これは今後の大きな課題である。



4. おわりに

OR並列型Prologで記述された7queensプログラムにより、プロセッサ割付け方式が実行マシンサイクルにどのように変化を与えるか実験してみた。実験機のSMにはPEのRPQに相当するバッファがないためSMへの通信(PE-SM間, SM-SM間)に際して交換網の開塞率が高くなりオーバーヘッドとなることがわかった。このため算術演算など解が唯一しか存在しないことがあらがじめわかっているような場合はストリームを利用しないことによりSMへのアクセスを減らすなどの工夫が必要であると言える。またハードウェアの点からはSMにバッファをもうける、あるいはSM間通信用のバスを設けること、などを検討する必要がある。

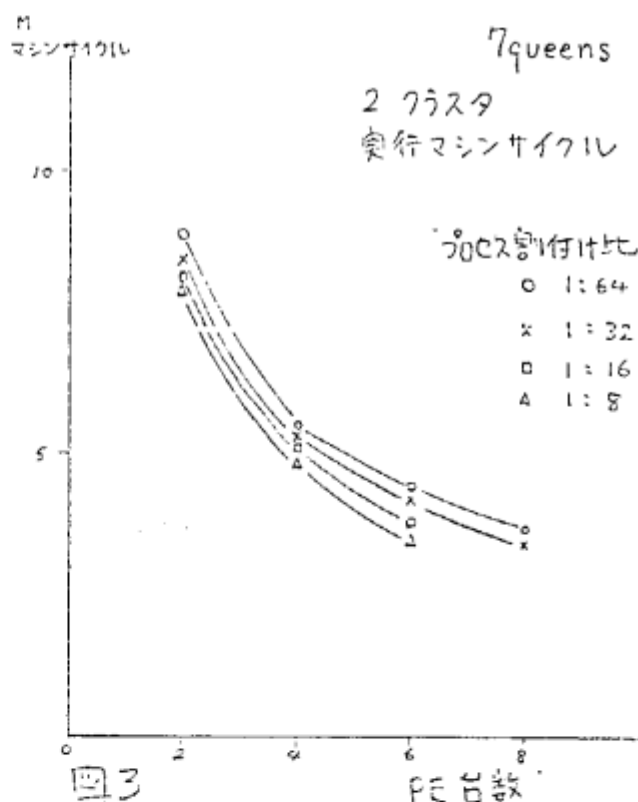
今後はさらにいろいろなプログラムを用いて詳細な評価をおこなっていく予定である。

謝辞

ICOT関係各位に感謝の意を表す。

【参考文献】

- 1) 伊藤 他, 「データ・フロー方式推論マシンのアーキテクチャ」
Proc. of Logic Programming Conf. ICOT Mar 1984
- 2) 伊藤, 「並列型Prologのデータ・フローマシンによる実行方式」
日経エレクトロニクス, 1984年11月5日号。
- 3) 久野 他, 「データ・フロー方式並列推論マシンにおけるAND/OR並列効果のシミュレーションによる測定」
情報処理第30回全国大会予稿集



久野 英治 大沢 一昭* 智 方之
(沖電気総合システム研究所)

伊藤 徳義**
(ICOT)

1. はじめに

PIM-D (Parallel Inference Machine based on Dataflow model) は第5世代コンピュータプロジェクトの一環として、ICOTと共同で開発している並列推論マシンの1つであり、データフローモデルをベースとしているところに特徴がある⁽¹⁾⁽²⁾。PIM-Dの目標言語は並列型論理プログラミング言語であり、現在、AND並列型Prologの代表言語であるGHC (Guarded Horn Clauses)⁽³⁾とOR並列型Prologの2つの異なる言語をサポートしている。筆者等は、PIM-D実験機を試作し、その性能評価を行ってきた。本稿では台数効果を中心としてその評価結果を報告する。

2. PIM-Dの概要

PIM-Dは、4台の処理要素PE (Processing Element) と4台の構造メモリSM (Structure Memory) を共通バスで結んだクラスタを構成単位とし、このクラスタを上位の階層バスで結んだ構造をしている⁽²⁾。

各PEは、データフローグラフで表現されたプログラムを並列に実行し、必要に応じて他PEへのプロセスの分散を行う⁽¹⁾。構造データは、SM群に分散して格納され構造データ操作が必要なとき、PEからSMへ操作指令が転送される。プロセスの割当てや構造データの割当ては、PEで制御される。このため、各PEは、全PEの空

きプロセス番号の管理表やSMの空きセルの管理表を持っており、プロセス起動命令や構造データ生成命令が実行されたときに、これらの管理表から新プロセス番号や構造データセルのアドレスを取出す。

現実実験機では、各PEのプロセス番号を前後のPE番号が昇順になるように格納している。SMの空きセルの管理表も同様である。従って、プロセスや構造データはPE群やSM群にほぼランダムに分散して割当てられる。この条件でPE (SM) 数を1~4台と変化させて実験機の性能測定を行った。

3. データ収集

表1はデータ収集のために用いたサンプルプログラムとノード (命令) あたりの平均実行マシンサイクル数である。同表に示すように、この値はプログラムによって大きな変化はなく、ほぼ2.5サイクル程度である。

次に、実験機の台数効果について示す。台数効果は、プログラムを実行する (全解を求める) のに要したマシンサイクル数、及び、マシンの性能指標の1つであるRPS (Reductions Per Second) で示す。ここで、RPSは単位時間当りに成功したゴール呼出し (Reduction) の回数である。

図1はOR型Prologプログラム (プログラム番号1~5) に対する台数効果を、図2はOR型PrologプログラムとGHCプログラム (プログラム番号1, 2, 5, 6及び7) の台数効果の比較結果を示す。

表1. 平均命令実行サイクル数

番号	プログラム	言語	命令当りの 実行サイクル
1	7-queens	OR型	25.93
2	7-queens (決定的宣言有)	OR型	24.12
3	BCP	OR型	26.26
4	DCG	OR型	21.16
5	qsort	OR型	25.57
6	7-queens	GHC	27.52
7	qsort	GHC	23.56

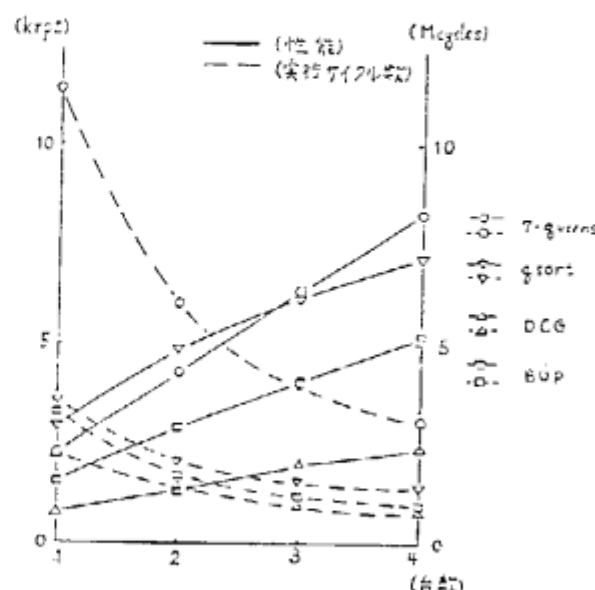


図1. OR並列型Prologプログラムの台数効果

* 1986年1月よりICOT ** 1986年1月より沖電気

ソフトウェアシミュレーション結果で示したように、プログラムに内在する並列性の大きいものほど性能の台数効果が顕著である^[5]が、実験機でも同様の結果を得た。例えば、7-queensは並列性が高く台数の増加に伴ってほぼリニアに性能が向上するのに対して、qsortは性能が飽和する傾向にある。この理由はqsortに関して並列性を引き出すためにAND並列を実現したが、台数の増加によるP E間転送のオーバーヘッドのほうが大きいためである。また、DCGの性能が低い理由は他と比べて1ゴール呼出し当りの起動される節の数は多いが、このうちの成功する節の数が小さいためである。

OR型PrologとGHCとを比較するにあたっては並列性の高い7-queens(決定的宣言の有・無について割定)と低いqsort(256個の要素の並べ換え)を選択した。そのときのP Eの総実行命令数とS Mの総操作指令数を表2に示す。決定的宣言はゴール呼出しを行って得られる解がただ一つとわかっているときに明示的に用いることができ、ストリームを介した通信を行うのではなく解を直接親ゴールへ返す。従って、S Mへのアクセス回数を大幅に減少でき、決定的宣言を行わなかった場合に比して命令総数で25%、指令総数で50%減少する。これによりマシンサイクルが20%短縮し、性能が50%程向上した。GHCプログラムはOR並列型より命令総数が50%、実行サイクル数も同様に50%程増加したが、性能は50%高くなっている。これはストリーム操作をユーザ定義述語として明示的に行っているためである(OR並列型では命令レベルでストリーム操作を行う)^[4]。

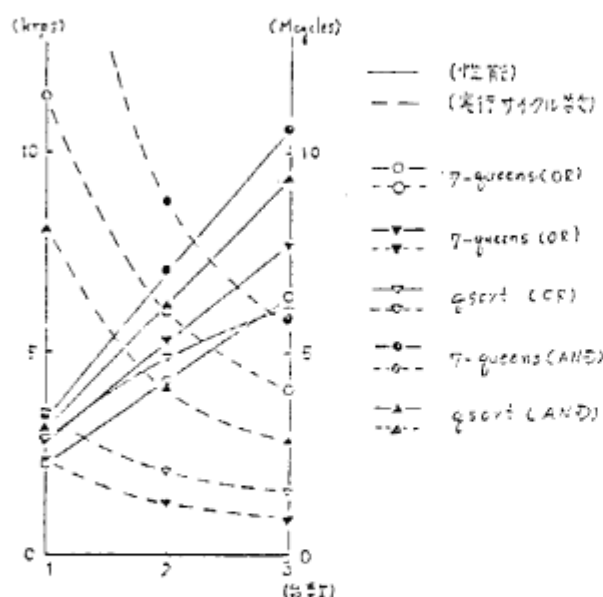


図2. OR並列型PrologとGHCプログラムの台数効果

qsortに関してはGHCの方がOR並列型Prologよりも高い台数効果が得られる。これはOR並列型がゴールレベルでプロセス間の同期をとるのに対して、GHCでは、ストリームを引数レベルで直接消費者プロセスに渡しているためである。

4. おわりに

並列推論マシン研究の一環としてPIM-D実験機を試作し、いくつかのサンプルプログラムに対する評価を行った。それによると、GHC及びOR型プログラムとも命令当りの平均所要サイクル数はほぼ一定である。P E台数を増加した際の並列処理効果はプログラム特性の影響を受ける。並列性の高いプログラムでは台数効果が得られ、処理方式の有効性が確認できた。今後、実験機の規模をP E16台、S M15台まで拡張する予定であり、更に大規模システムの評価を行う予定である。最後に、日頃ご指導いただくICOT第4研究室の内田室長はじめとする関係各位に感謝する。

表2. 命令実行数とS M操作指令数

プログラム	言語	命令数	指令数
7-queens	OR型	440050	190163
7-queens (決定的宣言有)	OR型	336829	96949
7-queens	GHC	660752	205145
qsort	OR型	138283	42906
qsort	OR型	103243	36380

【参考文献】

- [1]. Ueda, K., "Guarded Horn Clauses," ICOT TR-103, June, 1985.
- [2]. 桑住 他, 「データフロー方式並列推論マシン実験機の構成」, 情処30全, 6C-2.
- [3]. 大原 他, 「データフロー方式並列推論マシン」, 情処32全, 2R-4.
- [4]. 伊藤, 「並列型Prologのデータフローマシンによる実行方式」, 日経エレクトロニクス, 1984年11月5日号.
- [5]. 佐藤 他, 「データフロー方式並列推論マシン」, 情処32全, 2R-1.

データフロー方式並列推論マシン

2R-1

— ソフトウェアシミュレータによるGHCプログラムの評価 —

佐藤 正俊 伊藤 徳義 * (I C O T)
 末住 晶介 久野 英治 六沢 一昭 ** (沖電気)

1. はじめに

論理型言語Prologを並列で高速に実行するコンピュータの研究が現在さかに行われている。I C O Tでは、並列マシン上での実行言語としてAND並列型論理プログラミング言語を選択し、'84はConcurrent Prolog(CP)をベースに、'85はGuarded Horn Clauses(GHC)をベースに検討を進めてきた。筆者等はデータフロー方式に基づく並列推論マシン:PIM-Dの研究を行っており、PIM-D上でOR並列型Prolog及びCPの処理方式を検討すると共に、そのソフトウェアシミュレーションによる評価を行ってきた[Kuno 85]。新たに、PIM-D用GHCコンパイラを作成し、PIM-DシミュレータをもとにOR並列型PrologとGHC及びCPとGHCの比較評価を行ったので報告する。

2. PIM-D における各言語の特徴

OR 並列型 Prolog, CP, GHC はそれぞれ PIM-D のプリミティブをノードとするデータフローグラフにコンパイルされ PIM-D シミュレータ上で実行される。このため各言語の特徴はデータフローグラフの違いとして現れ、この相違が実行時間に影響すると思われる。以下各言語の相違点の概略を述べる。

i) OR 並列型 Prolog

複数の節から成る定義を呼び出す場合、節の間の実行順序関係が規定されていないならば起動された節は OR 並列に実行される。この時これらの節の実行によって得られた情報はシステムが用意したストリームを介して消費側のプロセスに渡される。

ii) CP

AND 並列型 Prolog であり AND 関係にあるゴールの間で論理変数を共有させると共に、ガード付コマンドと脱出し専用変数の機能を用いてこれらのゴールにより起動されたプロセス間で共有変数に対する結合情報をインタラクティブに通信できる。即ち、呼出された節の頭部またはガードにおける統一化で、ゴールから渡された共有変数が非変数項に結合された場合、コミットオペレータの実行を契機として結合情報をゴール側に知らせる。このため共有変数に関する結合環境は、節の頭部及びガードを実行している間はその節でローカルに保存する必要がある。

iii) GHC

CP と同様に AND 並列型 Prolog であるが GHC においては頭部またはガードにおいてゴール変数に対する結合を許さずその結合は中断されるためローカルな結合環境を作る必要はない[Ueda 85]。従って、CP より処理効率が良いと期待できる。

3. PIM-D シミュレータ上での言語の特性

a) 評価プログラム

典型的な探索型の問題であり、その特性がよく知られている点で N queens プログラムを評価対象として選択した。N の値はシミュレートする仮想マシンの規模がモジュール数 1 .. 64 台の範囲内であるので、この範囲内で並列性がほぼ飽和する 6 を選んだ。

b) 評価

6 queens プログラムの実行時間の台数効果について OR 並列型 Prolog, CP, GHC の測定結果を図 1 に示す。また PE 数が 16 台のときの 6 queens プログラムの実行時間(単位はマシンサイクル[110 84])と実行命令数の比較を表 1 に示す。

表 1 実行時間と実行命令数の比較

	実行時間	実行命令数
OR 並列型 Prolog	57446 (0.8)	100912 (0.6)
CP	185165 (2.5)	327449 (2.1)
GHC	73552 (1.0)	156116 (1.0)

()内はGHCを1とした時の値

以下シミュレーションデータに基づいて CP と GHC, OR 並列型 Prolog と GHC の比較を行う。

Kマシンサイクル

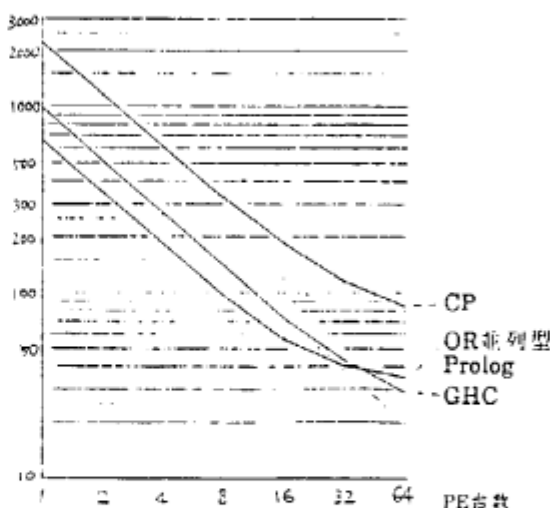


図 1 6 queens プログラムの実行時間比較

* 1986年 1月より沖電気 ** 1986年 1月よりI C O T

i) CP と GHC の比較

CP, GHC とともに AND 並列型 Prolog であるが、主な相違点はローカル環境を作成するか否かにある。つまり CP においてはローカル環境を LM (Local Memory) 上に作成するため、LM アクセス及びその制御のための処理がオーバーヘッドとなっている。PE 数が 16 台のときのローカル環境を作成のための LM アクセス時間と生成トークン総数の比較を表 2 に示す。表 2 よりローカル環境を作成のための影響がわかる。

表 2 LM アクセス時間と生成トークン総数の比較

	LM アクセス時間	生成トークン総数
GHC	0	156116 (1.0)
CP	217393	482874 (3.1)

()内は GHC を 1 とした時の値

ii) OR 並列型 Prolog と GHC の比較

OR 並列型 Prolog と GHC の相違は探索処理を OR 並列型で実行するか AND 並列型で実行するかの違いであるが PIM-D では、いずれの処理もストリーム通信を介したプロセス間の並列処理モデルによって統一的に実現している。但し、OR 並列型 Prolog ではシステムがシステムストリームを使用しプリミティブレベルで暗黙的に実行するのに対して、GHC ではユーザが明示的にプログラムしなければならない点で異なる。この相違は表 3 に示すようにゴールの呼出し回数及び実行命令数として現れる。

表 3 ゴールの呼出し回数及び実行命令数の比較

	呼出し回数	実行命令数
GHC	5616 (1.0)	156116 (1.0)
OR並列型Prolog	4262 (0.8)	100912 (0.6)

()内は GHC を 1 とした時の値

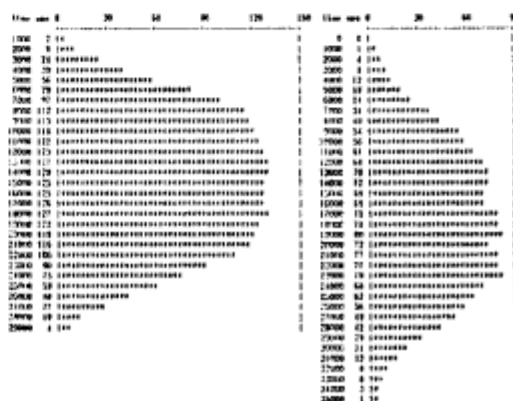
GHC の解の操作のための述語は merge 述語と append 述語でありこれらのゴールの呼出し回数(1576)を GHC の全ゴールの呼出し回数より差引いた値(4040)はほぼ OR 並列型 Prolog のゴールの呼出し回数(4262)と等しいと考えられる。また実行命令数については merge 述語、append 述語ともに約 20 ステップであり、再帰実行時の命令ステップはともに約 30 ステップである。これより merge 述語と append 述語の実行命令数はそれぞれ merge 述語で 7970、append 述語で 38790 となり合計で 46760 と推定できる。つまり実行命令数についても GHC の解の操作以外の実行命令数と OR 並列型 Prolog の実行命令数がほぼ等しいと考えられる。

iii) 飽和の傾向について

図 1 が示すように飽和の傾向は PE 台数 16 ぐらいから始まるが言語によってその傾向がやや異なる。特に OR 並列型 Prolog と GHC の傾向が異なっている。以下 OR 並列型 Prolog と GHC の場合を考察する。

OR 並列型 Prolog と GHC とともに PE 台数による実行命令数の変化はない。ただし、GHC の場合は PE 台数 64 の場合 PE 台数 16 に比べて若干増加しているが 0.1% 未満の増加でありほぼ同じと考えられる。この GHC での PE の台数による実行命令数の変化は merge 述語の実行回数が多少変るためである。即ち、主としてネットワークの遅延により、merge 述語における非決定性が実行回数の変化として現れるからである。

実行命令数の時間的変化を図 2 に示す。この図は PE 台数が 64 のときの 1000 マシンサイクルあたりに各 PE で実行された平均命令数(ops)を 1000 マシンサイクル毎に表示している。図 2 より GHC と OR 並列型 Prolog のシミュレーション上での並列性の現れかたの違いがわかる。つまり PIM-D においては GHC は引数レベルで OR 並列型 Prolog はゴールレベルで同期をとっているため並列性が PE 台数に比べて低い場合には並列性の現れかたの違いが顕著に現れてしまい、GHC の実行時間が短くなっている。



i) GHC ii) OR 並列型 Prolog

図 2 実行命令数の時間的変化

4. おわりに

①種類の並列型論理プログラミング言語についてその処理性能をPIM-Dシミュレータ上で 6 queens プログラムに対して評価した。その結果GHCはCPに比べて約2倍またOR並列型Prologに比べて約0.8倍の処理性能であった。またここでは触れなかったが他のプログラムについてもほぼ同様の結果を得ている。

今後はストリーム処理における最適化及びプリミティブレベルの引き上げによる最適化等によりPIM-DにおけるGHCの処理方式をさらに検討して行く予定である。最後に、日頃ご指導をいただく内田第4研究室はじめ並列推論マシングループ諸氏に感謝の意を表する。

(参考文献)

- [Kuno 85] 久野他, 「データフロー方式並列推論マシンにおけるAND/OR並列効果のシミュレーションによる測定」情報処理第30回全国大会予稿集。
- [Ito 84] 伊藤, 「並列型Prologのデータフローマシンによる実行方式」日経エレクトロニクス, 1984年11月5日号。
- [Ueda 85] Ueda, K., "Guarded Horn Clauses," ICOT TR-103, June, 1985.

関係データベースマシンにおける

3C-10 関係代数演算専用エンジンの並列処理評価

伊藤 英則

安部 公朗

角田 健男

(財) 新世代コンピュータ技術開発機構

1. はじめに

ICOTの第5世代コンピュータ・プロジェクトは、推論機構と知識ベース機構を兼ね備え、高度な知識情報処理システムの構築を目指している。前期3年間に於いて、知識ベース機構については、Prologなどの論理型言語との親和性を生かし、バックエンド型関係データベース・マシン“Delta”を開発した[1][2]。本稿では、関係代数演算専用エンジンを複数台持つような“Delta”を抽象化したバックエンド型関係データベース・マシンを想定する。それらのエンジンが並列に動作して関係代数演算を行う時に、実際の処理系の実測データ・経験データを用いて、シミュレーションにより評価を行い、応答特性について考察する。

2. 関係代数演算専用エンジンの関係代数演算処理方式

関係代数演算を高速に実行するためには、演算の対象となるデータを前もってソートしておくことが効率的である。処理負荷の重いJOINなどは、特にその効果が高い。そこで我々は、ソート処理を高速に行うために、ソートすべきデータをストリーム状に渡し、ストリームの転送時間とオーバーラップしてデータをソートするPipelined 2-way merge sort アルゴリズムを採用し、専用ハードウェア“ソータ”を開発した。これにより、ソーティングのオーダを $O(N \log N)$ から $O(N)$ に落とすことが可能となった。また、ソートされたストリームに対して、関係代数演算を高速に実行するために、ソータの後段にその専用ハードウェアである“関係代数演算処理部”を置いた。ソータは、さらに n 段のプロセッサと、各プロセッサの持つメモリ部からなる。実際に実現したソータは、 $n = 12$ で、最終段のプロセッサのメモリ部のサイズが64KBである。従って、64KB以下のデータを一度にソート出来る。関係代数演算処理部は、既にソートされたリレーションを格納する2つのメモリと、条件を満たすレコードだけを選択出力するプロセッサからなる。JOINやSELECTIONなどの関係代数演算は、関係代数演算処理部で実行される。実際に実現した関係代数演算処理部の2つのメモリサイズは64KBであるので、64KB以下のリレーションと64KB以下のリレーション(あるいは定数)のJOIN、SELECTIONが一度に実行できる。

3. データ分割並列処理方式

3.1 データ分割ソート演算

大量データをソートする時、ソートすべきデータを空きエンジンの台数分まで分割し、各分割データを空きエンジンに割当て並列実行することができれば、実行時間を著しく短縮することができる。これをデータ分割ソート演算と呼ぶ。データ分割ソート演算は次のステップで実行される。

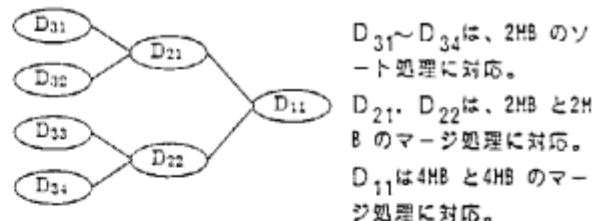
- ① ソートすべきデータの長さを L とし、空きエンジンの台数を m とした時、データを m 分割して L/m のデータをそれぞれのエンジンに割当て、並列にソートを行う。
- ② ソート済みのデータを2個づつ集めてP個のペアーを作り、エンジンの関係代数演算処理部で2-way mergeを行うことにより、P個のソート済みのデータを出力する。この処理を空きエンジンで並列に行い、 L の長さのデータがソートされるまで繰り返す。

3.2 データ分割JOIN(SELECTION) 演算

2つのリレーションのデータサイズをそれぞれ L_1 , L_2 とした時、両者のJOINを $L_1 \times L_2$ と表わす。これを m 台の空きエンジンで並列実行させるためには、片方のリレーションを m 分割し、 $L_1/m \times L_2$ のJOINを1台のエンジンに担当することを考える。これを、データ分割JOIN演算と呼ぶ。同様な方法で、SELECTIONを並列実行させる場合をデータ分割SELECTION演算と呼ぶ。

3.3 データ分割サブコマンド木

データ分割ソート/JOIN/SELECTION演算は、図1のようなデータ分割サブコマンド木に変換され実行される。

図1 8HBのソートのデータ分割サブコマンド木($n = 4$)

データ分割サブコマンド木は、葉から根に向かって処理が進み、同じ深さのノードは並列に実行可能である。各ノードは、前のノードの実行が全て終了しないと、実行を開始できない。

4. 関係データベースマシン・シミュレーションモデル

本稿で対象としている関係データベースマシンの概念図を図2に示す。インタフェース部は、ホストマシンから問合せコマンド C を受取り、 C をコントロール部に送る。また、 C のレスポンス・データをホストに返す。コントロール部は、インタフェース部より送られてくる C を解析しソート処理やJOINなどの関係代数演算レベルのコマンド C に変換し、スケジュール部に送る。スケジュール部は、空きエンジンの状況を監視して、エンジン並列制御戦略によりエンジンの並列実行制御を担当する。また、それにより選ばれた C をデータ分割サブコマンド木に変換する。階層構造メモリ(HM)部はリレーションの格納場所であり、演算の

対象となるリレーションをディスクからバッファメモリ上のステージング、リレーションの分割処理を担当する。

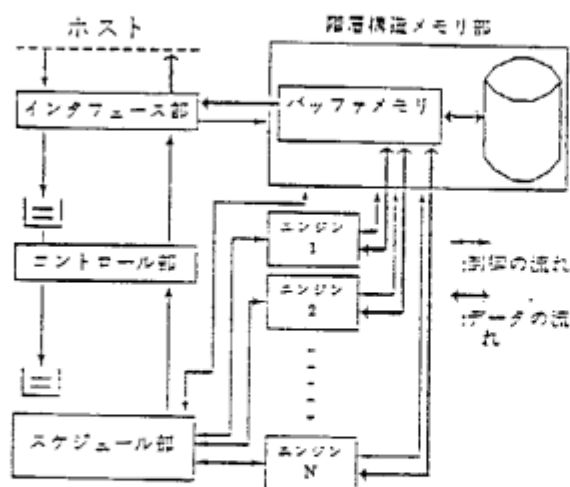


図2 関係データベース・マシンの概念図

5. エンジン並列制御戦略

エンジン並列制御戦略として、以下の3つを考える。

(1) データ非分割先着優先戦略

どのコマンドもデータを分割せずに、スケジュール部のキューに到着した順にm台の空きエンジンを割当てる。

(2) データ分割先着優先戦略

スケジュール部のキューから先着順にコマンドを1つずつ取り出し、空きエンジンの台数分データを分割し、データ分割サブコマンド木に変換し、実行する。ただし既にデータ分割サブコマンド木に変換され、処理が根にまで至っていない木が存在するならば、それらの木のうち到着の早い木のノードの処理を優先して割当てる。

(3) データ最多分割先着優先戦略

(2) でデータを分割する際、実装されているエンジンの総台数分で分割する。あとの条件は(2)と同じ。

6. 目的関数

目的関数として、次の2つを考える。

(1) 平均待ち時間（平均レスポンスタイム）

(2) 平均HHメモリ使用量

(3) 平均エンジン稼働率

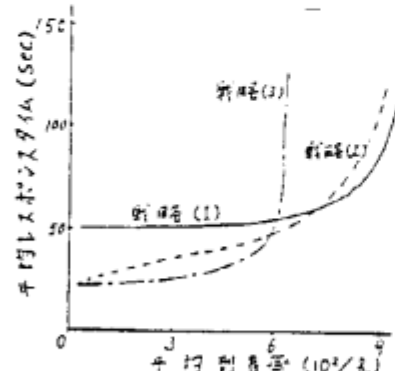


図3 平均到着率と平均レスポンスタイム

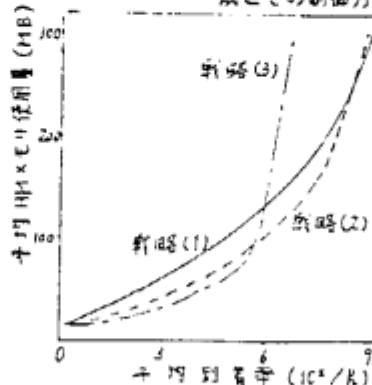


図4 平均到着率と平均HHメモリ使用量

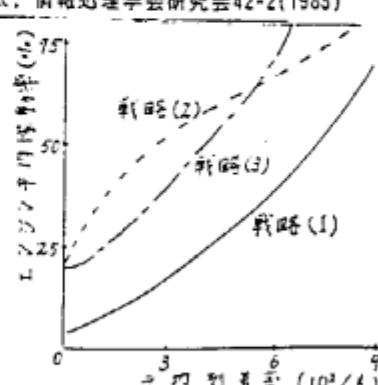


図5 平均到着率と平均エンジン稼働率

7. シミュレーションの条件

(1) 装置の条件

- ① インタフェース部、コントロール部、スケジュール部、HH部は、5MIPS の処理能力を持つ1CPU構成とし、エンジンは、その付加プロセッサとする。
- ② エンジンの台数は、16台とする。
- ③ HHのバッファメモリサイズは、128MB とする。
- ④ HHメモリとエンジン間、及びHHメモリとディスク間データ転送速度は3MB/S とする。また、HHメモリとエンジン間は、入力1本、出力1本、合計2本のチャンネルを具備するものとする。

(2) 処理時間に関する条件

Delta で実測可能な時間は、実際の処理系の時間を使用する。それ以外の処理時間は、実測値からの推定、ダイナミックステップ数の見積りにより求めた。

(3) ホストマシンからの問合せコマンドに関する条件

問合せコマンドの到着は、ポアソン到着を仮定し、ソート、JOIN、SELECTIONが等確率に1つずつ到着するものとする。また、演算の対象となるリレーションのサイズは、1MB から16MBまでランダムに選ぶ。

8. 評価結果

シミュレーション結果を図3、図4、図5に示す。

9. 考察とまとめ

図3より、トラフィックが疎の時には1つの演算をより多くの台数で実行した方が効率が良いが、密の時にはデータを分割しないで実行した方が良い効率が得られた。この傾向は、図4についても同様である。これらは、トラフィックの疎密により戦略を変える必要があることを示す。

今後、この評価結果をもとに、中期に開発を予定している知識ベースマシンのユニフィケーション・エンジン並列制御方式の検討を行う予定である。

【参考文献】

- [1] Kakuta T., et al. The Design and Implementation of Relational Database Machine Delta. Proc. IJDM'85, (1986)
- [2] 安部他, 関係データベースエンジンのハードウェア構成とその制御方式, 情報処理学会研究会42-2(1985)

単一化エンジンを用いた知識ベース・マシンの構成方法

横田 治夫

村上 昌己

森田 幸伯

伊藤 英則

(財) 新世代コンピュータ技術開発機構

1. はじめに

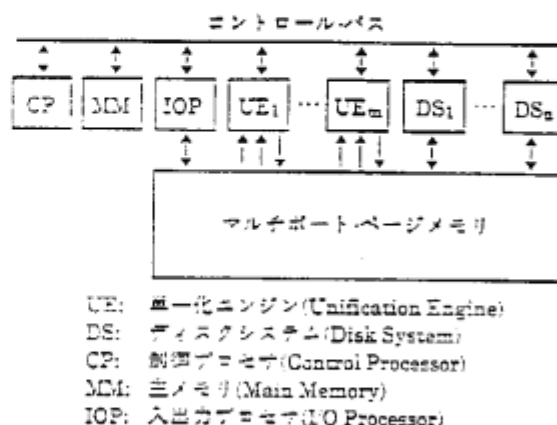
計算機で知識を取り扱う研究が、注目されている。そのような知識情報処理系では、大量の知識の中から必要な知識を効率良く検索する機能と、検索された知識を用いて高速に推論する機能とが重要となる。第五世代コンピュータシステム研究開発プロジェクトでは、それぞれの機能実現のため、知識ベース機構と推論機構の構築をめざしている。我々は、知識ベース機構実現の一つのアプローチとして、関係型知識ベース・モデルを提案した[1]。本稿では、関係型知識ベース・モデルに基づく知識ベース・マシンの構成方法と、そのマシンの中心となる構成要素であるストリーム処理を基本とする単一化エンジンの実現方法を提案する。

2. 関係型知識ベース・モデル

データと知識の違いは、データが必要な情報をそのまま表現しているのに対し、知識は相互に関連し合った必要な情報を適当な構造に置き換えて表現しているものといえる。我々は、大量の知識を格納するため、関係データベースの各アイテムを定数から明確に定義された構造物である項に拡張して、単一化(unification)を用いて項を検索するモデルを提案した。項が格納されたテーブルのことを項関係、関係データベースにおける結合および制約演算のイタリティ・チェックを単一化操作に拡張した演算のことをそれぞれ単一化結合(unification-join)、単一化制約(unification-restriction)、それらの検索処理全体をRBU(Retrieval By Unification)と呼ぶ[1]。

3. 知識ベース・マシンの全体構成

関係型知識ベースをサポートするハードウェアとして、図1に示すような構成の知識ベース・マシンを提案する。2次記憶に格納された項関係に対し、高速にRBU操作を行うため、ディスク・システムから単一化エンジンに項の集合をストリームとして流すことを前提としている。図中のマルチポート・ページメモリは、複数の処理装置からページ単位で同時にアクセスしても競合が起らないように設計された共有メモリである[2]。本システムでは、ディスク・キャッシュとして、あるいは単一化エンジン用のバッファ・メモリとしてマルチポート・ページメモリを用いている。紙面の関係上、マシンの全体構成の詳細は割愛するが、全体アーキテクチャについては[3]を、システム全体の制御については[4]を参照されたい。



UE: 単一化エンジン(Unification Engine)
DS: ディスクシステム(Disk System)
CP: 制御プロセッサ(Control Processor)
MM: 主メモリ(Main Memory)
IOP: 入出力プロセッサ(I/O Processor)

図1. 知識ベース・マシンの構成

4. 単一化エンジン

4.1 エンジンの構成

単一化エンジンの構成を図2に示す。単一化エンジンは、単一化結合を前提として構成するため、マルチポート・ページメモリとの間で、同時に2つの入力ストリームと1つの出力ストリームを処理できるように、3本のデータバスを持つ。単一化結合での単一化の項の組合せは、本来2入力ストリーム中の項の数の積になる。項間にある順序付けを行うことにより、この組合せの数を減らすことができる。可変長ソータとペア・ジェネレータはそのためのユニットである[5]。単一化ユニットは、ペア・ジェネレータにより生成された2項間の最適化作用素(mgu)を求める。また、前処理ユニットは項関係のタプルから処理対象のアイテムを切り出し、後処理ユニットは単一化ユニットで得られたmguをもとのタプルに適用するためのユニットである。

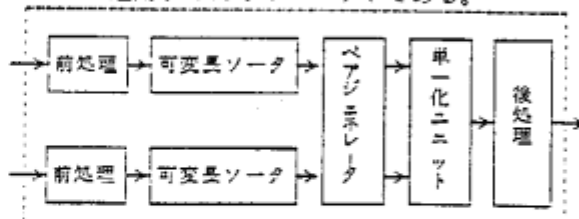


図2. 単一化エンジンの構成

4.2 単一化アルゴリズム

一般的なPROLOG処理系などでの単一化のアルゴリズムは、構造物共有方式にしろ構造物コピー方式にしろ、変数ごとにバインドすべき相手をポインタでつなげてゆくことを基本としている。そのようなポインタをたどる操作は、ストリーム処理を基本とする専用ハードウェアには向いていない。

もともと、単一化アルゴリズムは図3のように書くことができる[6]。図3のアルゴリズムの繰り返しの部分(Step2-Step4)をハードウェア化し(単一化エレメントと呼ぶ)、それを直列に接続して(図4)、複数の項の組合せを流すことにより、単一化処理をパイプラインで行うことができる。

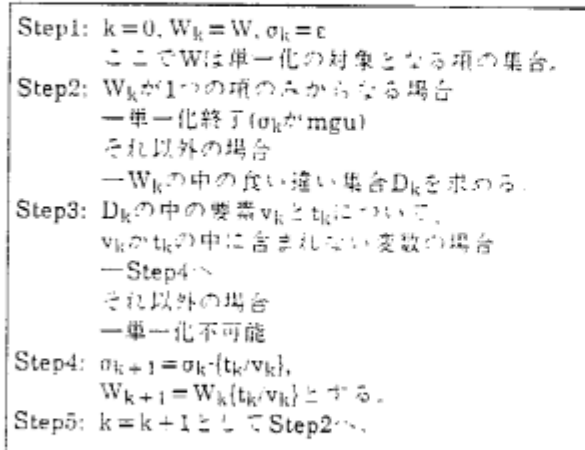


図3. 単一化アルゴリズム

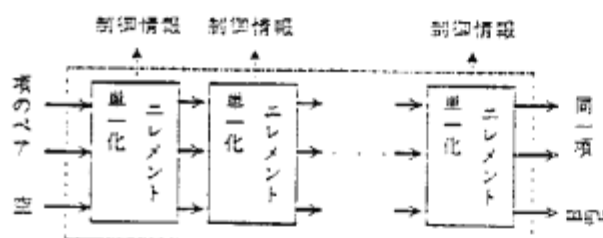


図4. 単一化アルゴリズムのハードウェア化

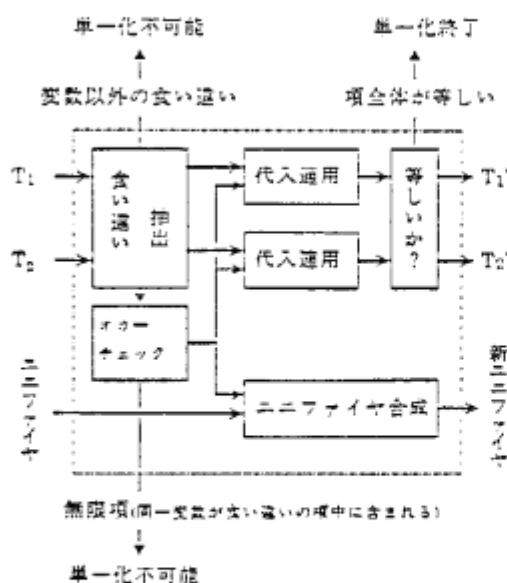


図5. 単一化エレメントの内部構成

4.3 単一化エレメント

図4中の1つの単一化エレメントの内部構成を図5に示す。これは、図3のアルゴリズムをそのままハードウェア化した形になっている。Stepに対応したブロック構成で、単一化エレメント中もパイプライン化される。

4.4 単一化ユニットの構成案

実際に単一化ユニットを構成する場合、項中の変数の数に制限がないと、無限個の単一化エレメント(ue)をつなげる必要がある。これは現実的に不可能なため、図6のaやbのように切り替えスイッチやスイッチング・ネットワークなどを使って、仮想的に無限長の接続を実現する。

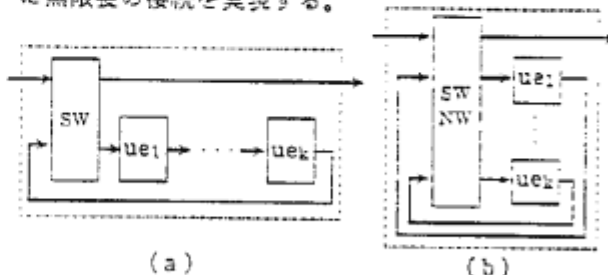


図6. 単一化ユニットの構成案

5. おわりに

関係型知識ベース・モデルに基づく知識ベース・マシンの構成方法と、その基本構成要素の単一化エンジンの実現方法について述べた。全体構成の方針は、二次記憶に格納された大容量の処理対象を専用装置で処理するようなシステム一般に適用できる。また、単一化エンジンも、知識ベース・マシンだけでなく、各種の知識情報処理システムで利用できる。今後は、知識ベース・マシンの実現方法についてより詳細に検討を進めていく。

〔謝辞〕本検討を進めるに当たり、有益な御示唆を頂いた知識ベースマシンWGのメンバーの方々に感謝します。

〔参考文献〕

- [1] 横田, 他: 単一化による知識ベース検索, 60年度通信学会部門別全国大会, S12-2.
- [2] 横田, 他: マルチポート・ページメモリの構成方法, 60年度通信学会部門別全国大会, 602.
- [3] Yokota and Itoh: A Model and an Architecture for a Relational Knowledge Base, ICOT TR-144 (submitted to the 13th Int. Sympo. on Computer Architectuer).
- [4] 物井, 他: マルチポート・ページメモリを用いた大規模知識ベースの制御方式, 情報第32回全国大会予稿, 50-5.
- [5] 森田, 他: 単一化結合の処理方式, 情報第32回全国大会予稿, 1M-7.
- [6] Chang and Lee: Symbolic Logic and Mechanical Theorem Proving, Academic Press.

単一化結合の処理方式

15-7

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

1. はじめに

第3世代コンピュータシステム研究開発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、その概念構造は柔軟性に富んでいることが望ましい。[1]で提案されている関係型知識ベースモデルは、知識ベースマシンに対する非常に柔軟な概念モデルとなっている。しかし、そこで提案されているRBU演算、特に単一化結合を単純に行うとマシン側の処理量が非常に大きくなる。本稿では、単一化結合を効率よく行うための処理方式およびその一実現方式について報告する。

2. 関係型知識ベースモデル

関係データベースでは、関係の各属性(attribute)の領域(domain)をアトミックなものに集合に制限している。

それに対して関係型知識ベースモデルは、関係の各属性の領域として項の集合を許した項関係を扱う。項関係上の操作としては、関係代数のイクェリティチェックを単一化操作(unification)に拡張した単一化結合(unification-join)、単一化制約(unification-restriction)を用いる。これらをRBU(Retrieval By Unification)演算とよぶ。射影(projection)は関係代数のそれと同じである。

3. 単一化結合の処理方式

我々は、非常に大量な知識が知識ベースマシンに格納されてい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つの項関係のタプル(tuple)のすべての組合せに対して指定されたアイテムが単一化可能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、1/O回数も増え非常に大きな処理負荷がかかる。そこで項関係のタプルを何らかの順序で順序付けることにより組合せを省略したい。

項を順序付けるためにジェネラリティの概念を導入する[1]。ただ、ジェネラリティの順序付けは、全順序ではないため、ジェネラリティの順序付けを保存しつつ、全順序にする方法を考える。そこで、項が木で表現されることに注意し、その木の線形化とその辞書的順序を用いることにする。

木の線形化としては、深さ優先のものと、レベル優先のものがある(図1)。本稿では深さ優先のものを

leftmost方式の順序と呼び、レベル優先のものをoutermost方式と呼ぶ。[1]で定義した順序はleftmost方式の1つである。但し、関数記号はそのアリティにより区別され、関数記号ごとは適当な順序があり、変数はその出現位置によって順序を付け、関数記号と変数は変数の方がよりジェネラルであるとする。

$f(g(Y), a)$	$f2, g1, Y, a$	$f2, g1, a, Y$
項	深さ優先	レベル優先

図1. 項の表現の例

このように(ジェネラルな順に)順序付けられた項の列に対して、文字列どうしの比較操作だけで単一化可能な可能性のある項のペアを選びだしたい。

項を文字列で表現したとき単一化可能なものは先頭から見てどちらかに変数が現れるまで文字列は一致しているはずである。文字列が順序付けられているので、変数より前方の文字列が一致しなくなるまで、その項をとっておきペアをつくりだす。異なったらそれ以降の項の列に対しては、単一化は不可能なので組合せを省略できる。

省略の仕方は順序付けの方式によって異なる。組合せの省略の例を図2に示す。□は単一化可能な項の組を、◇はペアとして選ばれたが単一化不可能な組合せを、◇は省略された組合せをしめす。図2-aがleftmost方式、図2-bがoutermost方式である。

オーダリングの概念は、2次記憶の項のクラスタリングなどにも有効な概念とおもわれる。

4. 単一化エンジンの一実現方式

4.1. 単一化エンジンの構成

単一化エンジンの構成を図3に示す。単一化エンジンは、次の5つのものからなる。

項関係のタプルから処理対象のアイテムを切り出す前処理ユニット、オーダリングに従い項をソートする可変長ソータ、ソートされた2つの項の列を受け取り、単一化の可能性のある項の組(ペア)を生成するペアジェネレータ、生成された項の組の最汎化作用素(mgu)を求める単一化ユニット、さらに最汎化作用素をもとのタプルに適用する後処理ユニットである。

本稿では、上記のうちソータおよびペアジェネレータについて述べる。この2つのユニットで、前節で述べた方式に従って組み合わせの省略を行う。

単一化ユニットに関しては、別稿[2]を参照されたい。

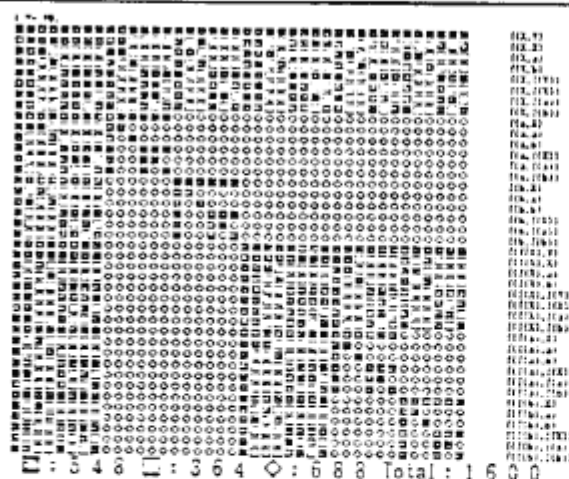


図2-a. 組合せの省却の例 (leftmost方式)

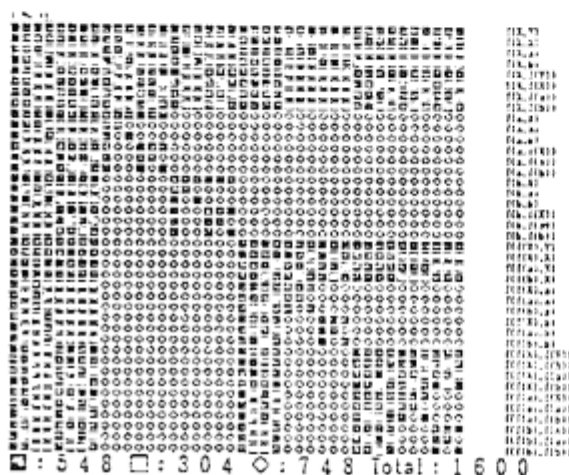


図2-b. 組合せの省略の例(outermost方式)

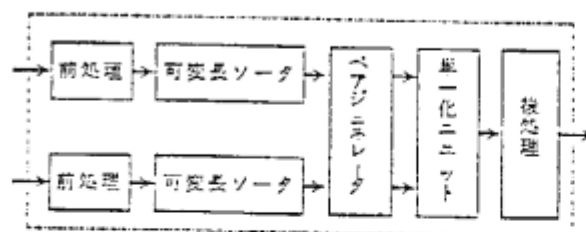


図3. 単一化エンジンの構成図

4. 2. $\gamma - \delta^*$

ソート部は、可変長の文字列を扱うソークで、方式としては2ウェイのマージソートを用いる。ハードウェアソータの方式としては、マージソート、バイトニックソート、ヒープソート等があるが、我々は可変長の2ウェイマージソート方式を採用する。項のソートのためには可変長文字列のソートを行わなければならない。可変長のハードウェアソータとしてはパイプラインドヒープソータ[3]が提案されているが、ストリーム型のデータ転送を基本と考えたときには、マージソータが適していると思われる。

可変長の文字列を文字単位に巻戻しなしにストリームタイプで処理するために、文字列の線形化 *Tr i e* 表現を用いる[3]。

線形化Trie表現とは、文字列を文字と位置情報
の組の列で表現し、直前の文字列と同じものは省
略して表わす方式である。項の列を線形化Trie
表現で表した例を図4に示す（例における項の順
度は、leftmost方式である）。

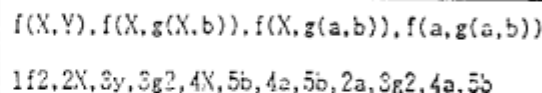


図 4. 項の列の線形化 *trie* 表現

4. 3. ペアジニホレータ

ペアジェネレータの構成を図5に示す。ペアジェネレータは、必要な項を単一化の可能性がなくなるまで、スタックに預けておき、入力された項と比較し文字列の比較で単一化不可能とされるもの以外の項のペアを全て出力する。

なお、出力の際には、項の表記方法を統一化ユニットで取扱いやすい形に変換する。

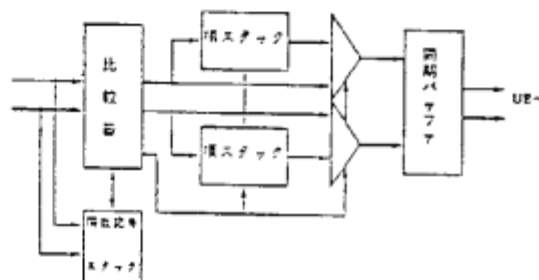


図5. 大アミエビ類の捕食率

5. 10 10

REU方式に基付く処理方式および実行方式の提案した。今後は提案したアルゴリズムやエンジンについてシミュレーション等により評価を行っていく予定である。

【謝辞】本検討を進めるに当たり、有益なご示唆を頂いた知識ベースマシンWGのメンバーの方々に感謝いたします。

[參考文獻]

- [1] Yokota, H. "A Model and an Architecture of a Relational Knowledge Base", submitted to the 13th in. Simp. Computer Architecture, 1986.
- [2] 横田他, 「単一化エンジンを用いた知識ベースマシンの構成方式」第32回情報処理学会全国大会 (1986)
- [3] Tanaka, Y. "A VLSI Algorithm for Sorting Variable-Length character Strings.", New Generation Computing, 1985.

単一化による関係型知識ベース操作

— 関係型知識ベースにおけるホーン節集合の更新 —

Manipulation of Relational Knowledge Base based on Unification Operation

— Update of Horn Set in Relational Knowledge Base —

坂田 千秋*

村上 昌己*

横田 治夫*

Chiaki SAKAMA

Hasaki MURAKAMI

Haruo YOKOTA

(株)東芝 情報システム研究所* (財)新世代コンピュータ技術開発機構(ICOT)*

1. はじめに

第五世代コンピュータプロジェクトでは、大量の知識を効率的に管理、運用するための知識ベース・メカニズムの開発を進めている。

その方法として、知識を関係モデルの上で表した関係型知識ベースモデルが提案されている[1]。また関係型知識ベースの上の演算として、単一化制約、単一化結合を導入し、これらの演算を用いて関係型知識ベースに格納されたホーン節集合の上の入力演算が実行できることが示されている。

関係型知識ベースにおける更新を考えたとき、その基本操作としての関係演算が必要となる。本稿では、新たに単一化可能抽出、単一化可能例なる関係演算を導入し、これらを用いた関係型知識ベースにおけるホーン節集合の更新について述べる。

2. 関係型知識ベース演算

知識を関係モデルの上で表現した、関係型知識ベース(以下、RKBと略す)は項関係

$T \subseteq K_1 \times K_2 \times \dots \times K_m$ (K_i : Term, T : term relation)

の集合と定義され、項関係は項タブルの集合

$T = \{(t_1, \dots, t_n) \mid t_i \in \text{Term}\}$

と定義される。(ここで、Termは項の集合とする。)

RKBに対する演算は関係データベースに対する演算である関係代数の拡張として、関係における各要素間のイクオリティ・チェックを単一化に拡張したものが考えられている。以下では、関係代数における集合演算の単一化可能性による拡張として2つの新しい演算を導入する。

Def.1 単一化可能抽出(unifiable extraction)

$R \triangleright S = \{(x_1, \dots, x_n) \mid (x_1, \dots, x_n) \in R, \exists (y_1, \dots, y_n) \in S, \exists \lambda \in \text{Sub}, \lambda(x_1, \dots, x_n) = \lambda(y_1, \dots, y_n)\}$

Def.2 単一化可能例(unifiable instance)

$R \triangleright S = \{(x_1, \dots, x_n) \mid (x_1, \dots, x_n) \in R, \exists (y_1, \dots, y_n) \in S, \exists \lambda \in \text{Sub}, (x_1, \dots, x_n) = \lambda(y_1, \dots, y_n)\}$

ここで R, S は項関係、 $(x_1, \dots, x_n), (y_1, \dots, y_n)$ は項タブル、Sub は置換の集合を表すものとする。また、項タブルの単一化はタブル中の項が単一化可能であるときに定義される[2]。

3. ホーン節集合の更新

関係型知識ベースにおいては、ホーン節はヘッド、ボディからなる2属性の項関係によって格納される。即ち、ホーン節のうち確定節(Definite clause)の集合 S

$S = \{P_1 :- Q_{11}, Q_{12}, \dots, Q_{1n_1}, \dots, P_2 :- Q_{21}, Q_{22}, \dots, Q_{2n_2}, \dots, P_m :- Q_{m1}, Q_{m2}, \dots, Q_{mn_m}\}$

は項関係、

$S = \{([P_1 \mid X_1], [Q_{11}, Q_{12}, \dots, Q_{1n_1} \mid X_1]), ([P_2 \mid X_2], [Q_{21}, Q_{22}, \dots, Q_{2n_2} \mid X_2]), \dots, ([P_m \mid X_m], [Q_{m1}, Q_{m2}, \dots, Q_{mn_m} \mid X_m])\}$

と表され、ゴール節(Goal clause) G

$?- G_1, G_2, \dots, G_n$

は項関係、

$G = \{([G_1, G_2, \dots, G_n])\}$

と表される。

以下では、前節で定義した関係演算によるRKBにおけるホーン節集合の更新について述べる。

3.1 追加

いまホーン節集合 S に、ある限定節を追加して S' とするとき、

$\forall R (S \vdash R \Rightarrow S' \vdash R) \quad (\vdash: \text{valid})$

であることが必要条件である。

即ち、 S から演算可能なすべての導出形 R は S' から演算可能でなければならない。項関係 S に項タブル C を追加して S' とするときのアルゴリズムは以下のようになる。

```
insert(S, C, S'):
begin
  P ← true
  for all T included in S do
    begin
      if (deducible(S-T, T) = false
        or deducible((S-T) U C, T) = false)
      then P ← false
    end
  if P = true
  then S' ← S U C
  else
    S' ← S
  end
```

ここで、 $U, -$ は関係代数における集合和、集合差を表す。

上で $\text{deducible}(S, T)$ は S から T が演算可能なとき true、そうでないときは false とする。そのアルゴリズムは以下の通り。

```
deducible(S, T):
begin
  S0 ← S
  S1 ← S
  R0 ← φ
  i ← 0
  if T ⊇ S0 ≠ φ
  then deducible(S0, T) = true
  else
```



```

while  $S_{i+1} \neq S_i$  do
begin
 $R_i \leftarrow \pi_{S_i}(\sigma_{S_i}(S_i \bowtie R_i))$ 
 $S_{i+1} \leftarrow S_i \cup R_i$ 
if  $T \models S_{i+1} \neq \emptyset$ 
then  $\text{deducible}(S_{i+1}, T) = \text{true}$ 
else
 $S_{i+1} \leftarrow S_i, i \leftarrow i+1$ 
end
end

```

ここで、 $\bowtie$ 、 π はそれぞれ関係代数における結合、射影を表し、 $A \cup B$ は A と B の単一化を表すものとする。
いま、R が S に対して冗長であるとは
 $S \models R \vee S \models \bar{R}$ ($\bar{R}$: satisfiable)
を満たすものとすると、追加によって項関係 S に冗長性が発生したとき、S から冗長性を除去して S' とするアルゴリズムは以下になる。
 $\text{redundancy}(S, S')$:
begin
 $S' \leftarrow S$
for all T included in S do
begin
if $\text{deducible}(S-T, T) = \text{true}$
then $S' \leftarrow S' - T$
else if $\text{satisfiable}(S-T, T) = \text{true}$
then $S' \leftarrow S' - (S \bowtie T)$
end
end

上で $\text{satisfiable}(S, T)$ は S において T が充足可能とき true、そうでないときは false とする。そのアルゴリズムは、 $\text{deducible}(S, T)$ において単一化可能例を単一化可能抽出に置き換えることによって得られる。

なお、冗長性の除去については処理効率の問題で知識ベース管理者の判断に委ねられる。

3.2 削除

いまホーン節集合 S から、確定節 C を削除して S' とするとき、

$$\forall R (S \models R \wedge S - C \models R \Rightarrow S' \models R)$$

であることが必要条件である。

即ち、S からの演繹過程において C を定理として必ず用いるようなすべての導出形 R は、 S' からは演繹不可能でなければならない。項関係 S から項タプル C を削除して S' とするときのアルゴリズムは以下になる。

```

delete(S, C, S'):
begin
 $S' \leftarrow S - C$ 
for all T included in S do
begin
if ( $\text{deducible}(S-T, T) = \text{true}$ 
and  $\text{deducible}((S-T)-C, T) = \text{false}$ )
then  $S' \leftarrow S' - T$ 
else
 $S' \leftarrow S'$ 
end
end

```

3.3 変更

いまホーン節集合 S の中のある確定節 C を変更して S

を S' とするとき、

$$\begin{aligned} &= R (S \models R \wedge S - C \models R \Rightarrow S' \models R) \\ &\wedge \forall R (S - C \models R \Rightarrow S' \models R) \end{aligned}$$

であることが必要条件である。

即ち、S からの演繹過程において C を定理として用いるような導出形 R で S' からは演繹不可能なものが存在し、それ以外のすべての導出形 R は S' からも演繹可能でなければならない。項関係 S 中のある項タプル C を C' に変更して S を S' とするときのアルゴリズムは以下になる。

```

change(S, C, C', S'):
begin
 $S' \leftarrow (S - C) \cup C'$ 
for all T included in S do
begin
if ( $\text{deducible}((S-T), T) = \text{true}$ 
and  $\text{deducible}((S-T)-C, T) = \text{false}$ )
then if  $\text{deducible}((S-T)-C') \cup C', T) = \text{false}$ 
then  $S' \leftarrow S' - T$ 
else  $S' \leftarrow S'$ 
else  $S' \leftarrow S'$ 
end
end

```

なお、実際の RK8 の更新にあたっては、回復、その他の諸問題を考える必要がある。

4. おわりに

関係型知識ベース操作として、従来の関係代数に加えて単一化可能性に基づく関係演算を導入することによって、関係型知識ベースにおけるホーン節集合の更新が実現できることを示した。

今後は、これらの操作をさらに拡張するとともに、実現方法について検討を進めていく予定である。

〔謝辞〕本研究を進めるに当たり、有益な討論をして頂いた研究室、並びに KBM 会議メンバーの皆様に感謝します。

〔参考文献〕

- [1] Yokota, H. et al: A Model and an Architecture for a Relational Knowledge Base, ICOT Technical Report, TR-144 (1985).
- [2] Murakami, M. et al: Formal Semantics of a Relational Knowledge Base, ICOT Technical Report, TR-149 (1985).
- [3] Miyachi, T. et al: A Knowledge Assimilation Method for Logic Databases, New Generation Computing, Springer, pp385-404 (1984).

単一化検索言語による 知識ベースソフトウェアの記述

151-9

村上 昌己 横田 治夫 伊藤 英則

(財)新世代コンピュータ技術開発機構

1. はじめに 筆者らは先に、知識ベースに大量に格納された知識を効率的に扱う方法として、関係型知識ベースの単一化検索による方法を提案した(1)。しかし関係型知識ベースモデルは、言わば「知識表現のアセンブラ」である。すなわち、高級な知識表現言語が関係型知識ベースモデルの上で記述されなければ、アプリケーション同士の知識を関係型知識ベースに格納し活用できない。

本稿では、意味ネットワーク的知識表現DCKR(2)の一部を関係型知識ベースの上に実現し、それらの保守・更新メソッドとして、has_a インヘリタンスの推移性の保持が関係型知識ベースの上で定義する項関係演算を用いて記述できることを示す。

2. 関係型知識ベース 関係型知識ベースにおける基本的なデータ構造は項関係である。

(定義1) 変数記号の集合をVar, 関数記号の集合をFun とする。ここでFun は有項集合で各元は固有のアリティをもつものとする。またFun とVar は共通部分をもたないものとする。Var の元とFun の元から通常のように定義される項の集合をTermとあらわす。Termの元に対する代入の集合をSub とあらわす。

$t_1, t_2 \in \text{Term}$ の単一化(unification)を通常通り定義する。

(定義2) $t_1, t_2, \dots, t_n \in \text{Term}$ f をFun に含まれない特別なn 属性関数記号とする。このとき、 $f(t_1, t_2, \dots, t_n)$ をn属性項タプルとよぶ。以下では f を省略して、

$$(t_1, t_2, \dots, t_n)$$

と書く。項タプル同士の単一化は次のように定義される。

(定義3) $t = (t_1, t_2, \dots, t_n)$, $t' = (t'_1, t'_2, \dots, t'_n)$ のとき t と t' が単一化可能であるとは、 $f(t_1, t_2, \dots, t_n)$ と $f(t'_1, t'_2, \dots, t'_n)$ の二つの項にmost general unifier: $\sigma \in \text{Sub}$ が存在することである。このとき、 $(\sigma t_1, \sigma t_2, \dots, \sigma t_n)$ を t と t' の最汎単一化タプルとよび $\text{mgu}(t, t')$ と書く。

(定義4) n属性項関係(n-attribute term relation)とはn項タプルの有項集合である。

(定義5) n属性項関係T1, T2の単一化積(u-intersection)を次のように定義する。

$$u\text{-intersection}(T1, T2) = \{t : t_1 \in T1, t_2 \in T2 \text{ が存在し, } t = \text{mgu}(t_1, t_2)\}$$

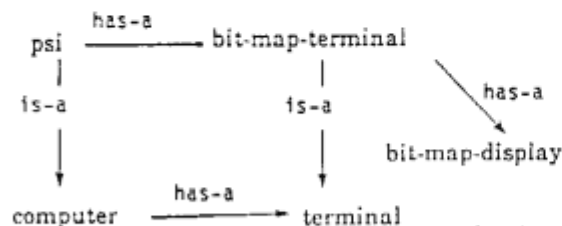


図 1

(定義6) n属性項関係T1, m属性項関係T2のi番目とj番目($1 \leq i \leq n, 1 \leq j \leq m$)の属性による単一化積合(u-join)を次のように定義する。

$$u\text{-join}_{ij}(T1, T2) = \{(\sigma t_1, \dots, \sigma t_i, \dots, \sigma t_n, \sigma t'_1, \dots, \sigma t'_j, \dots, \sigma t'_m) : (t_1, \dots, t_i, \dots, t_n) \in T1, (t'_1, \dots, t'_j, \dots, t'_m) \in T2 \text{ が存在し, かつある } \sigma \in \text{Sub} \text{ について, } \sigma t_i = \sigma t'_j.\}$$

また項関係に対する射影(projection)と和集合(union)は関係代数演算の場合と同様に定義する。項関係に対する演算を記述して項関係演算とよぶ。

項関係演算は、関係代数演算の拡張として定義されている。関係データベースモデルのタプルでは、対象とするドメインの要素は定数に号によってのみ表現された。したがって関係代数演算ではタプルの各要素に対して施される操作を表現することはできなかった。関係型知識ベースモデルでは、タプルの各要素を表現するのに項を用いることを許し、対象とするドメインの上での操作を表現できるようにしたものである。

3. DCKR 意味ネットワーク的知識表現DCKR(definite clause knowledge representation)(2)は、Prologの上にsem 述語:sem(X, P)を定義することによってis_a 階層及びhas_a 階層を表現している。sem(X, P)は直観的には「Xというオブジェクト(の典型的なインスタンス)はPという性質をもつ」と読む。例として図1に示した意味ネットワークは図2のように定義されたsem 述語によって表現される。図2で第1行目はis_a インヘリタンスの停止節である。2~6行目は図1の各枝に対応する。7行目以下はhas_a 関係のインヘリタンスに推移性をもたせるために用いられる。図3に「psiは何を持っているか?」という質問によって検索した例を示す。

ここでは、sem 述語の定義の更新の例として、図1の意味ネットワークに次の様な枝を付け加える場合について考える。

```

sem(A, is_a(A)).
sem(psi, A) :- sem(computer, A).
sem(computer, has_a(A)).
sem(psi, has_a(A)) :- sem(bit_map_terminal, A).
sem(bit_map_terminal, has_a(A)) :- sem(bit_map_display, A).
sem(bit_map_display, A) :- sem(terminal, A).
sem(computer, has_a(A)) :- sem(terminal, has_a(A)).
sem(computer, has_a(A)) :- sem(terminal, is_a(A)).
sem(psi, has_a(A)) :-
  sem(bit_map_terminal, has_a(A)).
sem(psi, has_a(A)) :-
  sem(bit_map_terminal, is_a(A)).
sem(bit_map_terminal, has_a(A)) :-
  sem(bit_map_display, has_a(A)).
sem(bit_map_display, has_a(A)) :-
  sem(bit_map_display, is_a(A)).
  
```

図 2

```

terminal  $\xrightarrow{has\_a}$  input_device
terminal  $\xrightarrow{has\_a}$  keyboard
keyboard  $\xrightarrow{has\_a}$  key

```

これらの枝に対応して次のような節が付け加えられる。
 $sen(keyboard, A) :- sen(input_device, A).$
 $sen(keyboard, has_a(key)).$

図 4

さらに has_a 関係の推移性の保持のために次のような節が

```

assert(なければならない).
sen(terminal, has\_a(A)) :- sen(keyboard, has\_a(A)).
sen(terminal, has\_a(A)) :- sen(keyboard, is\_a(A)).
sen(keyboard, has\_a(A)) :- sen(key, has\_a(A)).
sen(keyboard, has\_a(A)) :- sen(key, is\_a(A)).

```

図 5

これらの節は図 2 に対して次のプログラムを実行することによって求められる。

```

comp :-
  clause(sen(A, has\_a(B)), true),
  assert((sen(A, has\_a(C)) :- sen(B, has\_a(C)))),
  assert((sen(A, has\_a(D)) :- sen(B, is\_a(D)))),
  fail.

```

4. 関係型知識ベースの上での DCKR の実現 文献(1) で示した関係型知識ベースの上で Horn 節を表現する方法を用いて、DCKR を関係型知識ベースの上で表現し、項関係演算を用いて検索を行うことができる。

ここでは前節で示したような更新時の has_a 関係の推移性の保持が項関係の演算によって実現されることを示す。更新される前の知識ベースを項関係で表現したものを KB とする。簡単なプログラムによって図 2 の例から得た KB の例を図 6 に示す。付け加えられる枝に対応する節を表現する項関係を K1 とする。図 4 の例の場合、K1 は図 7 のようになる。このとき推移性保持の節に対応する項関係 K2 は K1 より次のような項関係演算により求めることができる。

```

KB = { ([sen(psi, A)]B), [sen(computer, A)]B),
        ([sen(bit_map_terminal, A)]B), [sen(terminal, A)]B),
        ([sen(computer, has\_a(A))]B), [sen(terminal, has\_a(A))]B),
        ([sen(computer, has\_a(A))]B), [sen(terminal, is\_a(A))]B),
        ([sen(psi, has\_a(A))]B), [sen(bit_map_terminal, has\_a(A))]B),
        ([sen(psi, has\_a(A))]B), [sen(bit_map_terminal, is\_a(A))]B),
        ([sen(bit_map_terminal, has\_a(A))]B), [sen(bit_map_display, has\_a(A))]B),
        ([sen(bit_map_terminal, has\_a(A))]B), [sen(bit_map_display, is\_a(A))]B),
        ([sen(A, is\_a(A))]B), B),
        ([sen(computer, has\_a(terminal))]A), A),
        ([sen(psi, has\_a(bit_map_terminal))]A), A),
        ([sen(bit_map_terminal, has\_a(bit_map_display))]A), A) }.

```

図 6

```

K1 = { ([sen(keyboard, A)]B), [sen(input_device, A)]B),
        ([sen(terminal, has\_a(keyboard))]A), A),
        ([sen(keyboard, has\_a(key))]A), A) }.

```

図 7

```

K2 = { ([sen(A, has\_a(B))]C), [sen(A, has\_a(D))]E), [sen(B, has\_a(D))]E),
        ([sen(A, has\_a(B))]C), [sen(A, has\_a(D))]E), [sen(B, is\_a(D))]E) }.

```

図 8

```

K2 = { ([sen(terminal, has\_a(A))]B), [sen(keyboard, has\_a(A))]B),
        ([sen(terminal, has\_a(A))]B), [sen(keyboard, is\_a(A))]B),
        ([sen(keyboard, has\_a(A))]B), [sen(key, has\_a(A))]B),
        ([sen(keyboard, has\_a(A))]B), [sen(key, is\_a(A))]B) }.

```

図 9

$K2 = projection_{1,3}$

$(u-join_{1,3}(TEMPLATE, u-intersection$
 $(K1, (([sen(X, has_a(Y))] , ([]))))$

ここで TEMPLATE は図 8 に示す項関係である。

図 7 の例から K2 を求めた例を図 9 に示す。これは図 3 に対して comp を実行した結果得られる節に対応する項関係と一致する。更新された結果に対応する項関係 KR は次の式で求められる。

$KR = union(union(KB, K1), K2)$

5. 知識コンパイラ 文献(1) では関係型知識ベースの「知識表現のコンパイル・コード」としての可能性を検査操作について検討した。本稿では新たに「知識表現の保守・更新操作のコンパイル・コード」としての可能性について見直しを得た。

関係型知識ベースは効率的な処理アーキテクチャの実現をめざすだけでなく、数学的形式化の可能性も追求しており、実用による形式化が与えられている(3)。ゆえに、高級な知識表現で書かれた知識ベース及びその上の操作を、関係型知識ベースで表現したものは、もとの知識のセマンティクスと考えることができる。我々は高級知識表現言語で記述された知識及びその操作を、項関係及びその上の演算処理に変換する知識コンパイラの開発を計画している。この知識コンパイラは高級知識表現言語と関係型知識ベースの間の逐次変換を実現したものとしてとらえることが出来る(謝辞) 有益な意見を頂いた K.B.M 会議及び知識ベース WG のメンバーの皆様には感謝します。

参考文献 (1) Yokota, H., et.al: A Model and an Architecture for Relational Knowledge Base, ICOT TR-144 (1985). (2) 田中, 小山: Definite Clause Knowledge Representation, Logic Programming Conference 85 (1985). (3) Murakami, M., et.al: Formal Semantics of a Relational Knowledge Base, ICOT TR-149 (1985).

図 6

図 7

図 8

図 9

KBMS PHI (1) 分散知識ベースシステムのシステム構成方式

伊藤 英則, 森田 幸伯, 大場 雅博, 山崎 晴明
(新世代コンピュータ技術開発機構) (沖電気工業)

1. はじめに

第三世代コンピュータシステム研究開発プロジェクトでは、知識情報処理に向けたコンピュータの開発のため、推論マシン・知識ベースマシンの実現を図っている。本研究は、複数の推論マシンと複数の知識ベースマシンがLANで結合された分散知識ベースシステムを対象とし、その実現に必要な分散制御技術や知識管理方式、協調問題解決方式などの技術開発を行うことが狙いである。

本稿では、分散知識ベースシステムとして試作を計画しているPHI (Predicate Logic Based Hierarchical Knowledge Management System)システムの構成方式について述べる。

2. 分散知識ベースシステムの研究課題

高度な知識情報処理のためには、広域に分散した大量の知識を効率よく管理する必要がある。また、知識は高価な資源であることから、その共有化も重要な要求となる。それらの要求を満たすことを目標として複数の推論マシンと複数の知識ベースマシンをネットワークで結合して分散知識ベースシステムを構築することがひとつの大きな研究課題である。

このような環境では、分散化された知識を統合化して、高度に共有化し、大量で高価な知識を高度利用することが望まれる。

そのためには、以下の条件を考慮する必要がある。

(1) 分散している知識の一貫性の保持----

知識の中には相互に関連するものがあり、それらが、ある状態 (一貫性) を保持する事を保証する必要がある。すなわち、一方を変更する場合、他方も変更を要求するような制御が必要である。

(2) ロケーショントランスペアレンシー----

分散知識ベースシステムの物理的構成の変更や、共有している知識の格納サイトの変更をおこなっても、他のユーザに影響させたくない。そのためには知識ベースシステムの物理構成を意識する事なく、知識をアクセスできることが必要である。

(3) 協調問題解決----

知識は、分散して格納されているので、一つの問題を各サイトで協調しながら解かなければならない場合がある。整合性のとれた解を得るため協調問題解決が必要である。

(4) アクセス権・セキュリティ----

知識の参照・変更は、それぞれ許されたユーザのみに限られ必要がある。

(5) ロバストネス----

あるサイトが障害が発生した際に、分散知識ベースシステム全体が機能を喪失すること

は望ましくない。障害に対する強靱性が必要である。

以上の他にも、分散知識ベースシステム特有の問題も多く存在するであろう。それらを明確にし、ユーザにとって応答性の良い、信頼性の高い、分散制御方式の体系を開発することが本研究の狙いである。

3. PHIシステムの全体構成

PHIシステムの全体構成は、図1に示す通り本稿の推論機能を果たすPSI (ホストPSI) および知識ベースマシンの機能を果たすPSI (PHI) とがICOT-LANにより結合された構成とする。

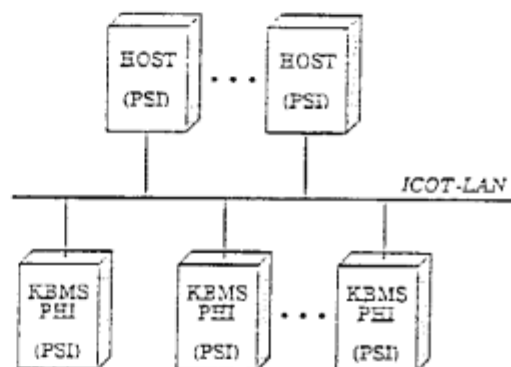


図1. PHIシステムの全体構成

我々は、前期3年間で、上記の環境を目標に、その第一歩として関係型データベースマシンDeltaを開発し、知識のうちファクトをデータベースに格納し、それと推論マシンを接続するいくつかの方式を検討した[1]。しかしながら、推論マシンと関係代数コマンドで検索するインタフェースをとると両者の間で非常に多くのインタラクションを必要とし効率的でない。PHIでは、推論マシンとのインタフェースをより高級なものにする。推論マシンとのインタフェースとしては、推論マシンが論理型言語を用いていることを考慮して、ホーン節 (図2) に設定する。

また、格納する知識の範囲も拡張し、扱う知識の対象をホーン節とする。知識ベースマシン側でもホーン節を処理することから、知識ベースマシンも推論機能をもつ必要がある。従って、PHIシステムにおいては、知識ベースマシンを推論マシンPSI上に実現する。

また、ネットワークに関しては、ICOT-LANを用いる。ICOT-LANは、その特徴として、任意のメンバからなるグループを形成しそのグループ内で同報通信を行うグループ通信機能を持つ。

PHIにおいては、知識ベース管理等の機能は、P.S.I上にソフトウェアで実現するが、知識ベース演算の高速化を図るために一部の機能をハードウェア化する。この専用ハードウェアを、知識ベース演算エンジンと呼ぶ(図3)。前期の関係代数エンジンも知識ベース演算エンジンのひとつのタイプと考

えることができる。

4. PHIの機能構成

PHIの機能構成は、知識ベース管理機能を実現する知識ベース管理部と、各サイトの知識ベース管理部を統合し、ユーザに対してひとつの知識ベースシステムとして見せる分散知識ベース制御部からなる(図4)。

PHIでは、ファクトとルールを分離して取り扱う。ファクトに対しては関係データベースを用いる。つまり、知識ベース管理部は、知識管理層と関係データベース管理層の2階層で構成する。ルールを取り扱う知識管理の機能については、さまざまなアプローチでいくつかの研究がなされているが、まだ、全容は見えていない。PHIでは、知識管理部の核として演繹データベース管理層を置き、それを基礎に知識管理方式の拡張を図っている。

分散制御部に関しては、知識ベース管理部の構成に対応して、分散データベース制御部を置き、それを核にして分散制御部の実現をすすめてゆく。

なお、知識ベース管理部の関係データベース管理層と分散データベース制御部をあわせて分散(関係)データベースを構成できる。

5. まとめ

分散知識ベースシステムPHIでは、現在関係データベース管理層、分散データベース制御部、および知識ベース管理層内の演繹データベース管理層を中心にして実現に向けた検討を行っている[3]-[5]。今後は、知識管理層および分散制御部についてさらに検討を進め、PHIの試作・評価をととして分散知識ベースシステムを実現する。

【謝辞】本検討を進めるに当たり、有益なご示唆を頂いた知識ベースマシンWGのメンバーの方々に感謝いたします。

【参考文献】

- [1] Yokota, H. et.al. : "A Model and an Architecture for a Relational Knowledge Base", submitted to the 13th Int. Simpo. Computer Architecture, 1986.
- [2] Yokota, H. et.al. : "Deductive Database System based on Semantic Resolution.", to appear in Proceeding of the 2nd Int. Conf. Data Engineer., 1986.
- [3] 伊藤, 他 : "知識ベースマシンDelta2", 昭和60年度信学全大
- [4] 宮崎, 他 : "KBMS PHI(2) 知識とデータの扱いに関する一考察", 第32回信学全大, 2M-6, 1986.
- [5] 吉田, 他 : "KBMS PHI(3) 分散知識ベース制御方式", 第32回信学全大, 2M-7, 1986.
- [6] 山中, 他 : "KBMS PHI(4) 知識ベースマシンシミュレータの機能概要", 第32回信学全大, 2M-8, 1986.

k \ l	0	1	>1
0		type 1 constant variable	type 5
1	type 2	type 4	type 6
>1	type 3		

一階法語構文

$P_1 \& P_2 \& \dots \& P_k - Q_1 \vee Q_2 \vee \dots \vee Q_l$

Deita

PHI

図2. 対象とする知識

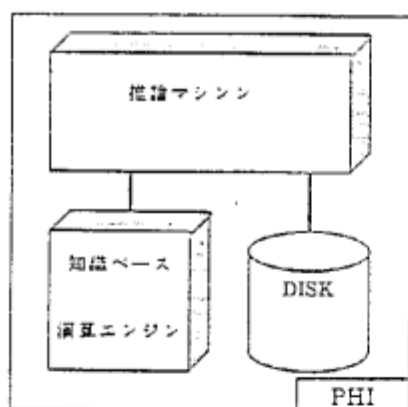


図3. PHIのハードウェア構成

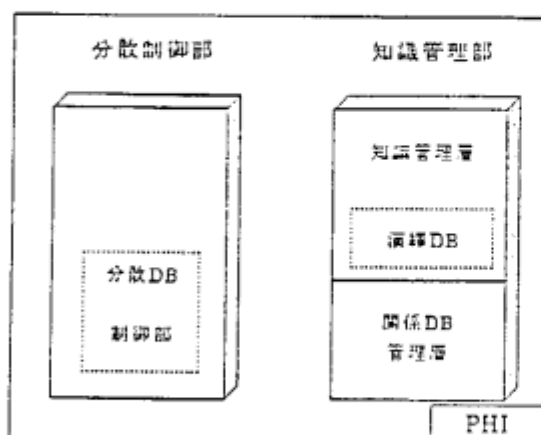


図4. PHIのソフトウェア構成

KBMS PHI (2)

2M-6 知識とデータの扱いに関する一考察

宮崎 収兄* 横田 治夫** 阿比留 幸民*

* 沖電気工業株式会社 ** 財団法人 新世代コンピュータ技術開発機構

1. はじめに

第五世代コンピュータ・プロジェクトにおける知識ベース機構の研究では前期3年間は関係データベースを基盤にした研究を行った。中期4年間の目標は知識ベースマシンプロトタイプの開発であり、知識ベースの扱いに関しては以下の2つのアプローチによる研究を行っている。

- (1) Term (項) を単位とする関係型知識モデルによる知識とデータの統一処理 [1]。
- (2) 関係データベース上に知識管理層を追加し2つの階層を合わせた知識ベース管理機構を実現。

この2つのアプローチは両者とも前期のアプローチの発展形態である。KBMS PHIでは後者を中心とした検討を行っている。この知識ベース管理機構の基本的考え方と構想を述べる。

表1. 論理型言語処理系と関係データベース管理システムの比較

操 能	論理型言語処理系	RDBMS
データ	ファクト	リレーション (タプル)
知 識	ルール	(ビュー)
異種結合性	○	×
否 定	negation as failure	無実定
異種体	○	×
変 数	○	×
問合せの表現	ゴール	問合せ言語
知識の構造化	△	×
格納/変更	プログラム/論理 アサート/リトラクト	更新言語
共有時の コンカレンシ	×	○
リカバリ	×	○
セキュリティ	×	○
インテグリティ	×	△
基 礎 実 行 方 式	ユニフィケーション	値の比較
実行方式	演算規則により 1つずつ求める (トップダウン)	関係演算による 全探索 (ボトムアップ)
実行の効率	カット等で最適可能	システムが行う
実行最適化	×	○
大規模化 (2次記憶使用)	×	○
その他の マルチメディア データ	×	△

2. 論理型言語処理系と関係データベース

使用する面からは、知識表現は問題に達したものを使うことが望ましい。他方システム作成者の面からは、できるだけ共通の処理方式によって実現できることが望ましい。また知識ベースの大規模化や共有化を図る必要がある。これらの要求を満たす知識ベース管理システムの構成を検討するため、理論的基礎が明確な論理型言語処理系と関係データベース

知識ベース管理システムに要求される機能には知識の表現能力や処理効率化、共有化などの知識管理などがある。論理型言語処理系と関係データベースの両者を比較したものはPrologに代表される論理型言語の処理系と関係データベース管理システム(RDBMS)がある。知識ベース管理システムに対する要求条件についてはこの2つのシステムを比較したものを表1に示す。ここで表現能力の比較はホーン型と関係データベースの自然対照関係、即ちホーン型のファクトと関係データベースのタプルの1対1対応を前提とした。

表1から明らかなように、表現能力の面では論理型言語処理系が、管理機能の面ではRDBMSが優れている。処理方式による差異についてはそれぞれ利点損失があるが知識ベースが大規模な場合はRDBMSが優れている。これは超大型Prolog等では処理の効率が悪化しているのに対し、RDBMSでは高度な最適化が可能なことや、2次記憶に対する効率的アクセスメソッドが使用可能による。また処理処理による最適化を行うシステムでも処理対象の範囲を限定して最適化を図れるRDBMSの方が全体の計算量を減らせる可能性が大きい。

3. 知識ベース管理システムの構成

前期の考察から知識ベース管理システムの構成として論理型言語によるプログラムとRDBMSの利点を生かし、両者を結合したものを近似的モデルとして考えることができる(図1)。このモデルを実現するシステムは1COTなどで試作されたが[2]、以下のような問題が指摘されている。

- (1) ニュープログラムと知識ベース管理システムの関係の明確化。
- (2) 知識の構造化や他の知識表現との関係。
- (3) 知識の管理、特に知識の検索方法や知識のインテグリティ管理。
- (4) 実行効率の改善。最適化処理や制約条件の実行効率のメタ知識による問題。

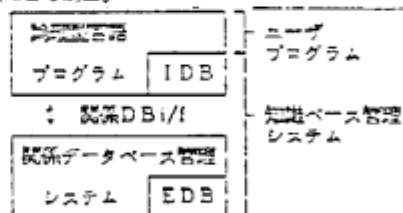


図1 論理型言語-関係データベース管理システムによる知識ベース管理システムモデル

これらの問題を解決するには知識ベース管理システムでデータ以外の知識も管理することが必要である。KBMS PH1では知識ベース管理システムをデータベース管理層と知識管理層の2階層から構成する(図2)。

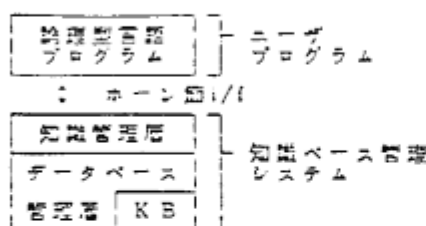


図2 KBMS PH1のモデル

知識ベースはインテンションデータ部(ldb)、エクステンションデータ部(edb)およびこれらを管理するメタ知識よりなりデータベース管理層に格納される。ldbはホーン節のルールに、またedbはホーン節のファクトに対応するものをそれぞれリレーションに格納したものである。

データベース管理層は知識ベースの検索や管理を関係データモデルに従って処理する。データベース管理層は関係演算を基本とするためldbの内容についても単純なデータとして扱い、ホーン節ルールとしての意味をもった処理は行わない。データベース管理層での処理はホーン節処理のうち以下の制限をつけたものと同等である。

- (a) (構文体を含まない) 純ホーン節の述語および not、比較、算術演算のみを含む。
- (b) 再帰定義を含まない。
- (c) 任意のゴールに対し解は、"yes/no"、またはファクト中に現われる値の組み合わせの集合のみである。

知識管理層は知識ベース管理システムに対する要求を処理するのに必要な知識をldbより取りだし推論を用いた処理を行う。この処理に必要なデータがedbにある場合は関係演算問合せによりデータをデータベース管理層に要求する。知識管理層の機能には以下のものがある。

- (a) 演算データベース機能(再帰的定義を含むホーン節問合せ(+ldb)の関係演算への変換およびデータベース管理層による検索実行の制御)。
- (b) 構文体を含むファクトとedbのリレーションとのフォーマット変換。
- (c) 検索空間の制限や最適化条件検索の処理および最適化。
- (d) ldbのリレーションへの格納と管理、ldbの検索、及び検索に用いたldb内容の出力。
- (e) 知識ベース管理に必要なシステム述語の処理

例えば祖先リレーションancestorの処理はPH1では概念的に以下になる。

【プログラムと知識ベース】

プログラム: ... get(ancestor(X, Y)), ...
 ldb : ancestor(X, Y) :- parent(X, Y),
 ancestor(X, Y) :- parent(X, Z),
 ancestor(Z, Y),
 parent(X, Y) :- edb(parent(X, Y)).
 edb : parentリレーション

【処理】

- ① プログラム実行時の問合せ発行、
"?-get(ancestor(taro, Y))".
- ② 知識管理層による関係演算生成

$$\text{ancestor}_{\text{db}} =$$

$$\pi_{\text{db}}((\sigma_{\text{db}}(\dots, \text{parent})) \cup$$

$$\pi_{\text{db}}(\sigma_{\text{db}}(\sigma_{\text{db}}(\dots, \text{parent}) \times_{\text{db}} \text{ancestor})))$$
- ③ 知識管理層の繰り返し制御下でのデータベース管理層による解の生成
- ④ 知識管理層での後処理後、解のプログラムへの受渡し

4. PH1の仕様

KBMS PH1では知識ベース管理システムを階層化し、知識ベースをldbとedbに分離して管理する。この結果論理型言語処理系のもつ知識表現機能とRDBMSの持つ管理機能の利点をあわせたシステムを構成可能である。さらに処理効率の面でも知識の性質に適した処理方式が適用できる。PH1では知識ベース処理の階層化により以下の特徴の実現を図っている。

- (1) 知識ベースのプログラムからの分離による知識の共有化・大規模化の実現。
- (2) 論理型知識表現を基盤とし、12ホーン節に加え構文体処理や検索制御機能を実現する。
- (3) ldbに対するトップダウン処理とedbに対するボトムアップ処理の組み合わせによる処理効率化。
- (4) ldbとedbの論理的分離による分散知識ベース処理の容易化。
- (5) ホーン節インタフェースによるPSI(ホスト)プログラミング言語との統一、およびユーザプログラムと知識ベース管理システムのインタラクション減少。(知識ベースマシンや分散知識ベース管理システムではインタラクション減少による効率化が必須条件である。)
- (6) データベース管理層によるldbレポジトリ込みを用いた知識管理層の検索空間の削減。
- (7) ldbとedbをあわせた知識ベースのインテグリティ管理の実現。

5. まとめ

KBMS PH1における知識ベースの扱いとシステム構成の基本的考え方を述べた。現在データベース管理層および知識管理層の核部分(演算データベース+構文体+検索制御)の実現方式の検討を行っている。今後はこれらの詳細化および知識ベースの構造化やldbの分散化などの検討を行う予定である。

【参考文献】

- [1] 伊藤他: "知識ベースマシンDelta2", 昭和60年度信学会情報システム部全国大会, 昭和60年11月
- [2] R.Yokota et al.: "An Enhanced Inference Mechanism for Generating Relational Algebra Queries", Proc. 3rd ACM Symp. on Principles of Database Systems, Apr. 1984

KBMS PHI (3) 分散知識ベース 制御方式

2M-7

吉田 勇*

高杉 哲朗*

大場 雅博**

*沖電気工業株式会社

**財団法人 新世代コンピュータ技術開発機構

1. はじめに 知識情報処理に不可欠な技術として大量な知識の効率的な管理及びその共有化がある。これらの技術課題の解決の課題技術としてICOTでは、複数知識ベースの統合化技術の確立を目指しPHIの試作を行っている[1]。現在開発等はPHIの知識ベースのうちデータ部分の分散化のための方式検討を行っている。その具体的な課題には、ロケーション・トランスペアレンシー、サイトの自律性、分散問合せ処理、分散サイト間のコンセンサレンシー制御などがある。本稿では上記項目を中心にPHIの分散知識ベース制御方式について報告する。

2. PHIの構成 PHIは物理的には複数のサイト(逐次型推論マシンPS1)とこれらを結合するブロードキャスト機能及びグループ内の通信グループを形成する機能を有する高速ネットワーク(LAN)とからなる[2]。また論理的には上記サイト内のトランザクション・マネージャ(TM)とデータ・マネージャ(DM)からなる(図1)。TMとDMの両方が1つのサイトにあっても良い。PHIではサイト間の通信目的に合せ以下の3つの通信グループを形成する。(1)システム通信グループ、(2)トランザクション通信グループ及び(3)DM間セッションである(表1)。これら3つの通信グループを適宜使用しトランザクション処理を行う。1つのトランザクションはユーザ端末が接続されているサイト(ユーザサイト)内のTMとそのトランザクション内で使用されるリレーションを含むサイト内のDMとで構成されるトランザクション通信グループ内で処理される。

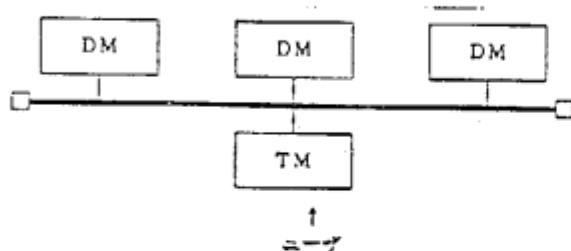


図1 PHIの論理構成

3. PHIの分散知識ベース制御方式

3.1 ディレクトリ管理 PHIでは知識ベースのデータ部分はリレーショナルモデルによって扱う。またデータの分散配置はリレーション単位とし、重複しての配置は行わない。ロケーション・トランスペアレンシーを実現するため、各々のリレーションにはユーザ毎に定義でき

表1 通信グループ

名称	メンバー	形成時期
システム通信グループ	システム内の全サイト	立ち上げ時
トランザクション通信グループ	当該トランザクションを受け付ける1つのTMと関連するDM	トランザクション発生時
DM間セッション	1対のDM	立ち上げ時

る、ユーザ定義名と(格納サイト識別子-サイト内管理識別子とからなる)システム内一意のシステム内IDとの対応表(ディレクトリ)が付与される。このためユーザ毎に各自のリレーション名が定義できる(図2)。図2において3つのリレーションR₁, R₂, R₃がサイトS₁に格納されており、ユーザ1とユーザ2に同じリレーションに対して異なるユーザ定義名を付与している。このディレクトリは各サイトが個々に自身のリレーションについてのみ管理する、非重複分散管理であるとする。この方式では、LANのブロードキャスト機能を活用することにより管理オーバーヘッドが少なく、且つ高度なサイト自律性を提供することが可能となる。

ユーザ定義名	システム内ID	ユーザ定義名	システム内ID
従業員	S ₁ ・R ₁	EMP.	S ₁ ・R ₁
部門	S ₁ ・R ₂	DIST.	S ₁ ・R ₂
給料	S ₁ ・R ₃	SALA.	S ₁ ・R ₃

ユーザ1

ユーザ2

図2 ディレクトリ例

3.2 トランザクション管理 通常、トランザクション実行時にはシステム内全サイト中、実行すべきトランザクションに関連する(該当するリレーションを保持する)サイトはユーザサイトを含み幾つかのサイトに限定される。これら関連するサイトで構成されるグループで1つのトランザクション処理を行うことにより他サイトに対するブロードキャスト通信時の余分な負荷を軽減することが可能である。PHIではこのグループをトランザクション通信グループとしてトランザクション開始時(スタート・トランザクション時)にユーザサイト内TMの指

示により形成する。このトランザクション通信グループは問合せ処理が終了するまで保持され、終了と同時にユーザサイト内TMの指示により消滅する。

トランザクションの管理はスタート・トランザクション、エンド・トランザクション及びアポート・トランザクションにより規定される。

3.3 分散問合せ処理 ユーザからの問合せはTMからトランザクション通信グループに対してコマンド木群としてブロードキャストされる。この問合せは単一サイト内で処理可能なものと、該サイトにはまだ処理（分散問合せ処理）とに分類される。前者についてはサイト内で閉じた処理であるため一意の処理であるが、後者についてはサイト間のリレーションの転送等を必要とするため、分散データベースにおける問合せ処理の性能向上の端にはこの分散問合せ処理の最適化が不可欠である。従来この最適化についてはセミジョイン方式などが用いられているが[3]、PHIでは事前にセミジョイン結果等の情報を管理していないため従来の方法では問合せを効率よく処理できない。そこで事前に得たリレーション情報をもとに問合せ処理時に最適な処理手順を決定する、以下の2つのダイナミックなアルゴリズムを検討した[4]。

(1) 個々のサイトが各分散問合せ中のコマンド木節にその対象のリレーションのサイズを比較することによりリレーションの転送方向を決定する。

(2) 分散問合せを複数のステージからなるコマンドグラフに展開し、各ステージにおいてサイト間で最適な転送プランを作成する。この方式(2)はさらに細かく分類でき、a)は各ステージ内全体に対する最適なリレーション転送プランを各サイトが個々に作成するもので、b)は自サイトに関するコマンド木のみの最適な転送プランを作成しその後ステージ内で各転送プランが矛盾しないことを確認するものである。

この(1)、(2)のいずれかで作成したリレーション転送プランを各サイトにおいて実行する。

3.4 コンカレンシー制御 従来よりコンカレンシー制御には、2フェーズロック方式とタイムスタンプ方式とがあるが[5]、PHIにおいてはこれらの内ブロードキャストLANの機能を活用するため2フェーズロック方式、特に事前宣言を用いた方式を採用する。

デッドロックの対処としてはDetection法、Prevention法、Avoidance法があるが、PHIにおいてはブロードキャスト機能を主に用いるのでメッセージの到着順序の歪転が稀である為これらの中でDetection法を採用する。Detection法は更に細かく集中型、分散型、階層型に分類できるが、単一型LANであり分散システムである点を考慮すると分散型がよく適合する。PHIにおけるコンカレンシー制御は具体的には、スタート・トランザクシ

ョン時にユーザサイト内TMはトランザクション内で使用するリレーション名を付加してシステム内全体にブロードキャストする。システム内全サイトはこのメッセージの内容を解析して、自サイト内に該当するリレーションがある場合該当するリレーションに対してロックを掛けその後、このリレーションに関する情報（ディレクトリ、そのリレーションのサイズ等）をユーザサイト内TMのみならず、システム内全サイトにブロードキャストする。ユーザサイト内TMはこのメッセージの内容をもとに必要とするサイトのみでトランザクション通信グループを形成する。このサイトがブロードキャストするリレーション情報により各サイトは他サイト内でのトランザクションのリレーション使用状況及びロック待ち情報を知ることができる。それによって各サイトがグローバルなトランザクション待ち状態グラフ(TWFG)を作成・管理し、グローバルなデッドロックを検出することができる。デッドロックを検出した場合、ひとつのトランザクションをアポートすることによりデッドロックを解消する。問合せ終了時（エンド・トランザクションもしくはアポート・トランザクション時）にはユーザサイト内TMは、トランザクション終了もしくは中止の旨をシステム全サイトにブロードキャストし、その後この通信グループの消滅を行う。

4. まとめ PHIにおける分散知識ベース制御方式についてデータ分散化を中心に報告した。本方式によればブロードキャスト通信機能を有効に利用することにより効率的な問合せ処理、コンカレンシー制御等が可能となる。現在は機能設計段階であり、今後はさらにこれらの事項について詳細化するとともに、知識全体の分散化の制御方式についても検討を加えてゆく予定である。

【参考文献】

- [1]伊藤 他: "KBMS PHI (1) 分散知識ベースシステムの構成方式", 第3回情報処大, 1986年3月
- [2]A. Taguchi et al.: "INI: Internal Network in the ICOT Programming Laboratory and Its Future", IC CC, Sydney, Australia, Oct. 1984
- [3]B. Williams et al.: "Introduction to a System for Distributed Databases (SDD-1)", ACM TODS, Vol. 5 No. 1 pp. 1-17, 1980
- [4]吉田 他: "KBMS PHIにおける分散問合せ処理方式", 昭和61年度信学会大, 1986年3月
- [5]C. Mohan: "Recent and future trends in distributed data management", Proc. NYU Symposium on New Direction for Data Base Systems, May 1984

KBMS PHI (4) 知識ベースマシン シミュレータの機能概要

2N-8

山中 幸隆*

山下 祥司*

森田 幸信**

* 沖電気工業株式会社

** 財団法人 新世代コンピュータ技術開発機構

1. はじめに

第三世代コンピュータ・プロジェクトでは既に前期過程を終り、その成果として逐次型推論マシンPSIと超大量データベースマシンDeilaが発表されている[1][2]。中期4年間に於いては知識ベースマシンのプロトタイプの開発がひとつの目標となっているが、筆者等はこの目標に於いてKBMS PHIの開発を進めている[3]。そして、現在知識ベース管理系の基本となるデータベース管理系として、Deilaの機能に準じた関係データベース管理システム(RDBMS: PHIでは知識ベースマシンシミュレータと呼んでいる)の開発を行っている。本稿ではこのRDBMSの機能概要について述べる。

2. PHIと関係データベース

PHIでは、

- ・プログラムからの知識の独立
- ・知識の共有
- ・大量知識の取り扱い

等を図っていて、その知識ベース管理系の構成を、

- (1) 知識管理層 (知識部)
- (2) データベース管理層 (基本部)

の2層に分けている。知識管理層では論理型言語のルールに対応するインテンショナルな知識の管理を行う。データベース管理層では知識の大量化を効率的に処理する為、ファクトに対応するデータを管理すると共に、インテンショナルな知識の二次記憶への格納も行う。尚、データベース管理層は単体でも動く事ができる。以下、論点を絞る為、データベース管理層が単体で動く場合を想定して話を進める。

3. システム構成

想定しているシステム構成を図1に示す。図1で、ホスト・マシン、RDBMSにはいずれも逐次型推論マシンPSIを想定していて、それぞれは高速度LANを介して接続されている。

4. コマンド・インタフェース

RDBMSとホスト・マシン上のホストとの間の論理的なインタフェースは、基本

的には、RDBMSがホスト側からいくつかのコマンドを受け、RDBMS内で処理を行い、そのレスポンスをホスト側へ返す事により構成される。

RDBMSの処理においては、次の様に処理単位を階層化している。

- (1) コマンド
関係代数を中心としたRDBMSの処理を指示する語句単位。
- (2) コマンド木
1つの意味のある結果を得る為のコマンドの集り。
- (3) コマンド木群
コマンド木の集り。
- (4) トランザクション
データベースの一貫性を保つ為の単位。1つもしくは複数のコマンド木群よりなる。トランザクション終了と共にロックも解除される。

5. E S P

RDBMSの記述は、Prologにオブジェクト指向の機能を付えたESP(Extended Self-contained Prolog)言語により行う。ESPはICOT(新世代コンピュータ技術開発機構)により開発された言語であり、RDBMSのインプリメントを行うPSIのシステム記述言語となっている為である。ESPを採用する事により、上位レイヤの知識管理層とのインタフェースが明快になり、上位レイヤの変更に関して大きな柔軟性を持つ利点が生じる。又、Prolog特有のユニフィケーション、バックトラック等の機構により、システム開発も容易になる。

6. リレーションとデータタイプ

RDBMSでは関係代数的処理を提供する為、ユーザが扱うリレーションを以下の3種類に分けている。

- (1) PR (Permanent Relation)
ユーザにより正規に定義され、トランザクション終了後も残る。
- (2) IR (Intermediate Relation)
ユーザの指定は名前のみ行われる。コマンド木に1つ対応し、トランザクション終了時に消去される。

(3) TR (Temporary Relation)

コマンド木の中だけで定義され、名前指定はない。コマンド木中の他コマンドからそのTRを生成したコマンド番号によって指示される。

又、RDBMSでは、データタイプとして文字型、数値型、無型を扱う。文字型には2バイト漢字JISと1バイトASCIIがあり、数値型には4バイト及び8バイト整数と浮動小数点実数がある。無型は演算対象として何にでもマッチする。

7. 機能

以下RDBMSの機能概要を述べる。又、一部のコマンド例を表1に記す。

7.1 データ操作機能

(1) 検索

インデクスとしてハッシュ、B+木が利用できる。コマンド木は内部で自動的に最適化される。

(2) 更新

更新情報をリレーションとしてまとめ、リレーション単位で更新を行う。これにより、タプル検索の処理効率を高める事ができる。

(3) 入出力

ホストとRDBMSとのデータ転送を行う。一回当たり最大4000バイトのデータ転送が可能である。

(4) その他の機能

ソート、コピー等を用意している。

7.2 データ定義機能

リレーションやインデクスの登録、削除等を行う。

7.3 非同期処理

RDBMSの異常状態に対処するもので、他の処理の実行状態にかかわらず優先して非同期な処理を行う。システムの稼働状態検査やトランザクション強制アボート等がある。

7.4 セキュリティ

セキュリティ機能として、ユーザのデータベース使用権をユーザ名とパスワードによりチェックする。パスワードはユーザが任意に変更できる。又、リレーション使用権をリレーション特有のアクセス許可モードによりチェックする。

ユーザはユーザグループとしても登録できるので、幅広い利用が可能となる。又、アクセス権付与コマンドにより、特定のユーザだけにアクセスを許す事ができる。

7.5 コンカレンシー制御

データの一貫性を保証する為、ロック

機構によるコンカレンシー制御を行う。ロックの単位はリレーションであり、範囲はリレーション群である。但し、スキーマに関しては階層管理を行ない、ロック単位をタプル群とする。これによりリレーションが過半数検索更新処理と定義処理の競合を許す事ができる。

7.6 リカバリ

データベース利用に関して、トランザクション障害、メディア障害等の障害からデータを守る為、シャドウ方式に基づく二重化管理を行い、トランザクション開始時の状態へ回復させる。

7.7 バッファリング

メモリ内での処理に比べて二次記憶とのI/O処理は負荷が大きい。そこで、二次記憶とのI/Oを少なくする為にLRU機構に基づくバッファリングを行う。

8. おわりに

RDBMSの機能概要について述べてきた。今後は、より詳細な検討を深め、より高性能なシステムの構築を目指す予定である。

【参考文献】

- [1] 角田 他: "RDBM "Delta" (I) ~ (III) ", 第26回情報全大, 1983
- [2] 西川 他: "パーソナル逐次型推論マシンφ", 第27回情報全大, 1983
- [3] 伊藤 他: "XBMS PHI (1) ~ (3) ", 第32回情報全大予稿集, 1986

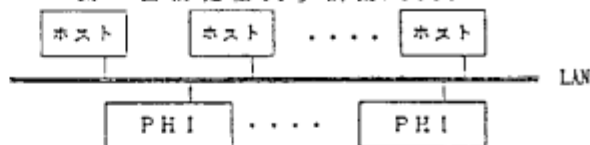


図1. システム構成

機能	コマンド
検索	PROJECTION SELECTION JOIN UNION COUNT SUMMATION ADDITION EQUAL
更新	INSERT DELETE UPDATE
入出力	GET GET-NEXT GET-END PUT PUT-NEXT
特殊操作	GROUP-BY SELECT-GROUP COPY SORT
データ定義	CREATE-RELATION PURGE-RELATION
アクセス権制御	REGISTER-USER AUTHORIZATION
トランザクション制御	START-TRANSACTION END-TRANSACTION
セッション制御	BEGIN-SESSION END-SESSION
非同期	SENSE-STATUS ABORT-PROCESSING

表1. コマンド例

50-5

マルチポートページメモリを用いた 大規模知識ベースの制御方式

物井 秀俊、横田 治夫、西田 健次、伊藤 英則 酒井 浩

(財)新世代コンピュータ技術開発機構 (株)東芝

1. はじめに

第5世代コンピュータシステム研究開発プロジェクトでは、知識情報処理向けのコンピュータシステム開発のため、推論・知識ベース機構の実現を図っている。我々は、その基本要素の一部として知識ベースマシンを開発中である。

知識ベースマシンは、大量の知識をディスク等の二次記憶装置に格納し、ホストから要求された知識の検索を高速に実行する専用マシンとして位置付けられる。このようなシステムでは、処理装置と主記憶間のボトルネックだけでなく、主記憶と二次記憶間のボトルネックが性能上の大きな問題となる。

我々は、この主記憶と二次記憶間のボトルネックを解消するため、複数の二次記憶装置およびデータ処理装置から同時にアクセス可能なマルチポートページメモリ(MPPM:Multi-Port Page-Memory)の、経済的かつ実用的な構成法を提案した[1]。さらに、大量のデータを扱うときに、ネックとなる処理に専用ハードウェア(エンジン)を用いる方式は、データベースマシンの一つの流れとなっている[2]。

本稿では、このMPPMと専用ハードウェアを用いた知識ベースマシンの制御方式について述べる。

2. 関係型知識ベース検索

我々が検討中の知識ベースマシンでは、項関係を構成要素とする関係型知識ベースを検索の対象とする。項関係とは、関係データベースにおける関係操作の対象を項、つまり変数を含む一種の構造体拡張したものである[3]。

関係型知識ベースに対する演算としては、項間の単一化を基本操作とする単一化結合と単一化制約および射影がある。これは、従来の関係代数の結合や制約演算の等号条件を単一化に置き替えた操作である。

3. アーキテクチャ

複数のプロセッサをバッファメモリを介して二次記憶と結合する場合、一本のデータバスで結合したのでは、データバス上での競合が頻発し、思うような性能を得ることができない。このようなデータ転送時の競合を抑えて高速なデータ処理を可能にするため、バッファメモリとしてMPPMを使用する。

検討中の知識ベースマシンの全体構成を図3.1に示す[4]。各要素は、以下のような動作をする。

・UE-項関係演算を実行するデータ処理専用ハード

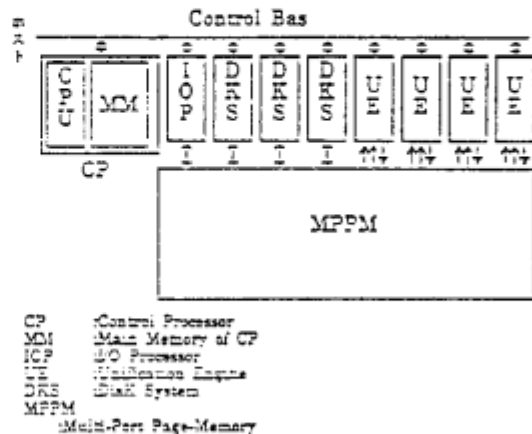


図3.1 知識ベースマシンの全体構成

ウェア[5]。MPPMからデータをストリームとして入力し、処理結果をストリームとしてMPPMに出力する。

・DKS-二次記憶とMPPM間のデータ転送を実行。
・IOP-ホストマシンまたはCPのMMとの間のデータ転送を実行。
・CP-システム全体の制御とホストマシン間のインタフェース制御を実行。

ここで、UEとDKSは独立にMPPMのページがアクセス可能なため、UEを実行させる一方でDKSによるデータの先読みを行うことができ、制御方式によっては効率的なデータ処理が可能になる。

知識ベースマシン全体の制御とホストとの間のインタフェース制御は、CPで行う。UE、DKSおよびIOPの各プロセッサは、CPに対する入出力装置として動作し、コマンドと割り込みにより制御を受ける。ここで、CPと各プロセッサ間の制御情報のやりとりは、コントロールバスを通して行われる。

4. 並列制御方式

機能分散型のマルチプロセッサデータベースマシンでは、各プロセッサのスケジューリングが、性能上の問題となる。このスケジューリング方式としては、例えばDIRECT[2]に見られるように、各プロセッサにスケジューリング機能を持たせる方式が代表的である。しかし、この方法ではトラフィックが増えた時、プロセッサの割当てオーバーヘッドが増加する恐れがある。

我々は、MPPMに接続された各プロセッサを入出力装置として動作させることにより、CPにおいて事象駆動型のスケジューリングを行い、効率のよい並列処理を行うことを考えている。

先ず $T_2 \leftarrow T_{body} \otimes head(T)$ (注) という単一化結合演算を考える。ここで、入力となる項関係 T, T_0 が

$$T = \{PT_1, PT_2, PT_3, PT_4\}$$

$$T_0 = \{PT_{01}, PT_{02}, PT_{03}\}$$

というページに分割でき、その内 PT_2, PT_{01}, PT_{03} が MPPM 上にステージング済みであるとする。この時、単一化結合はページ毎の演算を UE で行い、図 4.1 に示すような手順で実行できる。

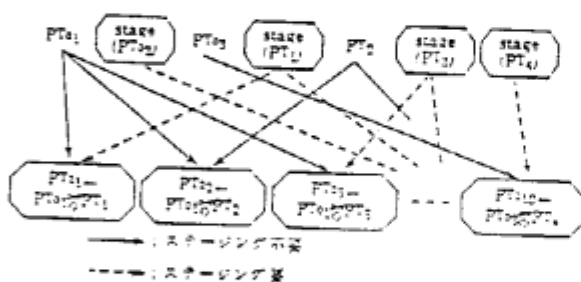


図 4.1 単一化結合の実行手順

図 4.1 の各ノードにおいて、 PT_2, PT_{01}, PT_{03} だけを使う単一化結合は、UE で即実行可能である。しかし、他のノードでは DKS におけるステージングの完了を待ち合わせる必要がある。このような待ち合わせを可能にするため、以下のバケット形式の命令を考える。

条件	演算	後処理
----	----	-----

ここで各部は、以下のような目的をもつ。

・条件部 — バケットが実行可能となる条件。

・演算部 — UE で実行する演算と使用するページ。

・後処理 — 演算結果を他のバケットに反映。

このバケットを用いると図 4.1 の処理は図 4.2 のように書くことができる。ここで、 $S(P)$ はページ P の状態を示し、 $S(P)=T$ の時ページ P がアクセス可能、 $S(P)=F$ の時アクセス不可能を示し、初期値は F とする。また、 $stage(P)$ は、ページ P の DKS によるステージングを示す。

プロセッサ	条件	演算	後処理
DKS	T	$stage(PT_1)$ $stage(PT_3)$ $stage(PT_4)$ $stage(PT_{02})$	$S(PT_1) \leftarrow T$ $S(PT_3) \leftarrow T$ $S(PT_4) \leftarrow T$ $S(PT_{02}) \leftarrow T$
UE	$S(PT_1)=T \& S(PT_{01})=T$ $S(PT_2)=T \& S(PT_{02})=T$ $S(PT_3)=T \& S(PT_{03})=T$	$PT_{01} \leftarrow PT_1 \otimes PT_{01}$ $PT_{02} \leftarrow PT_2 \otimes PT_{02}$ $PT_{03} \leftarrow PT_3 \otimes PT_{03}$	$S(PT_{01}) \leftarrow T$ $S(PT_{02}) \leftarrow T$ $S(PT_{03}) \leftarrow T$
	$S(PT_4)=T \& S(PT_{04})=T$ all $S(PT_{0i})=T$	$PT_{04} \leftarrow PT_4 \otimes PT_{04}$; Nop	$S(PT_{04}) \leftarrow T$ $S(T_2) \leftarrow T$
CP	$S(T_2)=T$	; <Terminate Process>	

図 4.2 バケットによる実行手順

(注) head と body という二つの属性を持った項関係 T, T_0 に対して、組毎に単一化操作を行い、単一化できた組の集合を項関係 T_2 に格納する。

以上に述べたバケットは、演算を実行する CP、UE および DKS 毎に分配し、プロセッサ毎に実行の可否を判定することができる。ここで実行可となった処理は、他の処理とは独立に DKS や UE を動かして演算部で示された処理を実行し、処理結果を他の処理に伝達する。この制御に対応するフローを、図 4.3 に示す。

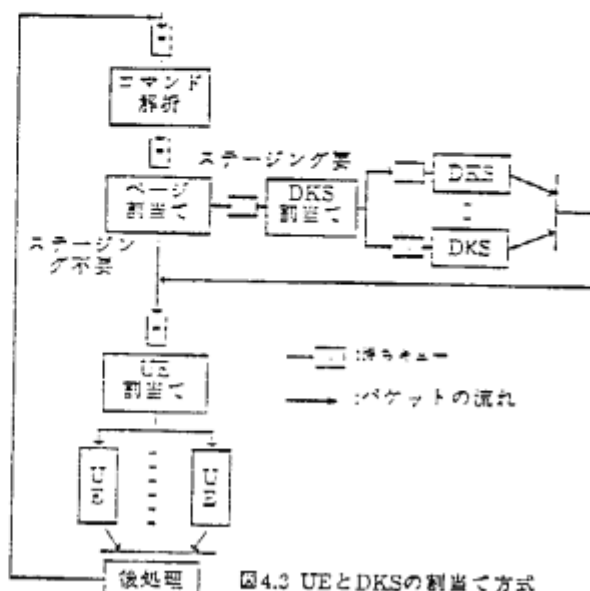


図 4.3 UE と DKS の割当て方式

5. まとめ

以上、MPPM と UE を用いた、知識ベースマシンの制御方法について述べた。マルチプロセッサ方式の関係データベースマシンでは、トランザクション毎、関係毎、また関係を形成するページ毎と、並列化のレベルが色々考えられる。本稿では、MPPM の特徴を生かすページ毎の並列処理について述べた。特に、事象駆動型の並列制御方式を取ることで、プロセッサの割当てオーバーヘッドの削減が図れ、効率のよい並列処理が可能となる。今後は、さらに制御方式の詳細化を図り、シミュレータを用いて、処理の単位や並列制御について、定量的な検討をすすめてゆく。

[謝辞] 本検討に当たり、有益な御示唆を頂いた知識ベース WG のメンバーの方々に感謝いたします。

【参考文献】

- [1] 横田, 他: マルチポートページメモリの構成方法, 昭和 60 年度通信学会部門別全国大会。
- [2] Borai, H., et al.: Implementation of the Database Machine DIRECT, IEEE Trans. on SE, Vol. SE-8, No. 6, pp. 533-543.
- [3] 横田, 他: 単一化による知識ベース検索, 昭和 60 年度通信学会部門別全国大会。
- [4] 森田, 他: 単一化結合の処理方式, 情報第 32 回全国大会予稿, 1M-7
- [5] 横田, 他: 単一化エンジンを用いた知識ベースマシンの構成方式, 情報第 32 回全国大会予稿, 1M-8

GHCプログラムの 2G-5 アルゴリズム・デバッグについて

竹内彰一

(財) 新世代コンピュータ技術開発機構

1 はじめに

計算機アーキテクチャの発展に伴い、並列計算機が実用化されていくに従って並列プログラムのデバッグはソフトウェア開発上極めて重要な技術となることが予想される。一方で並列プログラムのデバッグは逐次プログラムのそれと比較すると非常に困難な仕事であると言われてきた。その理由としては、(1) 複数の計算が並列に行なわれていること、(2) 並列に行なわれている計算が互いに作用を及ぼし合うこと、(3) 逐次プログラムには無いデッドロックのような新種のバグが出現することなどがある。

通常、プログラムのデバッグは実行トレースを基本とし動作を理解する過程で行なわれるが、この方法では多数の計算が錯綜している並列計算のデバッグの際、プログラムの負担が増大し並列プログラミング自身の専門化を阻害かねない。これに対しデバッグのアルゴリズム(バグの存在する範囲の絞り込み)が相違しているアルゴリズム的なデバッガは、プログラマが必要最小限の情報を与えてやればデバッガ自身がバグの位置を同定することができるという点で、プログラマの能力によらず並列計算の複雑さにも対応できる着実に負担の軽い手法である。

Shapiro は一連の質問をプログラマに発しながらバグを検出するアルゴリズム "Divide and Query" について報告している[2]。本稿ではこの "Divide and Query" アルゴリズムを拡張したものに基づいた並列論理型言語GHC[1]のためのアルゴリズム的なデバッガについて報告する。このデバッガは必要最小限の情報だけをプログラマに提示し、それに対するプログラマの応答に基づいて着実にバグの存在する範囲を決めていく。またこのデバッガはGHC自身により関係に実現されている。

2 GHCプログラムのアルゴリズム・デバッグ

GHCはユニフィケーションに基づいた非決定的な並列プログラミング言語である[1]。本デバッガは以下に示すような症状を持つGHCプログラムが与えられると、プログラマとの一連の質問応答の後バグを含んだ節のインスタンスを探す。

①ゴールの計算は成功して終了するが出力値(出力ユニフィケーション)が間違っている。

②ゴールの計算がデッドロックに陥る。

プログラマとの質問応答は行なわれた計算がプログラマ

の意図するものであったかどうかに関するものである。意図と異なる部分が実際に行なわれた計算の中にあれば、そこにバグを含んだ節を用いた計算が行なわれたことになる。プログラムSに対するプログラマの意図としては

(ア) S上で成功終了するすべての計算の入出力関係の集合

(イ) S上でサスペンドするすべての計算の入出力関係の2つを考え、これをM(S)で表わす。M(S)は非決定的な並列プログラムの意味としては不十分であるが、①、②の症状のデバッグには十分である。M(S)は常にすべてを書き表わす必要がなく、デバッグを行なうプログラマの頭の中にあって、具体的な計算(終了しているか、サスペンドしている)の示す入出力関係がM(S)に含まれるか(意図に合うか)、含まれないか(意図に合わないか)をプログラマが判定できるだけで十分である。

本デバッガではまず最初に①、②の症状を示した計算を木の形で再現する。この木はGHCプログラムのガード部を省略したAND木になっており、計算木と呼ぶ。計算木は次のような性質を満足する。

(i) 計算木の各ノードは exit(ゴール) または suspended(ゴール) のいずれかである。

(ii) Tを計算木、L0をそのルート、T1,...,Tnをルートのすぐ下の部分木、L1,...,Lnをそれらのルートとすると

$$G_0 := G_1, \dots, G_n$$

はプログラム中に存在するある節のインスタンスである。(ただし、GjはLj=exit(Gj)あるいはLj=suspended(Gj)を満足するものである)

exit(G)のGはこのノードをルートとする成功した計算の入出力関係を表わしており、suspended(G)のGは未だ終了していない計算の部分的入出力を表わす。計算木中のノードNの重みw(N)は次のように定義される。

$$w(N) = 0 \quad \text{if } N \text{ がシステム述語に対応する,} \\ = 1 + w(N_1) + \dots + w(N_n) \quad \text{otherwise}$$

ただし、Ni (i=1~n)はNの子ノードである。デバッグのアルゴリズムを以下に示す。

アルゴリズム

入力: プログラム S, ゴール P

M ← M(S) (最初は空)

出力: S の節のインスタンスである次のような節

$A: -B1, \dots, Bn.$

ただし, $A \in M(S)$, $B_i \in M(S)$

ステップ1

ゴール P の計算をシミュレートし計算木を生成する。

ステップ2

計算木のルートの重み w を M^* を法として計算する。

もし $w = 1$ ならば

$\langle \text{ルート} \rangle : -Q1, \dots, Qn.$

を出力として返す。ただし $Q1, \dots, Qn$ はルートの子ノード。 $w > 1$ の場合は計算木の中でその重みが $\lceil w/2 \rceil$ を超えない最大のノード m を求め、 $m \in M(S)$ かどうかをプログラマに問合せる。答が $m \in M(S)$ の場合は同じ計算木について M^* を $M^* - (m)$ としてステップ2を再帰的に適用する。答が $m \notin M(S)$ の場合は、 m をルートとする部分木に対して同じ M^* のままステップ2を再帰的に適用する。

本アルゴリズムと Shapiro のアルゴリズムとの違いは $M(S)$ の意味を GHC のデバッグ向きに拡張した点にある。ノード数が N 、最大分枝数が b であるような計算木に対して本アルゴリズムが高々 $b \cdot \log N$ 回の質問を發して後必ず停止してバグを含んだ節のインスタンスを返すことは Shapiro の証明と同様にして証明することができる。

3 実装

本デバッガは GHC 自身により GHC のメタ・インタプリタ的なものとして実装されている。デバッガは大きく分けて2つ、質問応答を管理するプロセスと計算木を管理するプロセスとからなっている。前者はプログラマとの質問応答を行ない、過去の質問応答の結果を記憶して同じ質問を繰り返さないようにしている。後者はデバッギングの前半(ステップ1)において計算をシミュレートしながら計算木を木状のプロセス・ネットワークとして生成する。デバッギングの後半(ステップ2)において、この木状のネットワークは木の各ノードの重み値の更新や中点の計算を並列に行なう。現在はデバッガをメタ・インタプリタ的なものとして実装しているが[3]にあるようにこのデバッガをプログラム S に関して部分計算すれば S の専用デバッガ S' を得ることができる。 S' は S に専用なだけもののインタプリティブなデバッガより効率がよいことが予想される。プログラム・リストを含む詳細は[5]にある。

4 議論

本稿では“Divide and Query”法にもとづく GHC プログラムのデバッギングについて述べた。高橋らは関数型プログラムを対象に“Divide and Query”法を基本としながらもプログラムの静的な構造についての情報を用いてデバ

ッギングの効率を改善する方法について報告している[4]。本稿で示したデバッガについても同様の改善の可能性はあると思われるが、我々はそれを前節で述べた部分計算の延長として行ないたいと考えている。

次のような症状を持つプログラムは現在のデバッガでは対処できないが、適当な拡張を行えば対処が可能となる。

- a 計算が失敗する。
- b 無限再帰を含み停止しない。
- c ガードの計算中に存在するバグ

次に示すようなバグは本デバッガのアルゴリズムでは原理的に対処できない。

- d 非決定的で再現性のないバグ

非決定的で再現性のないバグをデバッグ時に適宜につかまえる手段はない。

- e スタベーション

スタベーションは一般に不公平なスケジューリングにより生じるが、本デバッガではスケジューリング自身をデバッグの対象とはできない。

- f 計算は正しいが途中の動作が間違っている。

一般に非決定的かつ並列プログラムの意味としては、正しい入出力値関係だけでは不十分であることが知られている[6,7]。しかしながら本デバッガは基本的に正しい入出力関係をベースにしているためこのようなバグはとれない。

これらのバグは非決定的並列プログラムに固有なものであり最も対処の難しいものである。これらはデバッガの対象というよりもむしろプログラムの検証系の対象というべきものであろう。

文献

- [1] K.Ueda: Guarded Horn Clauses, Proc. of Logic Programming Conf. (Tokyo), 1985.
- [2] E.Shapiro: Algorithmic Program Debugging, MIT Press, 1983.
- [3] 竹内他: Prologプログラムの部分計算とメタ・プログラムの特殊化への応用, Proc. of Logic Programming Conf. (Tokyo), 1985.
- [4] 高橋他: 関数依存グラフの変換による関数型プログラムのデバッグ法, 情報処理学会ソフトウェア基礎論研究会 WGSP12-3, 1985.
- [5] A.Takeuchi: Algorithmic Debugging of GHC Programs, ICOT Technical Report, in preparation.
- [6] J.B.Brock, W.B.Ackerman: Scenario: A Model of Non-determinate computation, In Lecture Notes in Computer Science, vol.107, Springer-Verlag, 1981.
- [7] J.Staples, V.L.Nguyen: A Fixpoint Semantics for Nondeterministic Data Flow, JACM, vol.32, No.2, 1985.

GHCサブセット逐次型処理系の作成

2G-4

*1 江崎令子 *2 宮崎敬彦 *1 太田孝

*1 三菱電機(株)情報電子研究所

*2 (財)新世代コンピュータ技術開発機構

1. はじめに

GHCサブセット(FlatGHCに、簡単なモジュール構造をサポートしたもの、以下、単にFGHC)の処理系を、逐次型推論マシンPSIのヒープ上で実現した。使用した言語は、ESPである。

FGHCとは、GHC[上田85]に対して以下の制限を加えたものである。

- (1) passive partは逐次的に実行する。
- (2) guard 部に書ける述語は、組込み述語のみ。
- (3) 候補節のテストは上から下に逐次的に行なう。

FlatGHCのプログラムは、コンパイラにより、Warren流の中間コードに変換される。本稿では、この中間コードを実行するエミュレータの実現方式について述べる。

2. 処理の概要

2-1. 処理の流れ

あるgoalについての大きな処理は、次のように進む。

- 1) goalを ready_queueから取り出す。
- 2) goalに対応する各clauseの guard部を実行。
 - ・ successならば、trustしてbody部のgoal列を展開し、ready_queueに入れる。
 - ・ suspendならば、goalをsuspension_queueに入れる。
 - ・ そのgoalに対応するすべてのclauseを実行して、successも suspendもしないならば、そのgoalはfailなので、AND兄弟をfailさせる。
- 3) 1)へ戻る。

図2-1は、処理の概要を図式的に示したものである。

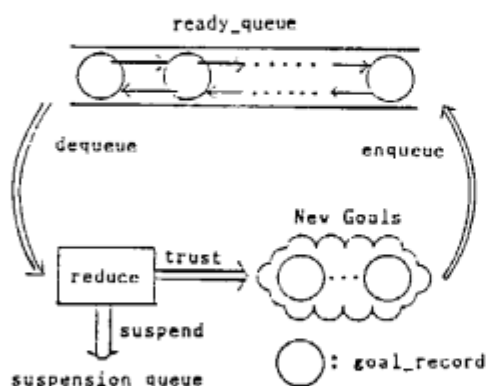


図2-1. 処理の流れ

2-2. ゴール・レコード

ユーザ・プログラム中の各goalに対応する。

ゴール・レコードは、スケジューリング用の双方向ポインタ、AND木管理用の双方向ポインタ、各モジュールに対応するclass objectと述語名アトム、引数表、どのメタ・コールに所属しているかを示す識別子、などの要素から成る。

2-3. メタコール・レコード

ユーザ・プログラム中のメタ・コールに対応する。

メタコール・レコードは、スケジューリング用の双方向ポインタ、AND木管理用の双方向ポインタ(兄弟用、子供用)、メタコールの結果を unifyする変数へのポインタ、メタコールの実行をコントロールするコードへのポインタ、識別子、などの要素から成る。

2-4. サスペンション・レコードと

サスペンション・キュー

サスペンション・レコードは、goalが未定義変数により suspendした時、そのgoalに対応して作成され、サスペンション・キューを構成する。

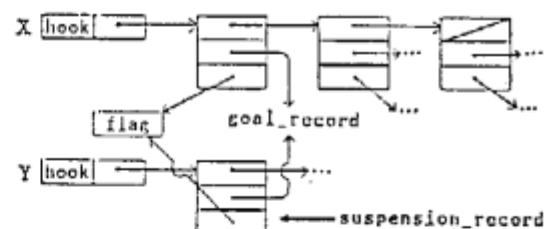


図2-2. サスペンション・キューの構成

サスペンション・レコードは、以下の要素から成る。

- (1) 同一の変数にhookしているサスペンション・レコード同志をつなぐ単方向ポインタ
- (2) suspendしているゴール・レコードへのポインタ
- (3) サスペンション・レコードの指しているゴール・レコードが、既にresume(hookしていた未定義変数が instantiateされて、suspendしていたゴールが、再度レディ・キューに入ること)されたかどうかを示すフラグ

図2-3に、resumeのアルゴリズムを示す。


```

if instantiate(Some_Cell) then
  ( Susp_rec = value(Some_Cell),
    repeat
      Next_record = (next(Susp_rec)),
      Flag = flag(Susp_rec),
      if value(Flag) == off then
        ( enqueue(goal_rec(Susp_rec)),
          value(Flag) := on ),
        die(Susp_rec),
        Susp_rec := Next_record
      until (Next_record == null) ).

```

図2-3. resumeのアルゴリズム

3. 中間コード

本処理系では、中間コードとして、Warrenの命令セット [Warren83] をFGHC用に拡張して使用した。

Warrenのものと基本的異なるのは、以下の点である。

- 1) clause中の変数は全てtemporary variableとする。
- 2) passive part中の unificationを特殊化し、効率化を計っている。

以下に、命令セットの概要を述べる。

3-1. ユニフィケーション系命令

(1) put系命令

引数レジスタから、ヒープ上に取られた引数表に値を格納する。

(2) put_to_reg系命令

引数レジスタから引数レジスタに値をロードする。TROのために、active partの適当なゴールに対して使用される。

(3) get0系命令

clause headの引数の unificationを行なう。get0_variable以外のget0系命令は、指定された引数レジスタの内容が未定義の場合は、suspendする。

(4) get系命令

active partの unificationを行なう。基本的にはsuspendしない。

(5) unify0系命令

構造体の要素に対する unification。get0_structureの後に現われ、read modeで実行される。unify0_variable以外のunify0系命令は、指定された構造体の要素が未定義の場合は suspendする。

(6) onify系命令

構造体の要素に対する unification。put_structureの後に現われ、write modeで実行されるか、または、get_structureの後に現われ、その結果に応じて、read modeまたは、write modeのいずれかで実行される。基本的には suspendしない。

3-2. 手続きの命令

:alloc_args ... ヒープ上に引数表をとる。

:terminate ... ゴール・レコード領域を解放して、AND木の操作を行なう。

:create ... ゴール・レコードの作成と、レディ・キューへの enqueueを行なう。

:first_create, :only_one_create, :last_create など、レコードの再利用を行なうよう最適化されたものもある。

以上の命令の他に、組込み述語などに対応する命令コードが考えられる。

3-3. コンパイルド・コード

FGHCプログラムは、コンパイラによって、ESPのプログラムに変換される。図3-1. は、コンパイルド・コードの例である。

```

*** FlatGHC source program ***
sift([PIL],S) :- true | S=[PIS1], shift(K,S1),
                                     filter(P,L,K).

*** compiled code ***
:sift(C, Vs, Bnd) :-
  :get0_list(C, 0, Vs),
  Bnd>0,
  :unify0_variable(C, 0, Vs),
  :unify0_variable(C, 4, Vs),
  !,
  :get_list(C, 1, Vs),
  :unify_value(C, 0, Vs),
  :unify_variable(C, 3, Vs),
  :alloc_args(C, 2, As),
  :put_variable(C, 2, Vs, 0, As),
  :put_value(C, 3, Vs, 1, As),
  :only_one_create(C, sift, As, Bnd-1),
  % make goal_record and enqueue
  :put_value_to_reg(C, 4, 1, Vs),
  :filter(C, Vs, Bnd-1);
  % direct call

```

図3-1. 中間コード・プログラム

図3-1. において、最後のgoal, filterは、レディ・キューに入らず、直接呼び出されることで、高速化が計られている。

4. 評価と課題

以上に示した仕様に基づき、逐次型推論マシンPSI上にGHCサブセットの処理系を作成した。

実行速度は、30要素の naive reverseで、0.87krps.であった。

現在の処理系では、逐次的に処理を行なう事を前提としているが、今後、数台のマシン上で実際に分散処理を行なう処理系へと拡張していく予定である。

「参考文献」

- [上田85] Kazunori, Ueda:
"Guarded Horn Clauses"
ICOT Technical Report 1985.
- [Warren83] Warren, D.H.D.:
"AN ABSTRACT PROLOG INSTRUCTION SET"
SRI International,
Technical Note 309, 1983.
- [宮崎85] 宮崎敏彦;
"KL1-B (第0版)の提案"
ICOT KL1-B検討会資料, 1985.

AN-12

ANALOGICAL REASONING AS PROBABILISTIC REASONING

TAKASHI YOKOMORI
IIAS-SIS, FUJITSU LTD.

1. Introduction

We note there are various contexts in which people say "one is analogous to the other". The expression "as like as two peas", for example, comes from the analogy of appearance, while "similar story" and "inherited temper" fall under another category of expressions corresponding to another type of analogy. One may classify analogies as follows:

- syntactic analogy (a)
 - {form, appearance, symbolism}
- semantic or functional analogy (b)
 - {temper, performance}
- static analogy (c)
 - {situation, sense (color, taste)}

In the study on analogical reasoning, [H85a] discusses and proposes one method for realizing the analogical inference based on its "theory of analogy" which only concerns the analogy of type (a). The key idea there is to formalize the concept of analogy using "inductive generalization" by [P70]. Provided that any object in analogical reasoning can be expressed as a finite set of ground atoms in Horn logic, the analogy between the two objects S_1 and S_2 is defined as follows:

S_1 and S_2 satisfy the partial identity condition if there exists a one-to-one mapping f such that $f(Sg_1) = Sg_2$, where S is the least generalization of S_1 and S_2 for substitutions g_1 and g_2 , i.e., S is the least general such that $Sg_1 \subset S_1$ and $Sg_2 \subset S_2$.

In the definition above, f claims a strict one-to-one correspondence between the two objects, which causes some difficulties in dealing with common cases of analogy. It cannot handle analogies between the two objects having distinct domains, further, even if in the same domain, a semantic analogy of the kind (b) or (c) above neither.

In the later paper [H85b] Haraguchi extends his formulation of analogy (partial identity condition) to overcome some of its drawbacks above and to deal with a kind of semantic analogy to some extent. There still remain, however, quite a few problems to be solved in the extended version. For example, the identification of a concept in different formulations cannot be treated, and more typically, no consideration has been given to the problem of handling the "everyday" analogy used in daily life such as "like something".

The goal of this paper is to show that taking Dempster's theory of probability as a formal framework analogical reasoning is regarded as probabilistic reasoning, and to argue on a possible way of handling semantic analogy in that framework.

2. Semantic Analogy and Analogical Reasoning

As mentioned in the previous section, by semantic analogy we mean the types (b) and (c) of analogy. However, our intension here is mainly the type (c), that is, static analogy. In what follows, we simply refer to semantic analogy as analogy, unless any confusion arises. We also assume that the objects in discourse of analogy are described in Horn logic formulas.

Consider the followings that give typical examples of analogy in daily life.

Let $\text{parent}(X,Y)$ be a predicate meaning X is a parent of Y , and $\text{couple}(X,Y)$ meaning X and Y are a couple. In general, we have

$\text{grand-parent}(X,Y) \leftarrow \text{parent}(X,Z), \text{parent}(Z,Y)$

--- (1)
in-law(R, X, Y) $\leftarrow R(X,Z), \text{couple}(Z,Y)$ --- (2)
where R is a relation name. (For example, in-law($\text{parent}, \text{tom}, \text{mary}$) denotes "tom is a parent-in-law of mary".) Further, it holds

$\text{transitive}(R, X, Y) \leftarrow R(X,Z), R(Z,Y)$.

Then, firstly, (1) is re-expressed as follows:

$\text{grand-parent}(X,Y) \leftarrow \text{transitive}(\text{parent}, X, Y)$

--- (1').
This implies that we must take the analogy (in fact, identity) of this kind into account. It is too numerous to count up instances of this type. (The identities of "ancestor" and "transitive closure of parent", of "male" and "antonym of female", etc.)

Secondly, since

$\text{in-law}(\text{grand-parent}, X, Y) \leftarrow \text{grand-parent}(X,Z), \text{couple}(Z,Y)$

we eventually have

$\text{in-law}(\text{grand-parent}, X, Y) \leftarrow \text{transitive}(\text{parent}, X, Y), \text{couple}(Z,Y)$

--- (3).
Comparing (1') and (3) instead of (1) and (3), it is suggested that there is a similarity in syntax between the two expressions which emerges through the "analogy relation" couple, while the comparison of (1) with (3) can never give us a chance to detect even a syntactic analogy between the two.

Note that the second example provides an interesting view of analogy in which one may say semantic analogy is transformed into syntactic one.

Now, considering analogical reasoning we encounter in everyday life, it can be, we believe, characterized by:

- (i) detecting analogy mostly in terms of subjective criteria,
- (ii) non-monotonic reasoning features, leading to
- (iii) approximate reasoning based on incomplete knowledge base.

Thus, it seems quite reasonable to look at analogical reasoning as non-monotonic reasoning in which approximate reasoning using subjective probability plays the central role.

3. Dempster's Probability Theory

In [Sh76] Shafer has discussed the way of treating uncertain knowledge such as "A bird can fly" or "People can see", and given a suggestive view that certainties of those are better thought of not as specific values but as ranges. He has refined Dempster's work [D68] on the theory of probability which extends the classical Bayesian theory so as to cope with a task of treating "ignorance" and "disbelief". (For the details on Dempster and Shafer's theory, see [I83] as well as the two above.)

As mentioned above, if we take the position to look at analogical reasoning as non-monotonic reasoning, then the analogical reasoning system we intend must have a

capability of inference based on "uncertain knowledge" in addition to that of belief revision. As shown below, Dempster's rule has many attractive properties in the sense above.

Here, we follow the notation in [G84a] and [G84b]. The truth values T and F are thought of as probabilities 1 and 0, respectively. A probability interval $[a,b]$ (where $0 \leq a \leq b \leq 1$) is denoted by the pair (a,b) , where b stands for $1-b$ and $a \leq b \leq 1$. Intuitively, " a " means the degree to which a given statement is confirmed (to be true) by the available evidence, while " b " the degree to which it is disconfirmed (or doubted to be false). Examples are a true statement with the value $(1,0)$, a false statement with the value $(0,1)$.

Incomplete information is represented as follows. "A bird can fly" might have truth value $(.9,.02)$, or "Humans are mammals" has the value $(1,0)$, while "Desk can see" which is an example of extremely incomplete information might have the value $(0,0)$. Thus, generally, the completeness of information represented by the value (a,d) corresponds to the sum $a+d$ called the mass of the truth value.

The rules are associated with their truth values as follows:

$Q \leftarrow (a,d) \quad P_1, \dots, P_n$
For example, $\text{can-fly}(X) \leftarrow (.9,.02) \text{ bird}(X)$. To perform one inference derivation using a rule, assume that $\text{bird}(x)$ has the truth value (a,d) . Then, $\text{can-fly}(x)$ is drawn with the truth value $(.9a,.02a)$. Note that " d " does not contribute to this value, because the rule cannot be used to draw any conclusion from the fact that " x is not a bird". (In case when a rule has multiple antecedents, a product of their each value should be used as the truth value of the total certainty for them.)

Suppose that
 $\text{bird}(\text{peeko}) (1,0)$
 $\text{can-fly}(X) \leftarrow (.9,.02) \text{ bird}(X)$.
Then we have a conclusion
 $\text{can-fly}(\text{peeko}) (.9,.02) \quad \text{--- (4)}$.
On the other hand, if $\text{isa}(\text{peeko}, \text{penguin}) (1,0)$ is obtained later, then using a rule
 $\text{can-fly}(X) \leftarrow (0,1) \text{ isa}(X, \text{penguin})$,
we eventually draw
 $\text{can-fly}(\text{peeko}) (0,1) \quad \text{--- (5)}$,
leading to a typical situation (a kind of contradiction) in non-monotonic reasoning.

Dempster has discussed this situation in detail and given a solution to this difficulty. If we denote by $(a,b) + (c,d)$ the inference obtained by combining rather contradictory two inferences (a,b) and (c,d) , then Dempster's rule gives us

$$(a,b) + (c,d) = \left(1 - \frac{ad}{ad+bc}, 1 - \frac{bc}{ad+bc} \right).$$

The operation $+$ has many preferable properties for our purpose. Those are

- [1] $(a,b) + (c,d) = (c,d) + (a,b)$
- [2] $(a,b) + ((c,d) + (e,f)) = ((a,b) + (c,d)) + (e,f)$
- [3] $(a,b) + (0,0) = (a,b)$
- [4] $(a,0) + (c,0) = (a+c, 0)$
- [5] $(1,0) + (c,d) = (1,0)$ if $(c,d) \neq (0,1)$
- [6] $(0,1) + (c,d) = (0,1)$ if $(c,d) \neq (1,0)$
- [7] $(1,0) + (0,1)$ is undefined
- [8] if we denote by $-$ the inverse of $+$, for $(c,d) \neq (0,1)$ or $(1,0)$, we have

$$(a,b) - (c,d) = \left(\frac{ad-bc}{ad+bc}, \frac{bc-ad}{ad+bc} \right).$$

This allows us to retract specific inferences without affecting others and provides many of the same features as a truth maintenance system.

4. Detection of Analogy

We shall discuss how analogical reasoning is incorporated in non-monotonic reasoning using Dempster's probability theory. Although Dempster's theory is suitable for dealing with subjective probability, we feel the detection of analogy should be performed by the mechanism of some kind that is based on objective measures. Such a measure is used for computing the degree of analogy.

Let $A(x)$ and $B(x)$ be predicates having their Herbrand domains $\text{Dom}(A)$ and $\text{Dom}(B)$, respectively. Assume that there is a one-to-one mapping F from $\text{Dom}(A)$ to $\text{Dom}(B)$ and all elements of $\text{Dom}(A)$ (or $\text{Dom}(B)$) are enumerable in some fixed order. Given an $A(x)$, define

$L(A) = \{w \in \text{Dom}(A) \mid A(w) \text{ holds true}\}$.
Further, let d_F be defined as follows:

$$d_F(A,B) = \text{dist}(F(L(A)), L(B)),$$

where $\text{dist}(L_1, L_2) = \text{norm}(L_1 \oplus L_2)$,

$$\text{norm}(L) = \sum_{i=1}^n m_i, \text{ and}$$

$$f_i = \begin{cases} 1 & (i\text{-th element of the domain is in } L) \\ 0 & (\text{otherwise}), \end{cases}$$

$$m_i = 2^{-i} (i > 1), \quad \sum m_i = 1,$$

$\oplus$ denotes the symmetric difference.

It is easily seen that a mapping dist satisfies the metric axioms, i.e. it provides a metric for measuring the analogy of the two concepts. Note that $d_F(A,B) = 0$ iff $A=B$ under F , $d_F(A,B) = 1$ iff A and B are complementary, and $0 < d_F(A,B) < 1$ holds.

Using a metric d_F , analogical reasoning is performed in the framework of probabilistic reasoning as shown below.

Suppose the domains D_1 and D_2 between which there exists a one-to-one mapping F are given. Further, let S_i be a set of Horn clauses with $D_i (i=1,2)$ in which each fact or rule of S_i is already assigned its truth value (a,d) . (If a clause is confirmed true, then $(a,d) = (1,0)$.) Then, consider the union $S = S_1 \cup S_2$ as the whole system for performing reasoning in $D_1 \cup D_2$.

There exist two types of rules to be possibly added to S .

- (1) If A is a ground atom in S_1 , then
 $F(A) \leftarrow (a,d) \quad A$ is possibly a rule,
where the truth value (a,d) is determined from F by subjective criteria.
- (2) If $d_F(A,B) = t$, then $B \leftarrow (1-t, t) \quad A$ is possibly a rule.

Thus, the reasoning system S with additional rules works as a probabilistic reasoning system using Dempster's rules.

This work is part of the major R&D of the FGMP, conducted under program set up by MITI.

References

- [D68] Dempster, A.P., A Generalization of Bayesian Inference, *J. Roy. Stat. Soc. B*, 30 205-247, (1968).
- [G84a] Ginsberg, M.L., Non-monotonic Reasoning Using Dempster's Rule, *Proc. AAAI-84*, 126-129, (1984).
- [G84b] Ginsberg, M.L., Implementing Probabilistic Reasoning, *Res. Rep. No. 84-31*, Stanford Univ., (1984).
- [H85a] Haraguchi, M., Analogical Reasoning Based on the Theory of Analogy, *RIFIS-RR*, Kyushu Univ. No. 105 (1985).
- [H85b] Haraguchi, M., Towards Semantics of Analogical Reasoning, paper at *Inf. Sys. Symp.*, IIAS-SIS, Fujitsu Ltd., Nov. 1985.
- [I83] Ishizuka, M., Dempster & Shafer's Theory, *IECE of Japan*, Sep. 900-903, (1983).
- [P70] Plotkin, G.D., A Note on Inductive Generalizations, *Machine Intelligence*, 5, 153-216, (1970).
- [Sh76] Shafer, G., "A Mathematical Theory of Evidence", Princeton Univ. Press, (1976).

4M-6

ON SEMANTICS OF LOGIC PROGRAMS WITH UNCERTAINTIES

YASUBUMI SAKAKIBARA
IIAS-SIS, FUJITSU LTD.

1. Introduction

A mechanism for dealing with uncertainties plays a central role in knowledge systems like expert systems. Especially, incorporating uncertainties into logic programs (e.g. Fuzzy-Prolog [2],[3]) is a current study. On the other hand, there are few semantic considerations for it. In this paper, we will first attempt to give a semantics for logic programs with uncertainties. This is a further work of [4] which is the only paper dealing with a semantics so far. Our basic definitions of logic programs with uncertainties follow it. Secondly, we will define a proof procedure for logic programs with uncertainties which corresponds to an interpreter for it and present a sufficient condition for certainty functions of clauses to achieve the completeness of the proof procedure.

2. Logic programs with uncertainties

A definite clause is a clause of the form $A \leftarrow B$, where A is an atom and B is a conjunction of zero or more atoms. A certainty factor c is a real number between 0 and 1. A certainty function f is a function from multisets of certainty factors to certainty factors. A logic program with uncertainties P is a finite set of pairs $(A \leftarrow B, f)$, where $A \leftarrow B$ is a definite clause and f is a certainty function.

The certainty function computes the certainty of the conclusion of a clause from the multiset of certainties of solutions to goals in the condition of the clause. Two requirements of f in [4] are that for every multisets S , $f(SU\{1\}) = f(S)$, and that f be monotonic increasing, which means that $S \leq S'$ implies $f(S) \leq f(S')$, where $\leq$ is the partial order over multisets, defined as follows. Let S and $X = \{x_1, \dots, x_n\}$ ($n \geq 0$) be two multisets. Then $X \leq S$ iff there is a multiset $Y = \{y_1, \dots, y_n\}$ such that $S \leq Y$ and $x_i \leq y_i$ for $0 \leq i \leq n$.

3. Semantics for logic programs with uncertainties

An interpretation I of a logic program with uncertainties P is a set of pairs (A, c) , where A is a ground atom and c is a certainty factor. An I contains at most one pair (A, c) for any atom A . A ground atom A is true in I with certainty c iff there is a pair (A, c') in I such that $c \leq c'$. An atom A is true in I with certainty c iff for every ground instance A' of A , A' is true in I with certainty c . Let $A \leftarrow B_1, \dots, B_n$ ($n \geq 0$) be a ground definite clause, f be a certainty function and S be the multiset of certainties $\{c_1, \dots, c_n\}$ such that (B_i, c_i) is in I for $1 \leq i \leq n$. Then $A \leftarrow B_1, \dots, B_n$ is true in I w.r.t. f iff A is true in I with certainty $f(S)$. A definite clause $A \leftarrow B$ is true in I w.r.t. f iff any ground instance $A' \leftarrow B'$ of it is true in I w.r.t. f . A logic program with uncertainties P is true in M iff for any pair $(A \leftarrow B, f)$ in P , $A \leftarrow B$ is true in M w.r.t. f . Such an interpretation M is called a model for P .

A partial order $\leq$, intersection $\cap$ and union $\cup$ on interpretations are defined in the natural way as follows. Let I and I' be interpretations. Then $I \leq I'$ iff for any pair (A, c) in I there is a pair (A, c') such that $c \leq c'$. (A, c_i) in $I \cap I'$ iff for two pairs (A, c) in I and (A, c') in I' $c_i = \min(c, c')$. (A, c_u) in $I \cup I'$ iff for two pairs (A, c) in I and (A, c') in I' $c_u = \max(c, c')$.

Lemma 3.1 (model intersection property) Let M be any nonempty set of models for P . Then $\cap M$ is also a model for P .

By the above lemma, the least model $\mathcal{M}(P)$ for P which is the intersection of all models for P exists.

We define a transformation Tc_p associated with P in the same way as standard logic programs. Let I be an interpretation. (A, c) in $Tc_p(I)$ iff $c = \sup\{f(S) \mid A \leftarrow B_1, \dots, B_n \text{ is a ground instance and } f \text{ is the certainty function of a clause in } P \text{ such that } (B_i, c_i) \text{ in } I \text{ for } 1 \leq i \leq n \text{ and } S = \{c_1, \dots, c_n\}\}$.

Lemma 3.2 Let I be an interpretation. Then I is a model for P iff $Tc_p(I) \subseteq I$.

Lemma 3.3 Tc_p is a monotonic mapping.

By the above lemmas and the Knaster-Tarski fixpoint theorem [1], we have the following theorem.

Theorem 3.4 The least model $\mathcal{M}(P)$ for P is equal to the least fixpoint of Tc_p ($\text{lfp}(Tc_p)$ for short).

We consider $\mathcal{M}(P)$ and $\text{lfp}(Tc_p)$ as the semantics for a logic program with uncertainties P .

4. Proof procedure for logic programs with uncertainties.

A (P, F) -derivation is a sequence of quadruples $\langle G_i, C_i, \theta_i, F_i \rangle$, $i=0, 1, \dots$ such that (1) G_i is of the form $\leftarrow (B_1, \dots, B_m)$ where $m \geq 0$ and B_j is an atom ($1 \leq j \leq m$) and F_i is of the form $\{X_1, \dots, X_m\}$, (2) C_i is a list of m clauses $(A^{(1)} \leftarrow D_1^{(1)}, \dots, D_{n1}^{(1)}, f^{(1)}), \dots, (A^{(m)} \leftarrow D_1^{(m)}, \dots, D_{nm}^{(m)}, f^{(m)})$, (3) θ_i is a most general unifier of $(B_1, \dots, B_m)$ and $(A^{(1)}, \dots, A^{(m)})$, and (4) G_{i+1} is $\leftarrow (D_1^{(1)}, \dots, D_{n1}^{(1)}, \dots, D_1^{(m)}, \dots, D_{nm}^{(m)})\theta_i$ and F_{i+1} is $\{Y_1^{(1)}, \dots, Y_n^{(1)}, \dots, Y_1^{(m)}, \dots, Y_n^{(m)}\}$ where $X_1 = f^{(1)}(\{Y_1^{(1)}, \dots, Y_n^{(1)}\}), \dots, X_m = f^{(m)}(\{Y_1^{(m)}, \dots, Y_n^{(m)}\})$. A (P, F) -derivation is successful if it is finite and its last goal (some G_i) is empty. The certainty of a successful (P, F) -derivation is F_0 in it. Then we have the following theorem which means the soundness of a successful (P, F) -derivation.

Theorem 4.1 For a ground atom A , if A has a successful (P, F) -derivation with the certainty c , then A is true in $\mathcal{M}(P)$ with c .

However we do not generally have the completeness of a successful (P, F) -derivation. With the restriction on a certainty function f that $f(\{c_1, \dots, c_m\}) \leq \min(c_1, \dots, c_m)$, the completeness can now be achieved.

Lemma 4.2 Assume that for each certainty function f , $f(\{c_1, \dots, c_m\}) \leq \min(c_1, \dots, c_m)$. Then for (A, c) in $\mathcal{M}(P)$, there exists some N such that (A, c) in $Tc_p^N(\emptyset)$.

Theorem 4.3 Assume that for each certainty function f , $f(\{c_1, \dots, c_m\}) \leq \min(c_1, \dots, c_m)$. For a ground atom A , if A is true in $\mathcal{M}(P)$ with c , then A has a successful (P, F) -derivation with certainty c' such that $c \leq c'$.

We believe this condition (i.e. $f(\{c_1, \dots, c_m\}) \leq \min(c_1, \dots, c_m)$) will play an essential role in the validity of quantitative reasoning.

This work is part of the major R&D of the FGCP, conducted under program set up by MITI.

References

- [1] K.R.Apt and M.H.van Emden, "Contributions to the theory of logic programming", JACM, 29, 3(1982).
- [2] M.Ishizuka and N.Kanai, "PROLOG-ELF INCORPORATING FUZZY LOGIC", Proc. 9th IJCAI, 1985.
- [3] M.Mukaidono, "Fuzzy Prolog no kokoromi" (in Japanese), Proc. 27th Conf. JIPS, 1983.
- [4] E.Y.Shapiro, "Logic Programs With Uncertainties", Proc. 8th IJCAI, 1983.

ON SEMANTIC ANALOGY IN OBJECT-ORIENTED REPRESENTATION

4M-13

YUJI TAKADA
IIAS-SIS, FUJITSU LTD.

1. Introduction

Analogy is the very powerful and basic human information processing in reasoning and learning. Therefore, various analogies have ever been discussed, but most of them did not have flexibilities or logical validities. In this paper, we discuss a study of analogy in an object-oriented language based on a logic programming language like PROLOG, and show some flexibility and validity of this analogy.

In such languages, each object represents an axiom set. For knowledge representation and processing, one of characters of a logic programming language provides some validities, while that of object-oriented language does some flexibilities. Here, we discuss analogies from two different viewpoints. One concerns logical semantics, that is, the one-to-one correspondence between models. The other concerns more structural relationships between objects.

2. Object-oriented Representation

In this paper, the objects in discourse of analogy are represented in an object-oriented language like KORE/KR [K85] or ESP [C84] based on a logic programming language.

Class is defined as follows:

```
class <class name> has
  isa <is-a definition> ;]
  part <part definition> ;]
  slot <slot definition> ;]
  [ <method definition> ]
end.
```

where underline indicates a reserved word and [...] indicates "... is optional. <is-a definition> comprises a sequence of class names, which form super classes for multiple inheritance mechanism, while <slot definition> a sequence of slot names. <part definition> comprises a sequence of relations of the form $P < C$, where P is a variable called part and C is a class name, meaning that the value of P is an instance belonging to a class C . Instantiating P , together with an instance belonging to the class <class name>, provides a PART_OF relationship. <method definition> is a collection of methods. Method M_i ($i=1, \dots, n$) is defined as follows:

$$MP_i \Rightarrow MC_{i1}, \dots, MC_{im(i)} \vdash (m(i) > 0)$$

where MP_i is a "term" in first-order language. MC_{ij} is called "method call" and denoted by $l \leftarrow m$, where m a message, l a variable supposed to be instantiated with an instance. We call the right side of $\Rightarrow$ body.

Suppose an object represents an axiom set. Then a method call $l \leftarrow m$ implies that m is

interpreted in the axiom set I . The interpretation of method M_i depends on all instances appearing in method calls of MC_{ij} s, therefore so does the interpretation of an object with M_i . Hence, it is possible to take that PART_OF explicitly represents a dependency of models between objects.

3. Definition of analogy as renaming

Analogy in this paper is based on renaming, which guarantees one-to-one correspondence between models. We take mutual recursively definitions in the following.

Definition 1 For given two terms t_1 and t_2 , renaming F is a non-empty set of pairs recursively defined as follows:

1) If both t_1 and t_2 are constants exclusively or variables, then $F = \{ \langle t_1, t_2 \rangle \}$.

2) If $t_1 = f(u_1, \dots, u_n)$ and $t_2 = g(v_1, \dots, v_n)$ then for each $i=1, \dots, n$ there exists renaming F_i for u_i and v_i , and $F = \{ \langle f, g \rangle \} \cup \bigcup_{i=1}^n F_i$ is one-to-one.

Definition 2 Given two classes C_1 and C_2 , let S_i be the set of all slot names in C_i ($i=1, 2$). Then, we say C_1 and C_2 are analogous on slots iff $S = S_1 \cap S_2$ is non-empty. Renaming $F_s = \{ \langle s, s \rangle \mid s \in S \}$ is called slot-analogy.

Definition 3 Let method call MC_i be $l_i \leftarrow m_i$, and suppose an instance denoted by I_i belongs to C_i ($i=1, 2$). Then, we say MC_1 and MC_2 are analogous on method call iff:

1) Classes C_1 and C_2 are analogous in the manner below. Let F_c be a renaming for C_1 and C_2 . (See definition 2)

2) There exists a renaming F_m for messages m_1 and m_2 .

3) $F_{mc} = F_c \cup F_m$ is a renaming.

We call F_{mc} analogy on method call.

We say two methods M_1 and M_2 are comparable iff both MP_1 and MP_2 have the same arity and $m(1)=m(2)$.

Definition 4 For two comparable methods M_1 and M_2 , let $MCS_i = \{ MC_{ij} \mid j=1, \dots, n \}$, where MC_{ij} is a method call in the body of M_i ($i=1, 2$). Then, M_1 and M_2 are analogous on method iff:

1) There exists a renaming F_{mp} for MP_1 and MP_2 .

2) If $n > 0$, then there exists a one-to-one correspondence R on $MCS_1 \times MCS_2$, where there exists an analogy on method call F_j for each pair in R , and $F_b = \bigcup_{j=1}^n F_j$ must be a renaming. In

the case when $n=0$, i.e. both MCS_1 and MCS_2 are empty, F_b is empty.

3) $F_m = F_{mp} \cup F_b$ is a renaming.

We call F_m a **method-analogy**.

Definition 5 Given two classes C_1 and C_2 , let F_s a slot analogy for their slots, F_{s1} a union of some of the method-analogies for their methods. Then, we say C_1 and C_2 are analogous if there exists a renaming $F = F_s \cup F_{s1}$. We call F **class-analogy**, or merely **analogy**.

Semantically, for analogous slots, only correspondence between their models is the predicate since detections of slot-analogy concern no slot-values, while for analogous methods, there exists a strict one-to-one correspondence between their models. Thus, the analogies here have various semantics, which may correspond to those of conditions for analogies like PIC and SPIC in [H85].

4. The role of PART_OF on analogy

On the basis of those analogies mentioned above, there exists another type of analogy, **analogy on parts**, where PART_OF relationships play the central role. This type of analogy concerns more structures of knowledge, rather than logical validities.

Definition 6 Let P_i be the set of all parts in C_i ($i=1,2$) and let $P \subset P_1 \times P_2$. Then, we say classes C_1 and C_2 are **analogous on parts** iff:

1) For any $q = \langle p_1, p_2 \rangle \in P$, there exists an analogy F_q for classes C_{p1} and C_{p2} which p_1 and p_2 belong to, respectively.

2) P is one-to-one.

3) $F_p = \bigcup_{q \in P} F_q$ is a renaming.

F_p is called **part-analogy**.

Now, taking the part-analogy into account, we extend the definition of analogy for classes so as $F = F_s \cup F_m \cup F_p$ to be a renaming. Suppose PART_OF relationship explicitly represents a dependency between objects. Then the part-analogy concerns the axiom sets an object depends on, rather than its own one.

PART_OF relationships and part-analogies provide some flexibilities. In learning, objects may often have no slot or method but parts. In such cases, PART_OF and part-analogies provide heuristic guides for acquisitions of not only new slots or methods but new classes. For example, consider Example-1 in which water-pipe and electric-chord, water and electricity are analogous. Then pipe-law and ohms-law, which would be acquired from the pressure method (i.e. voltage method) by transforming pipe-law's one, are part-analogy. Similarly, part-analogies provide the same guides in analogical reasoning.

Moreover, in acquiring methods, each of which has only analogous parts in its method

calls, there exists a semantic validity in the sense of method-analogy, that is, the acquired method in one class should be precisely method-analogous to the method in the other part-analogous one. Therefore, PART_OF relationship and part-analogy provide guides for acquisitions with that semantic validity.

5. Discussion

Consider an object in this language as the unit of knowledge representation, i.e. the schema like frame. Then, we may take that various semantic levels of analogy like slot- or method-analogy would be due to various qualities of knowledge. Actually, since knowledge has declarative or procedural properties, this fact suggests that it is possible to take slot-analogy as declarative analogy and method-analogy as procedural analogy.

Moreover, for reasoning and learning by analogy, part-analogy may provide some logical validities as well as flexibilities, because PART_OF relationship concerns not only the knowledge structures but dependencies of semantics between objects.

Hence, we believe that the approach discussed here is more fruitful, and developing a theory and a system based on it for analogical reasoning in this knowledge representation is our current work.

This work is part of the major R&D of the FGCP, conducted under program set up by MITI.

References

- [C84] Chikayama, T., ESP Reference Manual, ICOT Technical Report, TR-044, 1984.
- [H85] Haraguchi, M., Towards semantics of Analogical Reasoning, paper at Inf.Sys. Symp., IJAS-SIS, Fujitsu Ltd., Nov. 1985.
- [K85] Katayama, Y. et al., Mondai Kaiketsu Sien Kankyou KORE (Sono 3), 5L-10, in this proceedings, 1986.

Example-1

Note: *self is a pseudo-variable, getslot() is built-in method call which gets a slot value.

```
class pipe-law has
  part p<water-pipe, w<water;
  pressure(R*I) =>
    p<-resistance(R), w<-flow(I);
end.
class water-pipe has
  slot length, area;
  resistance(rho*(L/S)) =>
    *self<-getslot(length,L),
    *self<-getslot(area,S);
end.
class water has
  slot volume, time;
  flow(V/T) =>
    *self<-getslot(volume,V),
    *self<-getslot(time,T);
end.
class ohms-law has
  part c<electric-chord, c<electricity;
end.
class electric-chord has
  slot length, area;
  resistance(phi*(L/S)) =>
    *self<-getslot(length,L),
    *self<-getslot(area,S);
end.
class electricity has
  slot quantity, time;
  current(Q/T) =>
    *self<-getslot(quantity,Q),
    *self<-getslot(time,T);
end.
```


新谷虎松・片山佳則・平石邦彦・戸田光彦
(富士通国際情報社会科学研究所)

1. はじめに

KORE(Knowledge Oriented Reasoning Environment)はProlog上に構築された汎用の問題解決環境であり、固有なプログラミングパラダイムを提供するためのいくつかのサブシステムにより構成される。これらサブシステムは、共通な内部表現を有することによりKOREとして統合化される。有益な問題解決機能を提供するためには、システム内に人間の知識・問題解決能力を効果的にモデル化(知識ベース化)する必要がある。我々は、それらをモデル化するための表現技法として対象指向的表現法とルール指向的表現法を用いる。このモデル化に必要な統一的な内部表現技法は、Prologにおける関係テーブル表現技法〔1〕を導入する。このような関係テーブル技法は、システムに自然な知識表現能力と強力な推論機能を提供する一方、知識ベースとデータベースとの自然なインターフェイスの構築を可能にする。

2. KOREの基本構成

KOREは次の基本サブシステムにより構成される。(1)KORE/DB(Data Base subsystem)〔9〕、(2)KORE/IE(Inference Engine subsystem)〔3〕、(3)KORE/KR(Knowledge Representation subsystem)〔8〕、(4)KORE/EDEN(Extended Dependency network subsystem)〔7〕これらサブシステムは、単独でも独立的に機能し、各サブシステムは内部表現形式において統一化することができ、全体としてKOREに様々な問題解決支援機能を提供する。これらサブシステムの基本機能の概要は次のように記述することができる。(1)のKORE/DBはKOREにおける関係テーブルを操作管理する機能を担う一方、単独では関係データベースを操作するためのパーソナルな関係データベースシステムとして機能する〔2〕。(2)のKORE/IEはKOREにおけるルール指向的問題解決・知識表現機能を担う一方、単独ではOPSS〔3〕等に代表される強力な前向き推論型ブ

ログクッションシステムとして機能する。(3)のKORE/KRはKOREにおける対象指向的知識表現・問題解決機能〔4〕を担う一方、単独では柔軟な知識表現言語として機能する。(4)のKORE/EDENはKOREにおいて知識を管理・利用するために必要となる知識間の関係情報〔5〕を保存し、知識ベース管理機能の一機能を担う一方、単独ではネットワークを操作するための強力なシステムとして機能する〔6〕。これら基本サブシステム間の関係は図1のように図示できる。

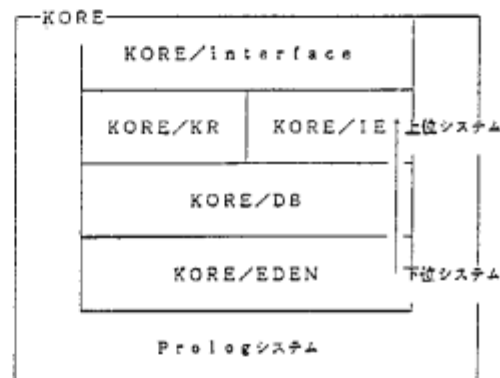


図1. KOREの構成

KOREは図1で示すように各サブシステムはその記述言語であるPrologとのインターフェイスを持つ。これらサブシステムにはレベルがあり、図1で示すような上位—下位概念が存在する。下位システムに位置づけられたサブシステムは、その上位のサブシステムの基本内部表現の管理のための基本機能を担う。例えば、KORE/IEの下位システムであるKORE/DBはKORE/IEにおけるグローバルデータベース管理機能を提供し、KORE/EDENはKORE/DBの関係情報管理機能を提供する。

KOREにおいて、上位システムは下位システムの機能を利用できるが、下位システムは上位システムを利用するための機能は与えられない。これは、各サブシステムの働きを明確にすることにより、各サブシステムを用いたアプリケーションの設計・構築法を明らか

にし、KOREの基本的なシステム設計とデバッキングを容易にすることに貢献する。

3. KOREの内部表現技法の概略

KOREは主に4つのサブシステムにより構成され、これらサブシステムは内部表現において共通の関係テーブル表現を持つことにより統合される(図2参照)。

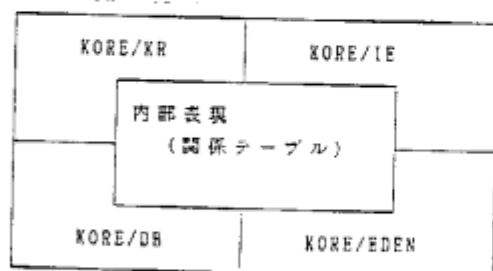


図2. KOREサブシステムの統合

関係テーブル表現は関係表現とその表現力において等価であり(1)、Prologにおいては垂直に表現ができる。例えば、以下のように関係テーブルはPrologのアサーションにより表現できる。

測定	名前	体重	身長
	太郎	80	172
	次郎	75	180
	花子	55	157

測定(太郎, 80, 172)

測定(次郎, 75, 180)

測定(花子, 55, 157)

このような関係表現における引数の位置はテーブルの列に相当し、各々の関係はテーブルの行に相当する。KOREでは関係テーブルをこのような関係表現により保存し、利用する。

4. サブシステム統合化の例

図3はKORE/KRとKORE/DBが内部表現において統合され利用されることを示している。図3においてStep1~Step4まではKORE/KRにおけるメッセージ・パッシングフローを示しておりStep1はクラスcarのロード、Step2はクラスcarからのインスタンスcivicの生成、Step3はインスタンスcivicへの属性値付加、Step4はインスタンスcivicの内部状態の表示ために用いられている。Step5はKORE/DBによる関係テーブルの検索が行われ

```

1 ?- file <- load,car. %Step 1
>> car loaded.

yes
1 ?- car <- new,civic. %Step 2

yes
1 ?- putvalue(civic,weight,90).
      :
      :
yes
1 ?- civic <- put,(power,83). %Step 3

yes
1 ?- civic <- pp. %Step 4

XXXXXXXXXXXXcivic object XXXXXXXXXXXX
classname car
>value set(
value      power 83
value      weight 90
value      displacement [1500,1300]
value      price 100
value      color [blue,red,white]
end of object
yes
1 ?- select * from civic. %Step 5
[value,power,83]
[value,weight,90]
[value,displacement,[1500,1300]]
[value,price,100]
[value,color,[blue,red,white]]

```

図3. KORE/KRとKORE/DBの実行例

ていることを示している。内部表現は以下のようなテーブルとして表現される。

civic	slot	facet	value
	value	power	83
	value	weight	90
	value	displacement	(1500,1300)
	value	price	100
	value	color	(blue,red,white)

5. おわりに

KOREのサブシステムにおける以上のような関係テーブルを用いた内部表現の統一化はKOREの強力な問題解決支援機能を実現する上で重要な働きを担うことになる。なお、本研究は第5世代計算機プロジェクトの一環として行った。

参考文献

- (1) Kowalski,R.: Logic for problem Solving, Elsevier North Holland, 1977,pp.31-44.
- (2) Deal,V.: On Database Systems Development Through Logic, ICF Transactions on Database Systems, Vol.7,No.1,March 1982,pp.102-123.
- (3) Foray,C.L.: OPS5 User's Manual, CHU-CS-81-135,July 1981.
- (4) Robrow,R.G. and Stefik,M.J.: The LOOPS Manual,XS-VLIIS1-81-13,1982.
- (5) London,P.E.: A dependency-based modelling mechanism for problem solving, AFIPS Conference Proceedings, Vol.47, 1978,pp.253-247.
- (6) 新谷・片山: 関係テーブル(その実質) — 知識ベースにおける知識表現方式とその応用 —, 情報31期全国大会, 1P-8, 1985.
- (7) 新谷, 他: 問題解決環境KORE(その2) — 知識型利用環境KORE/EDENとその応用 —, 情報32期全国大会, 5L-9, 1985.
- (8) 片山, 他: 問題解決環境KORE(その3) — 知識表現サブシステムKORE/IEとその応用 —, 情報32期全国大会, 5L-10, 1985.
- (9) 新谷, 他: 問題解決環境KORE(その4) — 関係データベースシステムKORE/DBとその実質 —, 情報32期全国大会, 5L-11, 1985.

照)。この行列転置圧縮は、主に、テーブルを表現するのに必要なアサーションの数を減らすために用いられる。



図3. KORE/EDEN の管理機能

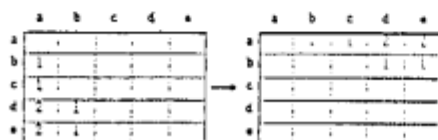


図4. 行列転置圧縮の例

図4の例では、行列転置圧縮によりアサーションの数を2つ減らしたことに相当する。つまり、行列転置圧縮前では行b, c, d, eに関するアサーションが必要であったが、その後では行aとbに関するアサーションを持てば良い。関係再配置圧縮は知識テーブルの行及び列における順番を入れ換えることによりその圧縮機能を達成する(図5参照)。この圧縮は知識テーブルそのものに必要な記憶スペースを削減する(これにより、アサーションの第二引数におけるリストの要素数を減らすことができる)。

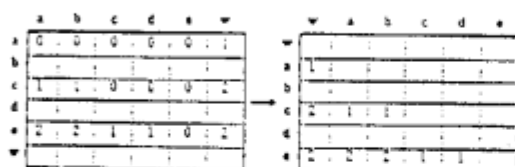


図5. 関係再配置圧縮の例

(5)のネットワークループチェック機能は知識テーブルに関係を付加した場合にそこから得られるネットワークにループが有るかどうかをチェックする機能である。このループ発見過程は図6のように示すことができる。

図6は図右上のループを形成するネットワークを知識テーブルに保存して行く過程を示している。ここで、到達可能性セットとは到達可能性を親から伝播させたり、子供に伝播させることをいう。到達可能性の伝播は関係タイプ2を知識テーブル要素に立てることに相当する。これら伝播

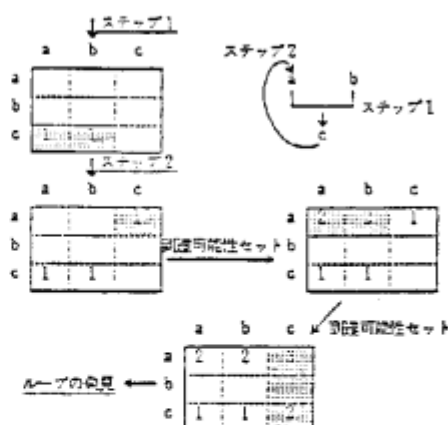


図6. ループ発見過程

操作のための操作(知識テーブル要素に2を立てること)はPrologにおける論理ビットORを用いて行う。そこで、このように自分自身に到達可能性がある知識の知識テーブルの列要素の中には整数3が生じる。なぜなら整数1(隣接関係タイプ)と整数2(到達可能性タイプ)の論理ビットORは整数3であるからである。つまり、このような整数3が生じるかどうかをチェックすることによりネットワークにおけるループ構造の発見を達成する。

4. KORE/EDEN の応用

問題解決において、一般にgenerate and test による問題解決法では、ある推論過程が失敗したら、最も最近の選択点までbacktrackingする必要がある。またそこから他の方法を試す必要がある。これは非常に効率が悪い。しかし、MOLGENで用いられているように推論過程(dependency records)を保存することにより、さらに効果的に失敗を修正することができる。この保存された推論過程に基づくbacktrackingは、dependency-directed backtrackingとよばれ、信念の修正等を実現するための有用な機能である。推論過程は、一般的に、ネットワーク構造を形成する。

KORE/EDENはこのような構造を素直に管理ができ、暗黙値推論やtruth maintenance system等を構築する上で重要な機能を提供する。

5. おわりに

KORE/EDENはネットワーク構造を効果的に管理するために設計され、知識利用の側面において効果的に機能することが期待できる。なお、本研究は第5世代計算機プロジェクトの一環としてICOTの委託で行ったものである。

参考文献

- (1) Hayes-Roth, F.: The Knowledge-Based Expert System: A Tutorial, *COMPUTER*, September, 1984, pp. 11-28.
- (2) 斎藤・片山: 知識テーブル(その実用) 一般データベースにおける知識表現方式とその応用, 第3回全国大会, 12-8, 1985.
- (3) Shitka, T.: Knowledge Table: An Approach to Knowledge Utilization Mechanism for Knowledge Base Failure (IAS-515 Research Report, 1986, vol. 38, (in preparation))

問題解決支援環境KORE (その3)

-知識表現サブシステムKORE/KRとその概要-

5L-10

*片山 佳則 新谷 虎松 平石 邦彦

富士通 国際情報社会科学研究所

1. はじめに

汎用を目指した問題解決環境を実現するためには、その問題・状況・処理方法などに於ける知識が(機能的にも構造的にも)柔軟に表現されていなければならない。

このように知識を表現するための表現・利用機構としてオブジェクト表現と論理型表現を中心としたシステムKORE/KRを開発している。これ自身は、単独な表現システムであるが、問題解決支援環境KOREにおいて、プログラミングパラダイムを提供する基本構成要素として利用可能である。

様々な知識を表現する場合には、その知識が持つ構造的な関係がシステム内で、どのように実現され、それがどのような意味を持つかなどの概念間の関係の意味を明らかにしなければならない。またどのような知識がどこに表現され、どのように利用できるかなどの対応が取れることも重要である。KORE/KRはこれらのための環境も備えている。

本稿では、この表現システムにおける概念間の関係を明らかにし、その機能の実現を中心として、表現・利用の環境について報告する。

2. KORE/KRにおける関係表現

知識表現を考える場合、その構造を表すためのネットワークとそれが持つセマンティクスが問題となる。システム側の構造に対して、人間が考えている概念を対応させることは容易なことではない。オブジェクト表現に限らず、他の表現においても、表現する知識の間の関係として様々なものを挙げることが出来るが、基本としてIS-A関係とPART-OF関係がある。これらは、一般には、包含関係と部分・全体の関係として考えられている。

KORE/KRでのオブジェクト表現では、概念は、全てクラスオブジェクトとして表現される。IS-A関係については、このクラス間(概念間)の階層関係表現として導入されている。

KORE/KRにおけるすべての関係は、概念の関係を表している。各概念の意味は、その結合関係の概念構造で規定される。

概念のレベルと実体のレベルに分けられるKORE/KRでは、次のように基本的な結合関係を与えている。

(1) 階層関係を表す結合は、IS-A関係である。各概念が持つ厳密な継承方法によってその概念自体の意味が決まっている。

(2) 機能的関係を表す結合は、PART-OF関係である。これにより、概念が持つ操作の受け渡し(分担)方法が決まり、概念自体の意味が表される。

これらの関係を用いて、知識を複数のクラス(概念)に分割して規定し、情報の隠匿などを利用し、矛盾無く表現する。階層関係を表す結合(1)には、多重表現が可能である。

これらの2つの関係は、(1)の結合の中では、満足のいく知識を見つけた時点で終了し、(2)では、結合されているすべてが対象となっていることから、KORE/KRでは、IS-A関係がORであり、PART-OF関係がANDと考えられる。

実際には、これらの意味は、その概念(クラス)が持つ個々の動作・実行内容により規定される。

概念レベルで構成したこれらの結合関係を利用する為には、概念に対応するオブジェクト(インスタンス)を実体レベルに持たなければならない。概念レベルと実体レベルの結合関係はINSTANCE-OFである。この関係によって実体レベルのオブジェクトが概念レベルで規定された機能・動作を利用出来る。包含関係など、全ては概念レベルで表されるため、INSTANCE-OF関係に多重な関係は必要としない。従って概念レベルでの結合関係で規定された概念オブジェクトの構造が重要となる。

2-1 概念レベルの結合関係

IS-A関係が表す概念の階層によって、知識は必要に応じて継承される。継承の対象となる内容は、上位の概念が持つslot・method内の知識や、PART-OF関係である。

上位を多重に持っていた場合は、検索の制限に従ってIS-A関係は初めに見つかった内容のみを受け継ぐ。複数の内容を受け継ぎたい場合には、原則としてPART-OF関係を用いて表現しなければならない。

但し、PART-OF関係を用いることで、概念間の意味的な対応が取れない場合には、各メソッドに対してIS-A関係での継承(検索)の制限を変更させることができる。

○ 継承(検索)の制限

・IS-A関係の定義によって表現される階層的ネットワークをleft-to-right・depth-firstで進められる。但し、各methodに対しての指定により次のように変更することができる。

① IS-A関係のネットワークで多重に表現された情報を全て継承する。この場合、継承結果の副作用は全てANDで処理される。

② left-to-right・right-to-leftの切り換えをする。

これらの方法は、Flavorのconvine-method⁽²⁾を参考にしているが、methodを実行する際にFlavorが持つdemon-methodの働きは持たない。

PART-OF 関係を引き継いだ実体オブジェクトに対して partに対応するすべての実体オブジェクトが必ずシステム側で生成される。

partとして生成された実体オブジェクトは、全体であるオブジェクトからの機能分担として動作し、他から直接メッセージを受け取ることは出来ない。

2-2 結合関係の実現

オブジェクト間の関係表現は、概念レベルで行われているため、これらのオブジェクトをどう実現するかによってほとんどが決まる。概念レベルのオブジェクトは次のように定義される。

```
defclass <オブジェクト名> ::  
  super-set <IS-A階層の上位オブジェクト集合> ;  
  parts <PART関係のオブジェクト表現> ;  
  value-set <スロット定義> ;  
  attach-method <付加手続きの記述> ;  
  method <メソッドの定義> .
```

super-set で、IS-A階層の直上の上位オブジェクトをすべて次の例のように記述する。

```
super-set OBJECT, WINDOW, CHECK ;
```

(英大文字はオブジェクトを表す)

KORE/KR の特徴として、PART-OF 関係の記述は、parts に、各部分をどのオブジェクトで分担するかをペラにする。

```
parts right-wing - WING,  
      left-wing - WING ;
```

parts を定義したクラスオブジェクト内の表現 (メソッド・付加手続きなど) では、parts で定義された各部分のオブジェクト (right-wing, left-wing) を、他のオブジェクトと同様に扱うことが出来る。各部分をどのように利用 (機能分担) するかは、parts 定義をしたオブジェクト内で定める。従って各部分への操作は、parts 定義のオブジェクトを通してしか行えない。オブジェクトのスロット・メソッド・付加手続きなども所定の方法で記述する。

3. 表現・利用の環境

前回、いくつかの機能を示した⁽¹⁾が、その他KORE/KRの利用に関していくつか取り上げる。

a) オブジェクト定義で示したようにスロットに付加手続きの記述を行える (attach-method)。メカニズムは、外部からのメッセージによりメソッドが起動されるのと同様、オブジェクト内のスロットから付加手続きの起動がメッセージ送信の要領で実行される。これは、partへの操作と同様にオブジェクト内でのメッセージ送信と考えられる。

b) すべてのオブジェクトは、外部ファイルとして予め作成しておき、file-manipulatorオブジェクトにメッセージを送ることで変換される。その際に副作用としての付加手続きの起動は、様々なオブジェクトへ影響を与える。そのためfile-manipulatorでは、全オブジェクトの変換・作成が終了してから、付加手続きが起動される。これによって未定義オブジェクトへの副作用などのエラーを無くす。

c) KORE/KR は、method-managerオブジェクトが組み込まれており、これによってメソッドに関する情報の管理を行っている。情報はいつでもメッセージにより利用できる。

▷メソッドがどのオブジェクトにあり、どのオブジェクトから利用出来るか。

▷オブジェクトはどのようなメソッドを持ち、その他にどのようなメソッドを利用できるか。

これらの情報により、オブジェクトの概念関係を利用して、メッセージによる実行を円滑に進めることが出来るようになる。method-manager は、file-manipulatorと連結し、メッセージを受けることで各メソッドの状態を記録していく。

現在、C-Prolog 上のKORE/KR がシステムとして持っているオブジェクト集合は次の5つに分類される。

- (1)<META> : クラスオブジェクト全体を統轄する役割。
- (2)<UTILITY> : クラスオブジェクトの階層の中で、上位レベルにあり、様々なユーティリティの機能を持つ。
- (3)<LOAD> : 外部からのオブジェクトの入力を行い、変換をする。

(4)<METHOD> : methodの情報を記録・管理する。

(5)<WINDOW> : windowシステムの働き。

これらがシステム起動時に組み込みとして入力される。

(4)と(5)に対しては、必要に応じて入力を選択可能。

4. おわりに

概念間の関係を明らかにし、KORE/KR の利用環境についてその概略を述べた。詳しい応用まで述べていないが、このシステムは幾つかのProlog上に実装されている。現在、PART-OF 関係を用いた構文の記述を取り上げている。

また、付加手続きを主に利用する表現は副作用としての複雑さを導くことになるため、注意が必要である。

尚、本研究は、第五世代計算機プロジェクトの一環として、ICOTの委託で行ったものである。

(参考文献)

- (1) 片山、新谷 “知識テーブル (その利用) - 知識ベースにおける対象指向表現とその処理機構の試作 - ”, 第31回情報処理全国大会, (1985) p.1233-1234
- (2) D.Weinreb, D.Moon “Lisp Machine Manual” 4th ed., Symbolics, Inc. (1981).

問題解決支援環境 KORE (その4)

—関係データベース・サブシステム KORE/DBとその概要—

5L-11

*平石 邦彦 新谷 虎松 片山 佳則

富士通㈱ 国際情報社会科学研究所

1. はじめに

論理型プログラミング言語であるprologは関係データベースの表現に適しており、それ自体で高度なデータベース問い合わせ言語として見る事ができる。本発表ではprologのこのような特徴を利用した関係型データベース操作言語について報告する。

2. prologと関係データベース

関係モデルでは、有限個の順序付けられたデータの組の集合を関係(relation)といい、この関係によりデータベースを表現する。また、各々のデータの組のことをタプル(tuple)と呼ぶ。

まず次のような関係suppliersを考える。

```
suppliers [sno,sname,status,city] =
  {(s1,smith,20,london),(s2,jones,10,paris),...}
```

ここで、sno,sname,status,cityは属性名をあらわしている。この関係はprologの事実(fact)により次のように表現できる[1]。

```
db_suppliers(s1,smith,20,london).
db_suppliers(s2,jones,10,paris).
:
```

このとき検索は次のように行うことができる。

質問: statusが20より大きくlondonにいるsuppliersのsnoは?

```
!- ?db_suppliers(X,_,Z,london),Z>20.
```

このように、prolog自体でデータベース言語としての機能を持っている。

3. システムの構成

データベースへの問い合わせにおいて、prologそのままだではユーザーにとって使いやすいものではない。そこで、本システムではprolog上にSQL[2][3]に似た文法を持つデータベース言語をのせ、それをprologプログラムに翻訳して実行するという形式をとっている。したがってユーザーはprologを意識することなくデータベースを操作することができる。

構成はFig.1のようになっている。コンパイラは、問い合わせのprologへの変換および述語の並び換えによる最適化を行う。データベース操作基本述語群は、タプルの追加・削除・更新など基本的なデータの管理

を行う。データベースに対する操作は、すべてこの述語群を通して行われる。

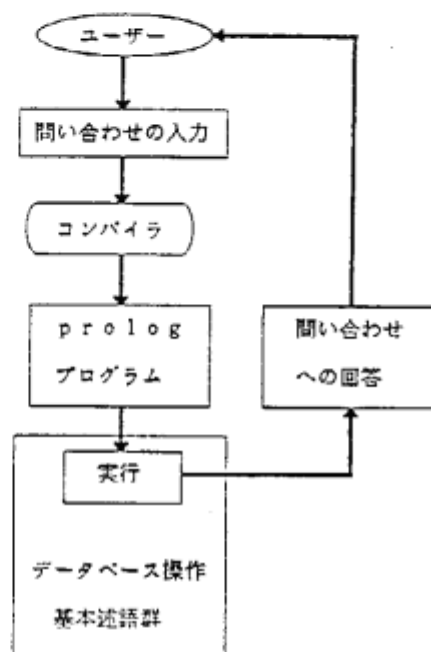


Fig. 1 KORE/DBの構成

4. 問い合わせの変換

次の問い合わせを考える。

```
select sno from suppliers
where city=london
and status >
  (select status from suppliers
   where sname=smith)
and sno in
  (select sno from shipment
   where qty>300).
```

これは次のようなprologプログラムに変換される。

```
ans(A) :-
  db_suppliers(_,smith,C,_),
  db_suppliers(A,_,D,london),
  db_shipment(A,_,B),B>300,D>C.
```

selectブロックにおけるwhere 以下の条件は連言形をしているものとする。このとき各条件を次のように分類する。

- 等号条件 (=, in) sname=smith, city=london,
sno in (select ...)
- 非等号条件
- 独立条件 1つの関係の属性についての条件,
qty>300
- 関連条件 複数の関係の属性を含む条件,
status>(select ...)

各属性値の取り出しは、述語db_関係名 (...) によって行われる。等号条件はdb_関係名 (...) の変数に統合化される。独立条件は、各関係の値を取り出す述語の直後に置かれる。関連条件は、すべての属性の値を取り出した後に置かれる。

5. 問い合わせの最適化

prologが問い合わせを評価する際の効率性、連言中に言われた目標(goal)の順序に依存する。したがって、問い合わせの目標の順序を入れ換えることにより、問い合わせの最適化を行うことができる。ここでは、述語db_関係名 (...) に対し次のような価格の式を導入している[3]。目標となる述語に対応する関係の要素数を、値を代入された引数の定義域の大きさの積で割ることにより価格が求まると仮定する。

$$COST = \frac{|R|}{|D1| \times |D2| \times \dots \times |Dn|}$$

ここで、|R|: 関係のタプル数

|D|: 引数がとる定義域Dの大きさ

である。この式によりdb_関係名 (...) の順序を動的に入れ換える。その手順は以下の通りである。

- ① 各db_関係名 (...) の価格を計算し、最小のものを選出す。
 - ② ①で選んだdb_関係名 (...) の変数すべてに値を代入し、残りの目標について再びこの手順を繰り返す。
- 次に評価すべき目標は、その時点での最小価格を持った目標となる。また、1つの目標の選択は残りの目標の選択に影響を与えることになる。

6. 規則の利用

prologは事実だけでなく規則(rule)も扱うことができる。ここでは規則を用いることにより実現されている機能について説明する。

(1) ビュー(view)

ビューは関係型データベース内の関係に対して必要な属性やタプルの集合を定義した仮想的な関係である。

これはprologの規則を用いることにより簡単に実現できる。たとえば2つの関係

```
suppliers [sno,sname,status,city]
shipment [sno,pno,qty]
```

からsupl [sno,sname,pno] という仮想的な関係を定義するには、

```
db_supl(A,B,C) :- db_suppliers(A,B,_),
                  db_shipment(A,C,_).
```

とすればよい。これはdefine_view コマンドにより実現される。suplはsuppliers とshipmentの等結合になっており、検索に関しては通常の関係と全く同様に扱うことができる。また、

```
db_a(X,Y,Z) :- db_b(X,Y), Z is (X+Y)/2.
```

のように、値を計算させることも可能である。

(2) 手続きの実行

ビューの場合と同様な方法により、あるタプルを検索したときに、自動的に手続きを実行させることができる。たとえば、

```
db_suppliers(s1,smith,20,london) :-
    method(s1,london).
```

とすれば、このタプルを検索したときに、手続きmethod(s1,london) が実行される。

7. おわりに

KORE/DBはprologをデータベースとして利用する際のユーザー・インターフェースとして設計されており、検索等の処理はすべてprolog自身が持つ機能を使っている。したがって、データベースとしての能力もprolog処理系に依存する。

なお、本研究は第5世代計算機プロジェクトの一環として、ICOTの委託で行ったものである。

(参考文献)

- [1] R.Kowalski, "LOGIC FOR DATA DESCRIPTION," In H.Gallaire(ed.), "LOGIC AND DATA BASES," PLENUM PRESS, New York(1978).
- [2] M.M.Astrahan, "System R: Relational Approach to Database," ACM Transactions on Database Systems, Vol.1, No.2, pp97-137(1976).
- [3] D. リー, "PROLOGデータベース・システム", 近代科学社(1985).
- [4] M.M.Astrahan, "Implementation of a Structured English Query Language," Communications of ACM, Vol.18, No.10, pp580-589(1975).

5U-4

方式設計段階における
性能評価支援の一検討

延田 規之 野田 泰徳 榎村 昌俊 平野 達郎
沖電気工業（株）

1. はじめに

論理装置の設計において、方式設計の段階で実機を作る前の正確な性能評価を行う事が、品質向上、設計期間短縮の面から望まれている。本論では、パイプライン処理方式の計算機をターゲットとし、その命令仕様、データバス構成、データバス構成上での命令仕様の実現方式（データ転送路）が決定した設計段階において、それらを用いて当該装置の性能を評価する為のデータを導出し、性能を向上させる為の改善案を見出す際の支援を行うツールについて検討した。更にパイプライン処理方式のモデル化について並列処理型装置を用いて検討した。

2. システム概要

パイプライン処理方式計算機の方式設計段階における性能評価項目の一つとしては、命令実行の並列度が考えられる。並列度の高い処理方式はハードウェアファシリティを有効に活用した、パフォーマンスの高い計算機を実現出来る可能性を持っている。その命令実行の並列度から計算機を評価し、性能を向上させる為の修正点を絞り込む為には、図1の様に各命令の各ステージが並列に実行されている状態（図1）を調べ、レジスタの利用に競合が起るクリティカルなデータバスを発見する事が有効であると考えられる。これを基に命令の実行待ちを最小にする修正を命令仕様やデータバス構成に加える事によって性能の向上を図る事が出来る。

本システムでは、上記の評価と性能の向上を図る際の支援を行うために、入力として、テストプログラム、データバス構成、各命令の各ステージでデータバスを構成するファシリティをどの様に用いるかを記述した命令定義記述を用い、データの流れのシミュレーションを行う。そして、各命令の各ステージが並列に実行されている状態（実行制御データ）と各ハードウェアファシリティにおいて利用競合が起る状態（競合データ）とを表すタイムチャート、Mips値、テストプログラム実行時の並列度、更に専門家の知識を組み込んだエキスパートシステムによってコンサルテーションされる設計変更の指針等が出力される。

3. システム構成

システムは図2の構成となる。シミュレータ部は入力データを基にシミュレーションを行い、リアルタイムでトレースデータをタイムチャート形式でコンソール等に出力する。post_processorは、シミュレーション終了後にMips値、命令実行の並列度等の統計データを出力する。adviserは設計変更の指針を示す知識処理モジュールである。

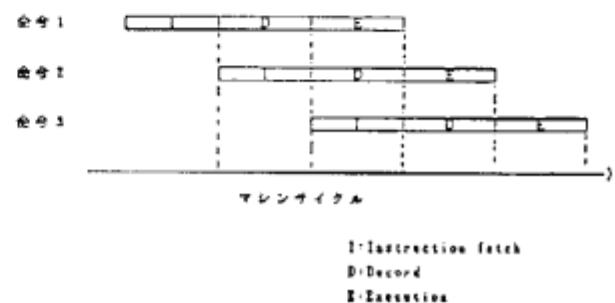


図1 パイプライン処理方式の並列実行

4. 並列論理型言語によるモデル化

4.1 G H C

G H C (1) は並列に存在するプロセスを動的に生成し、各プロセス間で通信や同期の制御を行う事が出来る。これにより並列実行するパイプラインの動作のモデル化が容易であると考えられる。更に論理型言語であるため、知識を直接表現することが可能である。以上よりパイプライン処理方式のモデル化に G H C を用いた。

4.2 シミュレータの構成

現在試作中のシミュレータ部の構成を図3に示す。命令実行モジュールは、命令定義記述に従って命令を実行し各ステージに於て利用するファシリティをリストにしてデータベース管理モジュールに渡し、実行終了メッセージの到着を待つ。データベース管理モジュールは、受け取ったリストに従ってファシリティの利用割合を解消しながら、データベースのシミュレーションを行い、それが終了すると実行終了メッセージを返す。

4.3 命令実行モジュールのプロセス構成

命令実行モジュールのプロセス構成を図4に示す。stage_process は命令を実行して行く主体、message_manager は各 stage_process 間のメッセージの管理、timer_process は各 stage_process の同期を管理するプロセスである。図は3段パイプラインの例である。これを G H C で記述したのが図5である。シミュレータの起動は述語 pipeline の呼び出しによって行う。起動されると、body 部に示された5つのプロセスが同時に生成され、処理を行なう。

4.4 記述

パイプライン方式の並列実行する各ステージを G H C のプロセスに割り当て、プロセス間の同期を G H C の同期機構を用いて容易にモデル化することができた。

5. まとめ

本システムは複雑な並列動作を行うパイプライン方式計算機の制御部とデータベース機構の設計を支援する事を目標としている。現在シミュレータ部のプロトタイプ化を行っている段階であり、他の部分の詳細決定等は今後の展望である。

尚、本研究は第5世代コンピュータプロジェクトの一段として行われている。西沢計取いた I C O T 第一研究室吉川重良及び第二研究室諸氏に感謝する。

(参考文献)

- *1 Kazunori Ueda "Guarded Horn Clause", ICOT Technical Report, TR-103 (1985)

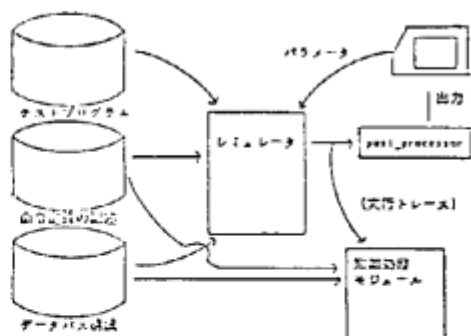


図2 システム構成

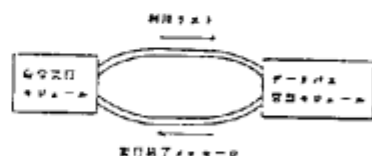


図3 シミュレータ部の構成

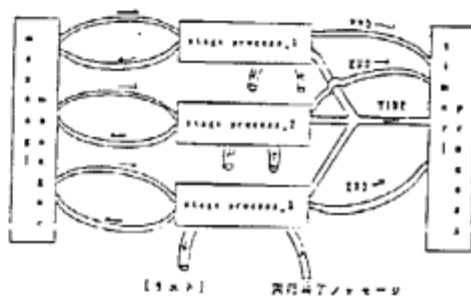


図4 命令実行モジュールのプロセス構成

```
pipeline(list(A,B,C).end(X,Y,Z)):-true |
  timer(0,T,F1,F2,F3),
  message_manager(in(U,V,Y).out(P,Q,B)),
  stage(1,[0|T],F1,U,([],P),A,Z),
  stage(2,[0|T],F2,V,([],Q),B,Y),
  stage(3,[0|T],F3,W,([],B),C,X).
```

図5 命令実行モジュールの記述例

知識表現／推論システムPeaceと その故障診断問題への適用

4L-9

古関 雅幸^{*} 和田 慎一^{*} 西田 哲朗^{**} 若杉 暢隆^{***} 近藤 真己^{***}^{*} 日本電気㈱ C&Cシステム研究所, ^{**}同 複合交換開発本部,^{***} 日本電気技術情報システム開発㈱

1. はじめに

電子装置の社会的役割は年々高まっており、例えば電子交換機では、障害発生時にはサービスに影響を与えずに速やかに復旧することが要求されている。このために、局用デジタル式電子交換機は、オンライン障害検出機能、装置の二重化、及び自己診断機能を備えている。しかし、その自己診断の能力には限度があり、診断できない障害が発生した場合には専門家による判断が必要である。また、他の電子装置に比べて寿命が長く、保守の専門家を長期に確保することが困難である。このような背景の中で、保守の専門家と同等の知識を持ち、その代わりとなるシステムの開発が望まれる。本稿では、現在研究開発を進めている電子交換機の故障診断エキスパートシステムについて、その基本となる知識プログラミングシステムを中心に述べる。

2. 故障診断における専門家の知識

電子交換機のような大規模なシステムの障害を完全に診断可能な診断プログラムの作成は困難な問題である。例えば論理ゲートのみからなる組み合わせ回路においてでさえ、障害を構成部品の単一縮退故障に仮定したとしても、完全な故障診断を行なう診断プログラムを作成する問題は複雑な組み合わせ問題となることが知られている。その上、実際のシステムははるかに大規模で複雑な順序回路であり、障害の種類もショート故障、開欠故障等、種々であり、かつ装置外部の状況も障害に影響する。

このように困難な問題に対して、人間は、診断プログラムでは診断できなかったような障害に対しても、障害時の状況、障害の出方、部品の壊れ易さ、回路構成、等を柔軟に考慮しながら種々のテストを行ない、故障箇所を推論していき、最終的には障害を取り除く能力を持つ。このような能力を持った保守の専門家の知識を計算機上に表現し同様の判断を

行なうシステムを実現するために、専門家のインタビューを通して診断の実態を調査した。

専門家は、まず、①症状、障害時の状況等の情報を集め、被疑部分を思い浮かべる。次に、②その被疑部分の確定あるいは切り分けのために有効な種々のテスト方法を考え、実行する。その時に、部品の壊れ易さ、テストのサービスへの影響も考慮に入れている。そして、③そのテストの結果をもとに被疑部分の更新を行なう。④、⑤を、交換可能単位までに切り分けが出来るかあるいはそれ以上切り分けが出来なくなるまで繰り返し行なう。そして、⑥その被疑範囲の部品の交換を行なう。最後に、⑦その交換により故障が完治したかどうか確認し、⑧完治しない場合は考え直す。このように人間は種々のテストを行ないながら被疑範囲を狭めていき、最終的には部品の交換によって診断を行なっているといえる。このような診断の方法は一般の電子装置について共通であるといえる。

このような診断をするシステムの実現には、従来の医療診断エキスパートシステムのように、症状と原因をルールで記述しただけの知識では不十分である。①の症状についての知識は各症状をオブジェクトとした表現形式が適している。②のテスト選択の知識は、条件部に被疑部分の分布を、動作部に選択するテストを記述した前向き推論ルールが適している。また、③のテスト結果判断の知識にも、条件部にテスト結果を、動作部に被疑部分を記述した前向き推論ルールが適している。④で使用する部品の知識は各部品、装置をオブジェクトとしたネットワーク構造の表現形式が必要である。

3. 知識プログラミングシステムPeaceの概要

多様な知識を表現し利用するための知識プログラミングシステムとして、Prolog言語上にPeace (Prolog based Expert Applications Environment)

を開発中であり、これを用いて診断知識ベースを試作している。これは、Prolog言語上に、オブジェクト指向の表現、前向き推論プロダクションルールの記述、マンマシンインタフェース作成用のライブラリを実現したものである。Prolog言語上にその自然な拡張として実現しており、Prolog言語の持つ利点をそのまま利用しながら、これらの表現形式を混在して用いることができる。

3. 1 オブジェクト指向の表現

オブジェクトは、一つの概念、事物等に関する性質をPrologの節の集合で表わす。これは他のProlog上のオブジェクト指向表現〔中島84、近山84〕と基本的に同様である。但し、

- ・ オブジェクト間の任意の関係及びそれを通したマルチプルインヘリタンスの定義を、関係を表わすオブジェクトを使用して自由に行なえる。クラス、インスタンス、メタクラス等の関係も一般的な関係の一つとして扱う。

- ・ オブジェクト変数及びメソッドの記述をオブジェクト内に定義された一引数の節及びボディ部を持った節に対応させており、Prolog言語との親和性がよい。メソッド起動はそのオブジェクト内に定義されたボディ部を持つ述語の呼び出しに等しい。

- ・ 逆関係の更新を自動的に行なう。

等の特徴を有している。例えば交換機のパッケージについての知識を以下のように記述する事が出来る。

```
equipment_A : category(category_1).
package_of : $inherit(category).
package_1 : a_kind_of # package;
             package_of # equipment_A;
             package_No(pkg001).
package_2 : a_kind_of # package;
             package_of # equipment_A;
             package_No(pkg002).
package : disp_No :- package_No(X) in origin,
              display(X),nl.
```

ここで例えば、categoryがcategory_1であるようなパッケージを全て求める処理は、Prologの組み込み述語 setofを用いて、

```
?-setof(X, (has_kind # X in package,
            category(category_1) in X), S).
S = [package_1, package_2]
```

のように求まる。ここで、述語inがオブジェクト内の節を操作する述語で、参照、メソッド起動等を行なう。これは、Prologで書かれたPrologインタプリタであるdemo述語〔國藤85〕を拡張して実現された。関係has_kindは関係a_kind_ofの逆関係であり、システムにより自動的に作成されたものである。関係package_ofを通して、オブジェクトequipment_Aのcategory(category_1)という性質がオブジェクトpackage_1とpackage_2に継承されている。関係package_ofはpackage_ofというオブジェクトに定義されており、\$inherit(category)により、categoryという述語を継承するように定義されている。

3. 2 前向き推論プロダクションルールの記述

前向き推論プロダクションルールは、同時に適用する複数のルールをまとめてルールブロックと呼ぶオブジェクトに記述する。このため、他のオブジェクトと同様に管理でき、また、効率のよい推論が可能になる。ルールの条件部及び動作部はPrologのゴール呼び出しの形式で記述する。例えば④のテスト選択の知識を表わすブロックは、

```
test_selection_rules :
    $control (do_1);
    rule1 if test1 を選択する条件
           then test1 を実行;      ....
```

のように記述できる。ここで、\$controlはこのブロック内でのルール実行の制御形態を表わす。

4. おわりに

電子交換機の故障診断における専門家の知識を整理し、その表現形式を考察した。Prolog言語上に、知識プログラミングシステム、及び知識ベースを試作中であり、今後その有効性を確かめていく。本研究は第5世代コンピュータ・プロジェクトの一環として進めているものであり、日頃御指導頂くICOT岩下室長に感謝します。

【文献】
 〔中島84〕 Prolog/KR における知識表現、第28回情報全大
 〔近山84〕 ESP Reference Manual, ICOT TR-044
 〔國藤85〕 論理型言語Prologによる知識ベースの管理, Logic Prog. Conf 85

仮定の修正を行なう推論方式の 故障診断への応用

4L-10

和田 慎一^{*} 古閑 義幸^{*} 若杉 暢庸^{**}^{*} 日本電気㈱ C&Cシステム研究所^{**} 日本電気技術情報システム開発部

1. はじめに

複雑な電子装置の保守の支援の必要性から、筆者らは電子装置の故障診断を行なうエキスパートシステムの研究・開発を行なっている[1]。故障診断では症状や装置に対するテスト実行の結果などの情報をもとに故障の原因についての判断が行なわれるが、「大体・・・だろう」という判断も行なわれる。通常はこのような不確かな判断も合わせて推論を進め、誤りがあるとわかったときにそれを修正するという方法がとられる。このような形の推論をエキスパートシステムとして実現するために、通常正しいと思われることを仮定して推論を進め、矛盾の検出された場合に誤った仮定を修正する推論方式の利用が考えられる。本稿ではこのような推論方式の故障診断における必要性、適用の方法、および矛盾が検出された際の仮定の制御の方法について述べる。

2. 故障診断における仮定に基づく推論の利用

2.1 故障診断における不確かな判断

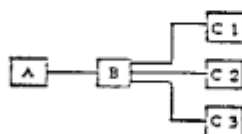
故障診断は、①テスト（故障に関するデータ・情報の収集のためのアクション）を選択する、②テストを実行する、③テスト結果として得られた情報をもとに故障原因として考えられる範囲について判断する、という処理の繰り返しとして進められる。この中で③の判断には種々の不確かさが含まれる。

有効なテストの数が限られている故障診断においては、得られたテスト結果のデータをもとにして故障原因として考えられる範囲をより狭く絞り込むために、不確かさを含む判断も合わせて用いる。判断に不確かさが含まれているとしても絞り込まれた範囲に該当する部品を交換して故障が直ればそれで診断を終了することができる。

2.2 不確かさを含む判断の例

不確かさを含む判断の例を挙げる。図1のような構成の装置においてC1、C2、C3が同一な構成の部

品であり、A-Ci (i=1, 2, 3) 図1 装置の構成間では定常的に信号の伝達が行われているものとす



る。A-Ci間の信号の伝達について異常検出の機能があるが、異常が存在するとともにタイミングや実行状況等の確率的条件が満たされた場合に限り異常が検出されるものとする。このような装置において、A-C3間のみで異常が検出され、A-C1、A-C2間では異常が検出されていない状況を考える。異常に関係している部分Aに関しては「Aが故障の原因であるならばおそらくA-C1、A-C2間でも同様に異常が検出されることになるはず」と考えられるが、これは対象としている状況と合わない。このことから「Aは故障の原因ではない」という不確かな判断が行われる。

2.3 仮定に基づく推論の利用

不確かさを含む推論を行なう以上、故障原因として考えられる部品を交換しても故障が直らないことがある。その場合への対処のためにdata dependency[2]あるいはtruth maintenance[3]の利用が考えられる。その利用は次のように行なう。

①判断において不確かなことを仮定として表わし、通常はそれを信じ、それに基づいて推論を進める。
②推論実行の履歴を記録しておく。これにより、導かれた各事実がどのような仮定に基づくのか（dependency）を調べられる。また仮定がくつがえされた際にそれに基づいて導かれた事実を消去できる。
③故障を起こしていると考えられる部品を交換しても故障が直らない場合や事実間に起こり得ない関係が生じた場合を矛盾とする。

④矛盾が検出された場合、前提となっている仮定を疑い、くつがえす（dependency-directed backtracking）。これによって矛盾が解消される。

2.4 結果を予想できるテストの実行の省略

仮定に基づく推論を行なうことによって、さらにテスト実行の制御に関して次のような処理を実現することができる。

①実行のコストが大きいテストのうち、その結果について予想できるものがある。通常時はテストの実行を省略し、代わりに予想されることを仮定として信じ、それに基づいて推論を進める。

②矛盾によってその仮定が疑わしくなった時に省略

していたテストを実行し、仮定の真偽を確認する。

2. 5 確率的推論に基づく処理との比較

故障診断における不確かさに対処するための方法として、確率的推論を用いて各部分あるいは各部品ごとに疑わしさを表わす値を評価し、その値の大きな部分から交換していく方法も考えられる。しかし、この方法では、推論の不確かな部分に戻ってその確認のテストを行なうこと、検出された矛盾に応じて、事実関係の一貫性を保ったうえで事実を修正することなどが困難である。

3. 仮定の制御

3. 1 仮定の制御の方針

仮定に基づく推論を採用する場合、検出された矛盾を解消するために、矛盾の前提となっている仮定のどれかをくつがえすことになるが、どの仮定をくつがえすかを決定しなければならない。また、ある仮定をくつがえすことにより、くつがえされていた別の仮定を再び信じることが出来る（信じて矛盾を生じさせない）場合がある。適用する問題の性質から仮定の制御の方針としては、次の2つに基づくものとする。

①より多くのより確度の大きな仮定を信じる

矛盾が検出された場合でも、互いに独立な仮定を矛盾を含まない範囲でより多く信じることによって、不確かながら有効であるような推論規則がより多く用いられるように保つ必要がある。このため、あらかじめ各仮定に確度を表わす数値を与えておき、矛盾が検出されたときには、矛盾を含まない範囲で、できる限りより確度の大きな仮定をより多く信じるように保つ、つまり確度の和が最大となるようにするものとする。

②疑わしい仮定の真偽について調べるテストの実行

テストの省略（2. 4で説明）のために代わりに立てられた仮定が検出された矛盾の原因として疑わしくなった場合、省略していたテストを実行して仮定の真偽を調べ、その結果を利用する。

3. 2 仮定の制御の実現法

仮定の制御の方針①の実現法として、単に「矛盾の前提となっている仮定の中で確度の最小のものをくつがえす」という方法が考えられるが、その方法は十分なものでない。このために、同時に信じ得る仮定の最大限の組み合わせ（矛盾を含まない仮定の組み合わせ）の集合を保持し、矛盾が検出された際にはその前提となっている仮定の組み合わせに応じてその集合を更新するという方法をとる。この集合

をC.B.S. (Consistent Belief Set) と呼ぶ。そして、常にその中で確度の和が最大になるような仮定の組み合わせにしたがって仮定を信じるようにする。

この処理の例を図2に示す。推論実行中にSTEP1で仮定A, Bの間で、STEP2で仮定A, Cの間で矛盾が検出された場合の処理が示されている。特にSTEP2において矛盾の前提となっている仮定A, Cのうち確度の小さいのは仮定Cであるが、仮定Aをくつがえし、仮定Bを再び信じることにより、確度の和が最大になるように保たれている。

実行のSTEP	<STEP0>	<STEP1>	<STEP2>
矛盾の前提となっている仮定	-	A, B	A, C
C.B.S.	{(A, B, C)}	{(A, C), (B, C)}	{(B, C), (A)}
C.B.S.で確度の和が最大となる仮定の組み合わせ	(A, B, C)	(A, C)	(B, C)
くつがえす仮定/ 信じる仮定	-/-	B/-	A/B

各仮定の確度は A : 12, B : 11, C : 10

図2 仮定の制御の実行例

4. おわりに

仮定に基づく推論の故障診断における必要性、適用の方法、そのための仮定の制御の方法について述べた。この方法を用いることにより、不確かなことが疑わしくなったときに、それに関するテストを行わせるという処理が実現できる。現在、Prologを用いて前向きプロダクションシステムおよびその上の知識ベースを試作しているが、処理の効率などに課題がある。また、このような推論の応用として、ユーザの推測する内容を一つの仮定として事実間の関係の一貫性を保った形で受け入れ、利用する、などの可能性も考えられるが、それらについて今後検討していきたい。

なお本研究は第5世代コンピュータプロジェクトの一環として進めているものである。

【参考文献】

- [1] 吉岡他, "知識表現/推論システムPeaceの故障診断問題への適用", 情報処理学会第32回全国大会講演論文集4L-9 (昭和61年) 1121-1122
- [2] Stallman, R. M. and Sussman, G. J., "Forward reasoning and dependency-directed backtracking in a system for computer-aided circuit analysis", Artificial Intelligence 9 (1977) 135-196
- [3] Doyle, J., "A truth maintenance system", Artificial Intelligence 12 (1979) 231-272

オンラインKWIC検索システムにおける 分かち書き処理

97-4

岡 夏樹 重永信一 玉越靖司
(新世代コンピュータ技術開発機構) (松下電器産業株式会社)

1. はじめに

我々は、実文例中での言葉の使われ方を容易に収集することができ、自然言語処理研究に有効なオンラインKWIC (Key Word In Context) 検索システムの開発を行ってきた。

本システムは、与えられたキーワードの検索をオンラインで効率よく行うために、文中に現れる単語の索引ファイルを前もって作成しておく方法をとっている。このための単語抽出は、漢字かな混じりのべた書き文から自動的にパッチ処理で行う。ここでは、この単語抽出用の分かち書き処理について報告する。

2. これまでの研究との比較

これまでに報告されている分かち書きの処理法は、かな漢字変換に関連した、かな入力を対象としたものが多い。いずれも最長一致法を用いた形態素解析を基本にし、2文節最長一致法[1]、文節数最小法[2]、語の出現頻度に基づく方法[3]、助詞に着目する方法[4]等のようにそれぞれ精度を上げる試みがなされている。

これに対し、漢字かな混じり入力では、簡単な方法でもかなりの精度が得られると予想されるので単純な最長一致法を採用した。ただし漢字かな混じり入力の場合は表記のゆれに対処する必要がある。

文献[5]は、字種の変り目に着目して、小さな辞書だけで解析することを提案しているが、これは精度の向上に限界があり、例外の処理が増えてかえって複雑になる恐れがあるので、本システムでは大きな辞書を使用することにした。

従来、文節内の接続可能性と文節を構成する条件とを考慮して解析を進める方法がとられてきた。これに対し今回報告する方式の特徴は、それらを手続き上区別せずまとめて隣接可能性としてとらえ、一貫して辞書と隣接表により処理を進めることで、これにより非常に簡明なプログラムを実現した。

3. 処理の概要

本処理は漢字かな混じり文(かな書き文も可能)の単語分かち書きを行い、合わせて各単語の品詞も決定する。

単語辞書および隣接表を用いた形態素解析のレベルで正しいと判断される結果のうち、最初に見つかったものを出力する。

4. 処理アルゴリズム

以下のように縦型探索を行う。

- 1) 解析点からの入力文字列と一致する単語すべてを辞書から探して、単語候補リストをつくる。
ただし英数字等は字種切りにより単語候補とする。
- 2) 付属語・活用語を優先した最長一致法による優先順位に従い、隣接表を用いて直前の単語との隣接検定を行う。
隣接可であれば、次の解析点にすすむ。
隣接不可であれば、次の候補について隣接検定する。
すべての単語候補について隣接不可であれば、一つ前の解析点に戻る。
解析開始点または解析確定点まで戻ってもだめだったら、解析済みの部分の次の文字を未登録語とする。
- 3) 一つ解を見つけたら終了する。

5. 辞書および辞書びき

かな表記および標準的な漢字かな混じり表記を見出しとし、送りがなのゆれにも見出しを追加する方針で対処した。

辞書びきの処理時間の観点などから辞書は次の3種類に分けた。

- ・ひらがな1文字表記の自立語辞書(約300語)
 - ・上記以外の自立語辞書(約30000語(表記だけが異なる語は含まない))
 - ・付属語辞書(形式名詞・補助動詞などを含む)(約250語)
- 辞書びきは索引表・索引木を使って高速に行っている。

6. 隣接表および隣接検定

接続、係り受けなどの関係の有無によらず、2つの単語（記号類なども含む）が隣接しうるかどうか（隣接可能性）のみに着目し、これを1枚の表の形式にまとめたものを隣接表と呼ぶ。

隣接可能であるのは次の場合である。

- 1) 文節内接続可能（自立語-付属語間、付属語-付属語間）
- 2) 複合語可能（自立語-自立語間、接頭語、接尾語、助数詞など）
- 3) 前の単語が文節の末尾に、後ろの単語が文節の先頭になりうる
- 4) 記号類（句読点、括弧、単位記号、文の先頭、文の末尾など）との隣接

7. 性能評価

性能評価用入力文として下記の2冊の本からそれぞれ100文ずつを使用した。

- ・「日本—その姿と心—」<第2版>、新日本製鉄株式会社能力開発部、学生社、1984
- ・「英語構文の研究」、高梨健吉、英誠社、1975

精度（正しく分かち書きされ、かつ、正しく品詞が与えられた単語数/総単語数）は、95%であった。品詞を問題にしなければ97%である。

本プログラムはPascalでインプリメントした。処理速度は、DEC2060のCPUタイムで、1文（平均23単語）当たり3.8秒であった。このうち辞書びきに必要な時間が約90%を占める。

8. 正しく分かち書きできない原因と考えられる対策

下記の原因のうち1)と2)が主要なものである。なお、辞書中のノイズになる見出しを削除するという場当たり的な方法は、入力対象を限定した場合には有効であるが、本質的な解決にはならない。分類や隣接検定をより細かくする方針をとるべきである。また、本システムはある入力文に対してどのような分かち書きが行なわれるか試験できるコマンドを備え、ユーザのキーワード選択の助けとしている。

1) 未登録語の存在

辞書の拡充で対処すべきである。表記のゆれにも対処しなければならない。未登録語が存在してもある程度単語の区切りは予測できると思われる[6]が、本システムでは、その様な処理は行っていない。

2) 単語候補の優先順位

処理時間や簡明さを考慮して、最初に見つかった解を一つだけ出力することにしたので、単語候補の優先順位が問題となる。精度の向上のためには、隣接検定を多レベルにすることが簡単でかつ有効な方法であると思われる。最も簡単には、隣接可能（出現頻度が多い）、隣接可能（出現頻度が少ない）、隣接不可能の3レベルにすることが考えられる。KWICのキーワード抽出という用途からみれば、複数の候補があればそれらを全部出力してやるのもひとつの方法であろう。

3) 隣接検定が不十分

4) 単語認定基準があいまい

9. おわりに

本システムは改良すべき点はあるが、現状でKWICのキーワード抽出用として実用に耐えうるものと考ええる。

なお、本研究は第五世代コンピュータ・プロジェクトの一環として、ICOTからの委託により、松下電器在職中に行ったものである。

10. 参考文献

- [1] 牧野他，“べた書き文の分かち書きとかな漢字変換”，情報処理学会論文誌，vol.20, No.4, pp.337-345, 1979.
- [2] 古村他，“文節数最小法を用いたべた書き日本語文の形態素解析”，情報処理学会論文誌，vol.24, No.1, pp.40-46, 1983.
- [3] 内田他，“自由入力方式のかな漢字変換”，情報処理学会自然言語処理研究会資料，27-3, pp.1-8, 1981.
- [4] 館林他，“自由文入力・かな漢字変換方式”，情報処理学会第26回全国大会講演論文集，論文番号2H-1, pp.1165-1166, 1983.
- [5] 植村，“電子計算機による自動索引の研究（上）”，電子技術総合研究所報告第743号，1974.
- [6] 古村他，“未登録語を含む日本語文の形態素解析アルゴリズム”，九大工学技報，vol.55, No.6, pp.635-639, 1982.

数式を含む日本文の構文解析

4S-4

和田 孝, 西谷 泰昭, 岩元 寛二

日本電気(株)ソフトウェア生産技術研究所

1. はじめに

我々は、プログラム自動生成システムPGENを開発中である。PGENは、日本文、図、表で記述されたモジュール仕様書からプログラムを自動生成する実験システムである[1, 2]。モジュール仕様書の日本文には、条件や計算を表現する数式が多く含まれる。これらの数式は日本文との明確な区切りなしで混入しているのが普通であり、また数式の中には日本語の表現が含まれる。そのため、日本文の構文解析において、数式の範囲、構造のあいまいさが問題となる。PGENにおける構文解析では、数式を日本文と同じ枠組みで解析する方式をとっている。本稿では、その方式について報告する。

2. 構文解析の概要

PGENの構文解析の全体構成を図1に示す。構文解析は入力日本文を、その構造を表現したネットワーク(N-ネットと呼ぶ)に変換する過程である。

文節情報抽出では、分かち書き解析で認識した文節を、次の係り受け解析での単位となる文節の形式へ変換する。

係り受け解析では文法に従って文の構造を解析し構文木を生成する。日本文の構造のあいまいさは構文木の部分木を同一の位置に複数保持することにより表現する。[1]

最速木選択では複数の部分木から最速な部分木を選択することにより最速な構文木に取り込む。最速な部分木の選択は、各部分木に対して各文節の係りの長さの値を並べたベクトルを計算し、それらを比較することにより行なう。[1]

N-ネット生成では、最速な構文木をルートからたどりながらネットワークを生成する。

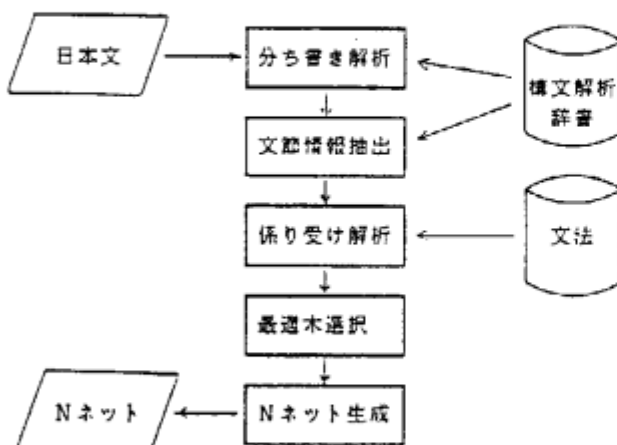


図1. 日本文構文解析部の全体構成

3. 数式の文法

日本文の中に混入を許す数式を図2のように定めた。二項演算子(BOP)と単項演算子(UOP)は辞書に登録することによりユーザが自由に定義できる。ユーザは各演算子について優先順位と結合性(左, 右, 否結合)を指定する。名詞句(NP-D, NP-F, NP-A)は日本文の文法に従う。名詞句は、中に数式を含んでもよい。

名詞句を終端記号とみなした場合、図2の文法はあいまいであるが、演算子の優先順位と結合性で解決される。

4. 数式を含んだ日本文の構文解析

プログラム自動生成システムの入力としては、ユーザに数式の構造を明示的に指定させないことが望ましい。その場合には、たとえば

(i) 「入力の通帳残高=元帳の通帳残高でなければ、…」

のように、数式に含まれる日本語のために数式の範囲があいまいになったり、

(ii) 「N+1番目の入出金明細」

のように構造のあいまいさ[(N+1)番目の入出金明細 か N+(1番目の入出金明細)か]が生じ得る。

このような日本語に起因するあいまいさを前述の最速木選択を利用して解決することにした。そこで、数式を日本文の場合と同じ枠組みで構文解析する方式をとった。

日本文は文節の並びであり、文節の間の係り受け関係を決定するのが日本文構文解析である。したがって、数式において文節に相当するものを決定し、その文節

```
<式> ::=
    <式> <BOP> <式>
  | <UOP> <式>
  | '(' <式> ')'
  | <NP-D>
  | <NP-F> '(' <式列> ')'
  | <NP-A> '(' <式列> ')'

BOP: 二項演算子
UOP: 単項演算子
NP-D: データを表現する名詞句
NP-F: 関数を表現する名詞句
NP-A: 配列を表現する名詞句
```

```
<式列> ::=
    (空)
  | <式>
  | <式列> ',' <式>
```

図2. 数式の文法

の間の係り受け関係を定める必要がある。

5. 数式における文節と係り受け

式は日本文における句に対応する。日本文の場合、たとえば「以下の処理を行なう」は2つの句、「以下の処理」と「行なう」が「を」という係り受け関係で結合してできた句である。二項演算子に関する式、たとえば

<式1> = <式2>

において、<式1>、<式2>は句であり、<式1>から<式2>へ係り受け関係があり、両者が結合して新しい式、すなわち句となると考える。二項演算子「=」は日本文における付属語（「を」等）に相当する。式における係り受け関係を図3に挙げる。

日本文の文節は係り受け関係の対象となる最小単位であり、句を構成する最小単位である（上述の例では「以下の」「処理を」「行なう」）。式においては、その係り受け関係から考えて文節に相当する最小単位は、

- (i) データを表現する名詞、または「(」と「)」で囲まれた式/式列
 - (ii) 二項演算子(BOP)が後ろに付いた(i)
 - (iii) 単項演算子
 - (iv) 関数/配列を表現する名詞
- である。また、名詞句(NP-D, NP-F, NP-A)の部分、日本文の文節の解釈に従う。

上記の文節の解釈では、1つの文節の中に「(」と「)」で囲まれた式/式列を含む場合、その文節の子として式の文節列が属することになる。すなわちカッコの構造が文節の入れ子構造として表現される。

文節の間の係り受け関係は日本文の場合と同じく、2句の間に係り受け関係があれば、両句のそれぞれ最後の文節の間に係り受け関係があるとする。

図4に式における文節とその間の係り受けの例を示す。

6. 構文解析方式

数式は日本文と同一の方式で構文解析する。以下、数式についてのみ述べる。

文節情報抽出では各入力について、数式の部分を前述の文節の列として解析する。数式における「(」と「)」の対応をとるのはこの時点であり、文節の入れ子構造が生成される。

係り受け解析では文法に従って式の構造を決定する。文法は2句の間の係り受け関係を記述したものである。式に関する文法は2句に含まれる演算子（正確には、左側の句の最後に現われる演算子、右側の句のトップレベルの演算子および最後に現われる二項演算子）に着目し、それらの優先順位と結合性の比較により記述

されている。

この文法規則により、たとえば、「A*」、「B+C」は優先順位が「*」>「+」であるので結合しない（係り受け関係がない）。また、「A+」、「B+C」は「+」の結合性が左であれば結合しない【(A+B)+Cが標準形である】。

最速木選択では、前述した数式に関連したあいまいさを除去する。選択は主に、演算子のオペランドのデータ型チェックにより行なわれる。演算子のオペランドおよび結果のデータ型はPGENの意味辞書に記述されている。「N+1番目の入出金明細」では「N+」が、「1」と「入出金明細」のいずれかへ係るので、両者への係りの良さをそれぞれ計算して比較する。「+」のオペランドが実数であると意味辞書に記述されていれば前者が選択される。最速木選択では部分式のデータ型が計算され、それが上位の式のデータ型チェックに利用される。

7. おわりに

現在、数式に関連したあいまいさの解決は演算子のデータ型チェックに多く依存している。PGENでは構文解析、意味解析の処理形態がバッテリであり、システムが最速と思うものを一意に決めている。この部分に金話機能を導入することを検討中であるが、これに伴い、数式のあいまいさの解決に、より多くのファクタを考慮してゆくつもりである。

なお、本研究は、ICOTからの委託（発注0706号）の一環として行なったものである。

参考文献

- [1] 西谷, 和田, 若元「知識利用のプログラム自動生成システム」情報処理学会第30回全国大会4K-10
- [2] 若元, 西谷, 和田, 庄司, 土田「PGEN-1: 仕様理解をベースとしたプログラム自動生成システム」情報処理学会第32回全国大会5M-6

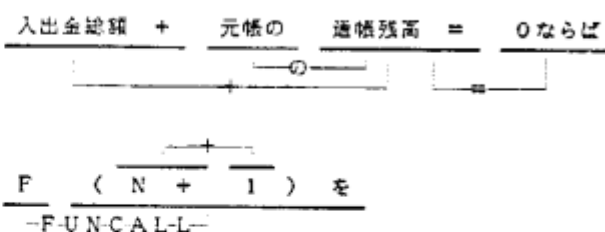


図4. 式における文節と係り受けの例

<式>	係り受け関係名	係り受け
<式1> <BOP> <式2>	<BOP>	<式1> → <式2>
<UOP> <式>	<UOP>	<UOP> → <式>
<NP-F> '(<式列>)'	FUNCALL	<NP-F> → '(<式列>)'
<NP-A> '(<式列>)'	ARRAYREF	<NP-A> → '(<式列>)'

図3. 式における係り受け関係

可能にするために導入したもので、その内容は図3に示すように、数値や図形コードで表現する形態情報、木構造の知識を指す構造情報、および、言葉で表現する機能情報からなっている。形態情報は、ハードかソフトかを区別したり、物の大小関係を判定するのに使え、構造情報は、新しいコンポーネントが言明されたときに木構造に加えるのに使い、機能情報は、計算するのか、記憶するのか、入出力をするのか等の区別をつけるのに用いる。機能を表現する用語はシステムで解釈が与えられている。またこの意味辞書には固有名詞の知識も記載されており、「〇〇〇の3倍の処理能力」といった場合にも具体的な値がとり出せるようになっている。

4. 要約処理

要約の出力となる表は、あらかじめ想定した構造知識の中から1つ選定する。この選択方法としては、見出し語の解析、記事の先頭部分の解析、記事中の専門用語の頻度からの推論等が考えられる。対象のルート概念が定まると、それに関する木構造知識をマクロ展開しながら作成して行く。これで、スコットを埋める準備ができたことになり、後は文章を解析しながら、与えられた構造知識との対応を取って行く。もし文章中の言葉（例えばプロセッサ）が木構造知識の中で複数箇所（例えばCPUのプロセッサとグラフィックのプロセッサが）存在すると、現在の文脈に近い方を選ぶ。逆にマクロ展開した木構造知識の中に一つも出現していない概念（例えばプ

ァイルサーバ）が文中に出現すると、その木構造への代入を断念するか、辞書項目を用いて、最も意味的に近い位置に新たにつけ加えるかのいずれかを選択する。文章が全て処理が終わると、表の中で記入された部分のみを出力する。

5. おわりに

要約システムで用いる世界知識の階層構造による表現と意味辞書について述べた。この方法では知識を要素概念で整理しており、変更・拡張に適した構造となっている。本システムは今後Prologを用いてインプリメントする予定である。特にポストエディットや知識ベースのデバッグ・更新を容易にするために、要約システムに適した知識エディタの開発も考えている。専門用語の数が増大して行くにつれて、処理速度の低下が予想されるため、高速化の工夫が必要となろう。

最後に本研究の機会を与えて頂いたICOTの横井俊夫第2研究室長、研究推進を支援して頂く総合システム研究所似島一彦部長、並びに、貴重なコメントを頂いた木下啓男、奥村晃の両室員に感謝する。

文献[1] 猪瀧、斎藤、堀：シナリオを用いる論文抄録理解・作成援助システム、情報処理学会論文誌、pp22-29 vol 24 NO 1, 1983

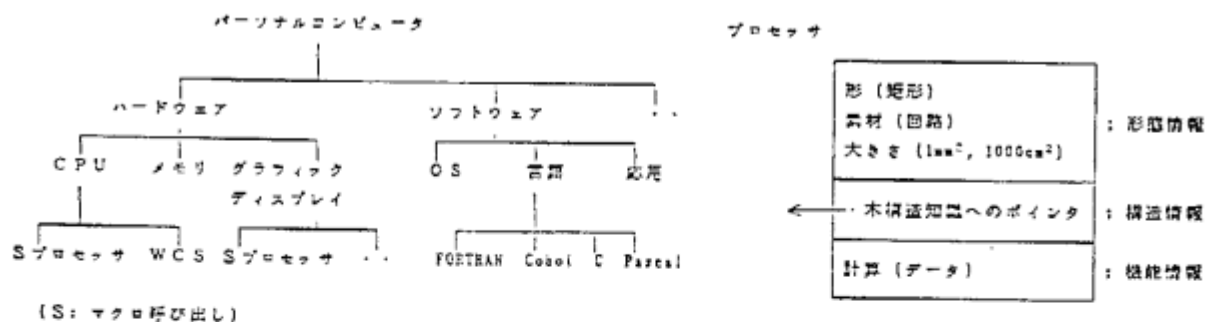


図1. 意味辞書の例

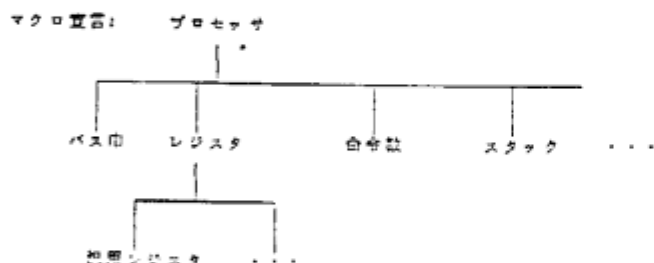


図2. 知識タイプの例

5W-6

PGEN-1:仕様書理解をベースとした プログラム自動生成システム

岩元 亮二, 西谷 素昭, 和田 孝, 庄司 功, 土田 賢省
日本電気(株) ソフトウェア生産技術研究所

1. はじめに

人間であるプログラマー向けに書かれたモジュール仕様書は、日本文、図、表の組み合わせで記述されることが多い。我々は、このような自然な形式で表現されたモジュール仕様書を理解しプログラムへ自動変換する方式について研究しており、その第一次プロトタイプシステムPGEN-0については既に報告済みである[1]。その後、図形の意味処理機能の追加、表の表現能力の強化、日本文のあいまいさ解消と省略語補完機能を強化し、文対応でなくモジュール仕様書単位で処理できる実験システムPGEN-1を開発した。本稿では、意味処理機能を中心にPGEN-1の特徴を述べる。

2. PGEN-1の概要

本システムの開発に当たって、対象分野を銀行オンライン業務プログラムとし、実際に書かれたモジュール仕様書を参考にして出来る限り制限を設けずに処理できるようにする方針をとった。

図1にPGEN-1のシステム構成を示す。システムは、図2、図3に示すような入力に対して、COBOL/S(NECの構造化COBOL)のソースコードを出力する。入力のモジュール仕様書(日本文)は、節・段落構造を持ち、節は節番号([.])で囲まれていると節タイトル(省略可)からなる節ヘッダで始まるものである。

日本文、表、図の各入力は、入力処理部で構造解析され、それぞれNネット、Tネット、Gネットと呼ばれるネットワークに変換される。これらのネットワークは入力に対応した構造を表現しており、冗長な情報、意味的なあいまいさ、情報の不足などをそのまま表現している。意味処理部は、これらのネットワークを参照して仕様書

の意味を表わす意味ネットワーク(Mネット)を作成し、更に適当な知識を用いてMネットをプログラム生成に十分な形に変形、補完していく。Mネットが正しく生成された後、プログラム生成部によりプログラムに変換される。

以上の処理では、種々の知識が利用されるが、これらの知識は、対象システムに固有なものとそうでないものに大別され、主に辞書に保持されている。対象システムに固有な知識は、主にデータに関連した情報であり、データの構造と属性、コード表、それらと日本語名との対応などの情報を持つ。これらの知識は、システムのデータ設計結果に対応し、PGEN-1ではデータ型辞書、変数辞書という形でユーザが定義するようになっている。一方、対象システムに固有でない知識として、日本文、表、図の文法的規則、あいまいさ解消や不足情報補完のための手続き的知識、プログラミング言語についての知識、主に用言(「代入」、「加える」など)に対応したプログラム上の意味知識などがある。これらの知識はあらかじめPGEN-1に保持されている。

- (1) 入力処理部
 - 全体が "00" でない入力部について、以下を繰り返す。
 - (1) 最初に入力を加える。
 - (2) 入力部をカウントアップする。
 - (2) 指定コードエリアのセット
 - 有値のとき、以下の処理を行う。
 - (1) 使用区分が1のとき、以下を行う。
 - (1) コード数をカウントアップする。
 - (2) コード数が以下になる。コード数番目のコードを "321" とする。
 - 0のとき、繰り越す。エラーコードを "321" とし、リターン。
 - (2) 指定コードエリアのセット
 - 有値のとき、以下の処理を行う。
 - (1) 使用区分が1のとき、以下を行う。
 - (1) コード数をカウントアップする。
 - (2) コード数が以下になる。コード数番目のコードを "321" とする。
 - 0のとき、繰り越す。エラーコードを "321" とし、リターン。
 - (2) 指定コードが "20" のとき、指定コードに "N0" をセットし、"20" のときは、"N1" をセット。
 - (3) 連続がフル指定範囲より大きいとき、表2に従い指定コードに値をセットしてリターンする。
 - (4) 変数辞書のとき、表1の読み合せ以外なら、エラーを409としてリターンする。
 - (5) 図1のとき、リターンする。
 - (6) 図1に従い、値をセット。

図2 入力日本文の例

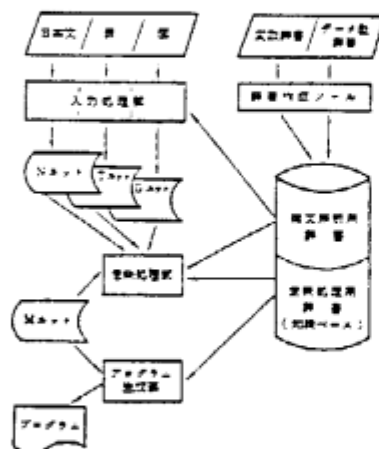


図1 PGEN-1のシステム構成

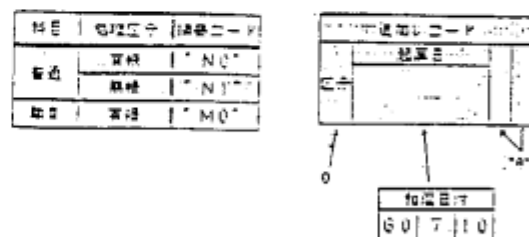


図3 表、図の入力画面

3. 日本文の意味処理

意味処理の入力であるNネットは、格文法の考え方にそったネットワークであり、構文上のあいまいさは構文解析の妥当性ベクトル[1]の考え方により既に除去されている。意味処理部では、冗長表現の削除、意味的あいまいさの解消、不足情報の補充、プログラム上の意味への変換を行なう。

(1) 冗長表現の削除

「Xの組み合わせ」という表現は、Xが表や図の参照、あるいは変数の並びを表わす表現のとき、単に「X」という表現で十分であり、「の組み合わせ」は冗長である。日本文では他に、「Xの値」、「X番目」、「Xの処理」などのように、Xの種類によっては単に「X」としてもよい表現がある。PGEN-1では、Xの種類を見極めた上でこれらの冗長表現を削除する。

(2) 並列句の解釈

簡潔な表現法として仕様書のなかには並列句が多用されている。例えば、「A, Bにa, bをセット」、「A, Bに1をセット」、「A, Bがa, bである。」などが典型的な例である。構文解析結果のNネットでは、並列句は図4(a)のように表現されている。この例は「A, Bにa, bをセット」に対応し、並列句「A, B」に並列句「a, b」をセットするということを表している。意味処理部ではこれを「Aにaをセットし、Bにbをセットする」という意味のMネット(図4(b))に変換する。また、「A, Bに1をセット」は「Aに1をセットし、Bに1をセットする」を表わす意味ネットに変換する。このように、並列句については、述語の種類とその格の内容の組み合わせによってどのように意味ネットに変換するかを手続的知識として組込んでいる。

(3) あいまいさの解消

日本文の単語に対して、プログラム仕様書ではいくつかの意味が対応する。例えば、「である」(be動詞)に対しては、等しい(eq)、含まれる(member)、代入する(set)、などの意味がある。また、データを示す名詞についても一つに変数を特定できない場合がある。

あいまいさ解消に仕様の論理構造を規定するMモデルを利用する。MモデルはMネットのひな型であり、Mネットのノード間の結合条件を規定している。例えば、Mモデルの一部(図5)により、setのobj-rとしてaddが許されることがわかる。このMモデルによって妥当な解釈の範囲を絞ることができる。この方法は特に日本文の条件部と実行部の区別が有効で、条件部にはpredicateが要求され、実行部にはstatementが要求されることからbe動詞のあいまいさを除くことが可能である。

データを表わす名詞のあいまいさ解消には、データの構造と属性に関する知識を利用するが、さらに節タイトルも利用する。例えば、図2の2節の中では、「コード数」を監査キーコードエリアのサブフィールドの「コード数」と同定し、3節の中では注意コードエリアのサブフィールドの「コード数」と同定する。

(4) 省略語の補充

データ型とコード表の知識を利用して、「普通預金」ときが「科目コードが普通預金(=21)に等しいとき」の意味に対応するMネットに変換されるのは先に発表した通りである[1]。このほかに、繰返し表現に伴



図4 並列句の表現

図5 Mモデルの一部

う省略語補充と局所的文脈による補充を行う。

繰返し表現では、図2の4節の例のように同型のパターンが繰返されるときに後のパターンの中で前のパターンと同一のものが省略される現象がみられる。この例では、後のパターンで「科目コードが」と「摘要コードに」が省略されている。PGEN-1では直前に類似のパターンがあるかどうかみてパターンによる補充を行う。

局所的文脈利用の補充の例は、「AとBを加え、Zに代入する」を「AとBの和をZに代入する」の意味への変換である。これは「代入」に対する補題obj-rを補充している。この場合、(3)で述べたMモデルに合致する対象を局所的文脈から同定している。

4. 図表の意味処理

図の意味処理では変数に関する知識を参照してデータ構造図としての意味を同定する。関係を表わす線分の両端の接する範囲については、両端でデータ属性が矛盾しない最小の範囲を選択する。関係の種類はeq(等号)、set(代入)のいずれかを決定する。このとき日本文の意味表現を参照し、条件部で図を参照していればeq、実行部で参照しているときはsetと判定する。

表の意味処理では、日本文において条件としてのみ表を参照しているときには(and, orでつなげた)条件の意味表現を、それ以外の場合は(if-then-elseの連鎖における)条件とアクションの意味表現を生成する。表のアクション部については、見出しと日本文の意味表現との間の関連を同定し、その結果により表の意味内容を日本文の意味表現のどこに埋めるべきかを決定する。

図と表の意味処理部は日本文の意味処理の途中のフェーズで起動され、その結果は日本文の意味表現と結合されてMネットを形成する。このMネットは日本文意味処理部の後段のフェーズによって繰返し文としての構造化、あいまいさの解消、省略語補充の処理を受ける。

5. おわりに

はじめに述べたように、PGEN-1はできるだけ現状の仕様書を理解できることを目的としているが、今後は確実に理解できる水準を定め実用的なものにすること、及び部品に対応する高水準の慣用表現の定義機能を追加することが重要であると考えられる。最後に本研究に対し励ましと御指導を下された日電・ソフトウェア生産技術研究所の森澤所長と藤野支配人に感謝する。尚、本研究は、(財)新世代コンピュータ技術開発機構からの委託(完仕3401-04号)の一環として行ったものである。文献[1] 西谷、和田、岩元「知識利用のプログラム自動生成システム」昭60情報学会全大第30回

並列論理型言語による

2U-2

LSIのフロアプランシステムの検討

— 実験システム概要 —

吉永和弘 青柳洋介 白水昇

(沖電気工業株式会社)

1. はじめに

VLSIの集積度の増大及び機能の多様化が日々進行しており、設計作業の効率化のために知識工学の応用が、又高度化のために並列処理の導入が求められている。そこで、第5世代コンピュータ計画での核言語の母体である並列論理型言語Guarded Horn Clauses (GHC) [1]を用い実験システムの試作中である。特に、従来の手法では、高品質・高速のフロアプランシステムは実現が困難であると考えられている。

本稿では、実験システム構成及び配置機構概略構成について述べる。

2. 実験システム構成

フロアプラン決定に関しては、設計者のノウハウに多くを依存しており実現が困難であるので、まず、既存の手法でもある程度自動化されている配置機構について試作した。また、既存CADシステム中の設計データの流用や実験システムの評価が必要であることからDBインタフェース機構を、大量の図形データを取り扱うためにグラフィクスシステムが不可欠であることからユーザインタフェース機構についても試作した。(図1参照)

2.1 配置機構

設計データを入力することにより、GHCの並列プロセスとして配置モデルを生成し、配置状況を表示しながら配置処理を実行し、配置結果を出力する。

2.2 DBインタフェース機構

CADデータベース中に格納されている大量の設計データをGHC上へ抽出し、逆にGHC上の配置結果をデータベース上に併合する。

GHC上から直接、データを入出力する直接方式と、データをGHC形式へ一旦変換し入出力する間接方式が考えられる。直接方式では、GHCとデータベース間で整合をとることが困難であるので間接方式でのみの試作を行った。

2.3 ユーザインタフェース機構

配置対象の状態を随時グラフィック端末に表示する。

GHC上に全体を実現する組こみ方式と、GHC上には図形処理コマンドの入出力インタフェースのみを実現する通信方式が考えられる。組こみ方式は、並列性を考慮した効果的インタフェースが可能となるが実現が困難であるため通信方式のみの試作を行った。

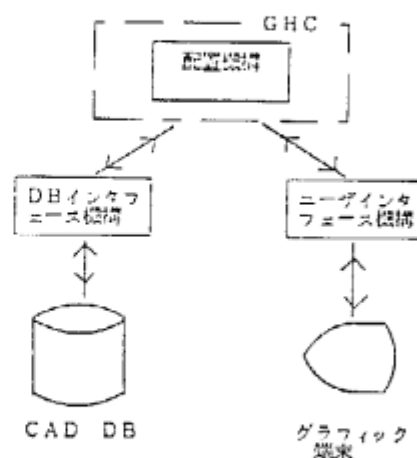


図1. 実験システム構成

3. 配置機構概要構成

GHCは、AND並列を並列プロセスの記述に、OR並列を非決定性プロセスの記述に用いることができる。また、プロセス間の共有変数を通信路として用いプロセスの同期を取ることができる。そのため、Concurrent Prologの場合と同様に、プロセスをtail-recursiveに走らせることでオブジェクト指向的なプログラミングが可能となる。〔12〕,〔13〕そこで、基本的にオブジェクト指向の考えをふまえ、並列性を導入することにより配置機構のモデル化を行った。(図2参照)

3.1 設計データ

設計データは、プロセスとして生成されると内部情報を出した後に消滅する、状態を持たず事実のみを表すオブジェクトとしてモデル化した。

データベース中には、各種ブロック(配置対象、ライブラリ等)がメンバ(ブロック毎の属性で、例えば外形や内部端子等)毎に格納されている。そこで、記述の単位をメンバ毎にし、形式を

db ([ブロック名, メンバ名],
出力) :-

true | 出力 = 内部情報。

に統一しプロセスの生成を単純化した。

3.2 プロセスジェネレータ

プロセスジェネレータは、設計データから受けとった設計情報を適宜入力することで、配置モデル内の全てのプロセスを具体化し、最後に消滅するオブジェクトとしてモデル化した。

具体化は、プロセス毎に、プロセス名、属性(当該プロセスの固定的情報)、通信チャネル、初期状態を設定することにより行われる。

2.3 配置モデル

配置モデルは、ブロック(素子、端子)及びエリアから構成される並列プロセス群が相互に通信を繰り返しながら初期状態から最終状態へ変移し消滅するオブジェクトとしてモデル化した。

なお配置モデルの詳細については、別稿〔4〕で説明する。

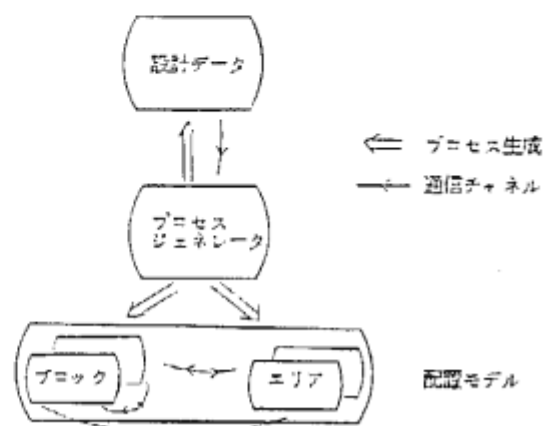


図2. 配置機構構成

4. おわりに

実験システムのプロトタイプが構築できた。しかし、配置モデルの変更に伴いプロセスジェネレータやDBインタフェース機構を変更しなければならない。そこで現在は、これ等に独立性をもたせ効率よく実験できるように検討している。

なお、本研究は第5世代コンピュータプロジェクトの一環として行われた。御討論をいただいたICOT第1研究室吉川室長に感謝する。

〈参考文献〉

- [1] 核言語第1版説明資料, ICOT
- [2] How to solve it in Concurrent Prolog, ICOT
- [3] 野田他, 並列処理方式論理シミュレータの一検討, 第30回情報全大
- [4] 吉永他, 並列論理型言語によるLSIのフロアプランシステムの検討ーデータ表現手法ー, 第32回情報全大

6U-2

並列論理型言語による LSIのフロアプランシステムの検討 - データ表現手法 -

吉永和弘 青柳洋介 白木昇
(沖電気工業株式会社)

1. はじめに

LSIのフロアプランシステムへの並列性の導入に関し、今回並列論理型言語GHCを用い配置機構について試作、実験を行った。システムをオブジェクト指向的考え方でモデル化し、その構成要素である設計データ、ブロック、エリアなどをGHCの並列プロセスで構成し、その表現方法や通信方法について検討を行った。

本論文では配置機構のモデル化、各プロセスの機能、問題点について報告する。

2. 配置機構

配置機構は、初期配置部と配置改善部に分割して考えた。今回は初期配置部について検討を行い試作した。また配置対象はポリセル方式のLSIの最下位層とした。以下、配置対象、初期配置部のモデル化について述べる。

2. 1 配置対象

配置対象は、ブロック（端子、素子）、内部に含まれるブロックの種類、大きさにより算出された数段のエリアより構成される。

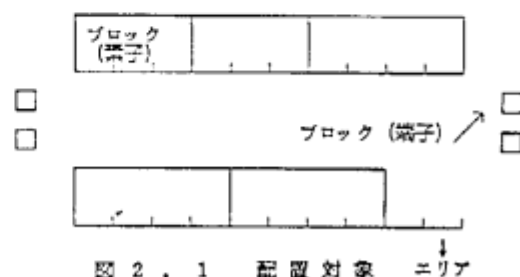


図 2. 1 配置対象 エリア

2. 2 配置機構のモデル化

入力端子側から信号の流れに沿ってセルが配置されてゆく様なモデルを検討した。以下、ブロックモデル、エリアモデルについて述べる。

(1) ブロックモデル

ブロックは隣接ブロックと同期を取るための同期プロセス、移動規則を持ちエリアと通信を行いながらブロックを移動させる移動規則プロセス、及びこれらのプロセスを制御するブロックマネージャから構成される。

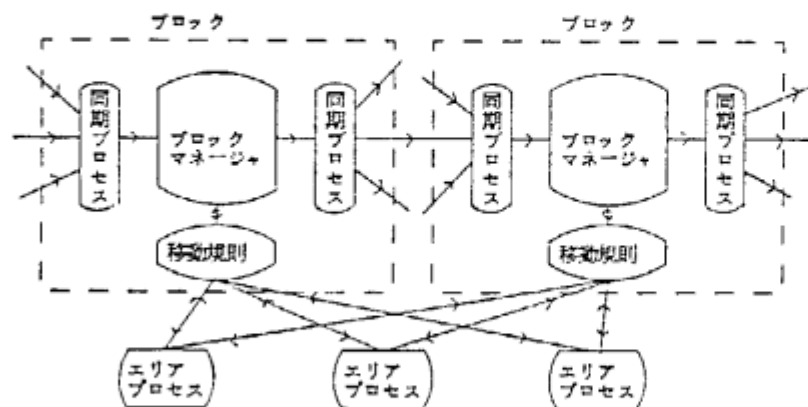


図 2. 2

① ブロックは入力側の隣接ブロックの移動終了で同期をとり、移動を開始する。移動開始の条件は隣接移動済ブロックとの結合密度の合計値が一定値に達した場合か、隣接ブロックのすべてが移動終了となった場合である。

② エリアと通信を行いながら配置可能な空エリアのうち、

$Cost = \sum (\text{距離} \times \text{結合密度})$ を最小にするエリアを占有する。そのとき、他のブロックから参照されているエリアがあれば、そのエリアが解放されるか、配置済になるまで待つ。

③ 移動が終了したら、出力側隣接ブロックにメッセージを送り消滅する。

```
block(BLKXNM, St, In, Out, Ar_chs) :-
    St=[state(BLKXNM, trans, Area, Pos) |Ns],
    block_move_rule(Ar_chs, Pos),
    block_area_sansho(BLKXNM, Pos,
        Ar_chs, Ar_chs_t, Res),
    block_area_haichi(BLKXNM, Res,
        Ar_chs_t, Ns),
    Ns=[state(BLKXNM, State, Area,
        (X1, Y1)) |Ns1],
    block_st_change(Res, State),
    block_trans_st(Res, Out, Out1),
    block(BLKXNM, Ns, In, Out1, Ar_chs).
```

ブロックの記述例

(2) エリアモデル

ポリセル方式のLSIではブロックはその種類により異なる幅を持つ(2～数十グリッド)。そのため、エリアを格子状に規格化された網目(メッシュ)に分割して表現した。各メッシュごとに1つのプロセスが存在する。

① エリアはブロックから参照メッセージを受け取ると自分の状態(配置済、参照中、未配置など)を返す。

② ブロックから配置要求メッセージを受け取ると配置可の場合自分の状態を配置済の状態に変更する。

```
area(GRPNM-Mno, Ar_ch) :-
    Ar_ch=[area(GRPNM, Mno, Type, (Y, X),
        State, Mode, Ack) |Ar_ch1],
    Mode=access(BLKXNM, haichi),
    State=state(BLKXNM, trans) |
    State1=state(BLKXNM, done),
    Ack=ack(GRPNM-Mno, State1),
    Ar_ch1=[area(GRPNM, Mno, Type, (Y, X),
        State1, Model, Ack1) |Ar_ch2],
    area(GRPNM-Mno, Ar_ch1).
```

エリアの記述例

2. 実験結果

今回は図3.1の回路を対象に実験を行った。結果は図3.2の様になった。

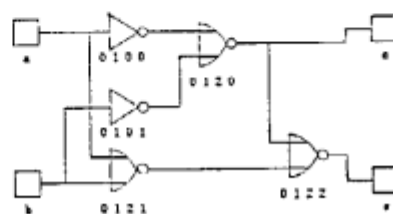


図3.1

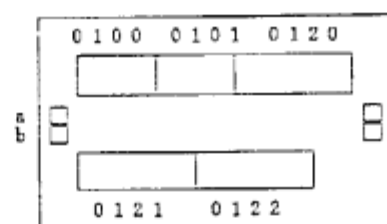


図3.2
配置結果

この実験によりブロックの配置を並列プロセスとして実行可能ことが確認できた。結果を人手による配置と比較すると最適とは言えない。これはブロックが局所的な情報しか知ることができないことと、移動規則が十分でないためであると考えられる。今後の配置改善部には、全体の監視を行うようなプロセスや移動規則の改良が必要である。また、より大規模な実データで実験を行い既存の配線システムにかけなどの実用評価をする予定である。

なお、本研究は第5世代コンピュータプロジェクトの一環として行われたものである。御討論をいただいたICOT第一研究室古川室長に感謝する。

5Q-8

Multi-PSIシステムの概要

瀧 和男 木村 康則 横田 実 近山 隆 内田 俊一
((財) 新世代コンピュータ技術開発機構)

1. はじめに

日本の第五世代コンピュータプロジェクトでは、論理型プログラムを並列に高速実行する並列推論マシンの研究を主にアーキテクチャの観点から行なってきた[1]が、並列ソフトウェアの面では、多くの問題が未だ未着手のまま残されている。そのおもなものには、

- (1) プロセッシング・エレメント間の負荷分散をソフトウェアで制御する方式
- (2) 分散環境におけるメモリ管理やガベージコレクション方式
- (3) 大規模な並列プログラムをいかにして書くか。
- (4) オブジェクトコードの分配、管理方式
- (5) 並列推論の世界での入出力や制込みの取扱い
- (6) 分散環境におけるプログラムのテストとデバッグの方式

等があり、アーキテクチャの研究を加速するためにも、並列推論に関するソフトウェア上の研究を進展させる事が急務になっている。

そのためのアプローチとして、われわれは研究開発ツール用に別途開発を進めてきたパーソナル逐次型推論マシンPSI[2]を活用することとし、これを複数台接続して実際に並列計算機環境を作り、その上に並列推論言語の実行環境を構築することにした。これをMulti-PSIシステムと呼ぶ。その目的は、並列推論マシン用の、

- (1) 言語 KL-1 (Kernel Language version 1)
- (2) オペレーティングシステム PIMOS (Parallel Inference Machine's Operating System)
- (3) 応用ソフトウェア

の研究開発を進めることであるが、開発にあたっては、ある程度大規模な並列ソフトウェアを書き、デバッグし、走らせるに足るだけの信頼性、使い易さ、性能を備えた並列推論言語の実行およびデバッグ環境を構築するよう、特に注意を払っている。

以下では、Multi-PSIシステムとその上で行なおうとしている研究開発の全体像について概要を述べる。

2. 研究開発スケジュール

次の3つのシステムを予定している。

- (1) Pseudo-multi-PSI (61年度初版完) : 1台のPSI上に作られるMulti-PSIのシミュレータで、ソフトウェアのデバッグに用いるほか、負荷分散戦略の調整に

必要なデータの収集も行なう。処理系は(2)に準ずる。

- (2) Multi-PSI(-V1) (61年度完成) : 現在のPSIを入力装置も含めて6台程度接続したシステムで、ESPベースのKL1処理系を実装する。PIMOSの各部を試作し、簡単な評価用プログラムを実行できるようにする。分散環境におけるデバッグ方式を検討する。

- (3) Multi-PSI-V2 (仮称) (63年度完成) : 改良小形化されたPSIのCPU部分を数十台接続し、その上にファームウェアで記述したKL1処理系を実装する。フロントエンドマシンとしてPSIを1台設置する。PIMOSの改良、拡張を行ない、本格的な応用プログラムを実行させ評価を行なう。

3. Multi-PSI (V1) システムの概要

PSIマシン6台程度を接続ネットワークハードウェアにより格子状に接続した構成をとる。共有メモリは置かず、メッセージパケットの交換によりCPU間通信を行なうが、メッセージはユーザプログラムが属に取扱うのではなく、言語の実行メカニズムの一部として自動的に生成される。実装予定の言語KL1は並列推論言語GHC[3]をベースに置き、AND並列、AND結合されたゴール間のストリームによるデータ通信、ガード部による同期機能などを特徴とする。他のCPUへのゴール送出や他CPU内のデータとのユニフィケーションは、メッセージパケットの形に変換されて送出される。CPU間の負荷分配がゴール1個の単位で行なわれ得ることから、かなりfine-grainな並列性を取扱うシステムといえる。

Multi-PSI(V1)システムはPSIマシンをベースとして、図1の階層構成図中の下線部分を新たに製作する。主な構成部分について以下に述べる。

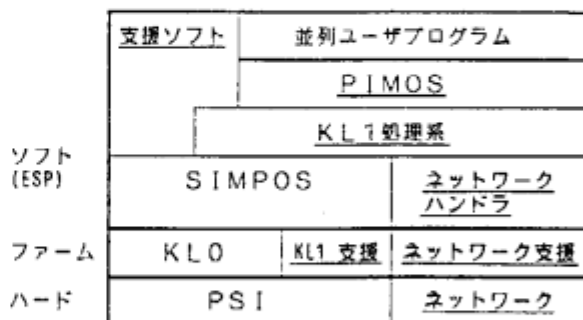


図1. Multi-PSI (V1) システムの階層構成図

4. 言語と処理系

前述のGHCをベースとした言語にモジュール化機能を取入れたユーザ言語をKL1-u(user)と呼び、PIHOSおよび評価用プログラムの記述にはこれを用いる。一方、言語機能を更に分解して機械語相当にしたものをKL1-B(base)と呼び、マシン上の処理系はKL1-B用に作られる。[4]分散環境用のユニフィケーションや負荷分散制御のサポート機能はこの中で実現される。これらの他にユーザが負荷分散制御や実行優先レベルを指定するためのnotationでプログラムの意味とは独立に付加できるものをpragma(KL1-P)と呼んでいる。KL1-uで書いたプログラムにpragmaを付加したものが、コンパイラによりKL1-Bに変換され実行される。

V1システムでは処理系はPSIのシステム記述言語ESPで記述の予定であり、効率向上の為に必要に応じてファームウェアサポートを行なう。

5. PIHOS

並列プログラムを効率良く実行するための基本機能として1.であげた(1)(2)(4)(5)の項目を実現する。中でも重要な試みは、距離の存在するネットワークすなわち通信のローカリティの概念を持ち込む必要のあるネットワーク上で、ソフトウェアにより負荷分散と通信の局所化を制御することである。これらは言語処理系とも深い関連を持つ。

PIHOSの初版では、並列推論マシン向きの実行管理と資源管理が主な研究テーマであり、ユーザに対するサービス機能はPIHOSの改良の段階で検討していく。

6. 支援ソフトウェア

V1システム用の支援ソフトウェアは、PSIのプログラミング・オペレーティングシステムSIHPOSを活用してESPで記述し、各PSI上に同じものを配置する。

(1) Multi-PSIシステムモニタ: KL1処理系と接続ネットワークの初期化、起動、停止、モニタリング、ハードウェア保守サポートを行なう。

(2) 分散環境デバッガ: 各プロセッサに割付けられたゴールや、プロセッサ間の通信に着目した低レベルのデバッガである。

7. 接続ネットワークハードウェア

主に実装設計の容易さから格子状のネットワークを選択した。格子状ネットワークはCPU間距離の概念が存在するネットワークの代表的なものであり、PIHOSにおける負荷分散制御の実験に適するが、V1システムではCPUの推論速度に比べてネットワークの転送速度が遅いため、異なるタイプのネットワークのシミュレーションにも使用できよう。

ハードウェアは、CPUの内部バスのオプションスロットに実装され、隣接の4台のCPUとそれぞれ接続する為の4

本のケーブルが引出される。これらの引出し口をそれぞれチャンネルと呼んでいる。チャンネルには送信用、受信用に各々独立した信号線の組が含まれる。データはパリティビットを含め10ビット単位で並列転送され、チャンネルの一方向当たりの転送能力は約500kB/secである。

データは可変長パケットの形で転送され、先頭には行先CPU番号の情報を持つ。接続ネットワークは行先CPU番号を認識し、それが自CPUであればデータパケットを取込む。異なる場合には立ちあげ時にプリセットされたCPU番号と送出先チャンネルとの対応表を引いて再送出する。各々のチャンネルから到着したパケットがそれぞれ異なる送出先へ再送され、行き先チャンネルがぶつかり合わない場合は、転送は同時並列的に実行される。この方式では、横方向優先などの簡単なルーティング戦略を固定的に用いるだけで、デッドロックが回避される。ネットワークとCPUのインタフェース部分には、4k語の容量をもつ送信用、受信用のFIFOバッファを設置しており、また各チャンネルの送出部分にも256語のFIFOバッファを置いた。

8. 負荷分散の方式

プロセッサ台数が非常に多い場合にも適用できること、プログラムからプロセッサ台数を意識する必要が少ないことを重視しつつ、通信の局所性と負荷バランスの制御をプログラマに委ねる立場をとった方式を検討中である。

プログラマには計算パワーが一様に分布した単位平面を想定させてそれを分割しては部分問題に与えることを書かせ、計算パワー平面上で部分問題間の通信の局所性を保つように分割と部分問題の配置を制御させる。この制御はプログラムにpragmaを付加して行なう。一方システムは、計算パワー平面と物理プロセッサとのマッピングを管理し、部分問題をどの物理プロセッサで実行するかを決める。長い間負荷の不均衡が続く場合には、隣接プロセッサとの通信によりマッピング関係を調整し、プログラマが制御しきれなかった負荷バランスを補正する。

9. おわりに

Multi-PSIシステムの概要について報告した。現在V1システム用の接続ネットワークハードウェアを製作中であり、ESPベースの処理系も製作が進んでいる。今後システムの構成要素の各論につき順次報告する予定である。

文献 [1] Murakami, K et al.: Research on Parallel Machine

Architecture for F.G.C.S., Computer, vol.18, no.6, June '85

[2] 西川他: PSIの性能評価(1), 第30回情報全国大会, no.1C-2

[3] Ueda, K: GUARDED HORN CLAUSES, Proc. on the Logic Programming Conf. '85, July 1-3, in Tokyo, pp.225

[4] 江崎他: GHCサブセット逐次処理系の作成, 本大会, 2G-4

高性能Prologマシン(CHI)の機械語命令先取り方式についての考察

3Q-4

梶田 伸一 中崎 良成 梅村 澄

(日本電気(株) C&Cシステム研究所)

はじめに 通産省第5世代計算機プロジェクトの一端として開発中の高性能Prologマシン"CHI" [1] は、Prologで記述したプログラムの高速実行環境の提供を目的としたProlog専用マシンである。Prologプログラムの高速処理を実現するために、コンパイラによる最適化 [2]、[3] とアドレス演算・データ演算・命令解釈の並列処理 [4] を採用している。この結果、大部分の機械語命令は、実行に必要なステップ数が2から4ステップとなり、機械語命令先取りが性能に与える効果は大きくなっている。本稿では、CHIの研究開発の過程で行なった機械語命令先取り方式の性能評価について報告する。

評価項目 CHIにおける機械語命令先取り方式の効果を評価するために、キャッシュメモリの構成とパイプラインのステージ数が異なる5種類のハードウェア構成の性能比較を行なった。パイプライン処理の評価では、命令のアドレス演算、命令の読み出し、OPコードの解釈とオペランドの切り出し、オペランド・アドレスの演算、命令の実行をサブ処理とし、1個のマイクロプログラム制御系が全体を制御する集中制御方式を対象とする。

キャッシュメモリの構成 キャッシュメモリを1個とし、データと命令共用にする場合と、データ用と命令用の2個に分割する場合を比較する。

データと命令共用とする場合、ハードウェア規模を小さくでき、制御も簡単である。しかし、命令の先読みによるキャッシュ・アクセスと命令実行によるキャッシュ・アクセスが同時に発生することによる処理性能の低下が予想される。

キャッシュメモリの構成を比較する理由は、キャッシュメモリを2個にすることで得られる性能が、ハードウェア規模の増加、制御の複雑化に見合うかを、評価することにある。

パイプライン・ステージの段数 CHIの機械語命令処理は、命令のアドレス演算、命令の読み出し、OPコードの解釈、オペランドの切り出し、オペランド・アドレスの演算、命令の実行からなる。機械語命令の先取りによる実行ステップ数の減少と、パイプラインの各ステージで行なう処理の簡単化によるクロック幅の減少が高速化に与える影響は大きい。そこで、次の5個の処理をパイプラインのサブ処理の候補とする。

- 1) 命令のアドレス演算
- 2) 命令の読み出し
- 3) OPコードの解釈とオペランドの切り出し
- 4) オペランド・アドレスの演算
- 5) 命令の実行

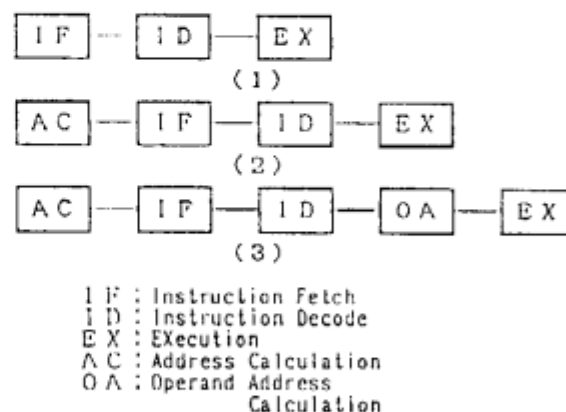
1) から4) のサブ処理は、1ステップで終了する。5) の命令の実行は、命令により必要なステップ数が異なり、2から4ステップ必要である。

ハードウェア規模、性能の2点から、第1図に示す3種類のパイプライン構成を比較する。(1) のパイプライン構成は、IF、ID、EXの3個のステージからなる。IFは、サブ処理2) を実行する。IDは、サブ処理3) を実行する。EXは、サブ処理1)、4)、5) を実行する。

(2) のパイプライン構成は、AC、IF、ID、EXの4個のステージからなる。ACは、サブ処理1) を実行する。EXは、サブ処理4)、5) を実行する。IFとIDは、(1) のパイプライン構成の場合と同じである。

(3) のパイプライン構成は、AC、IF、ID、OA、EXの5個のステージからなる。OAは、サブ処理4) を実行する。EXは、サブ処理5) を実行する。AC、IFとIDは、(2) のパイプライン構成と同じである。

EXステージの処理が少ない(3) のパイプライン構成が、最も性能が良いと予想できる。しかし、ハードウェア規模、制御の単純さは、(1) のパイプライン構成が優れている。パイプライン・ステージの段数を比較する理由は、段数を増加することで得られる効果とそのために支払うハードウェア量、制御の複雑さの関係を評価することにある。



第1図 パイプラインの構成

命令先取り方式 キャッシュメモリの構成とパイプラインの構成の組み合わせから、次の5種のハードウェア構成が考えられる。

- a) キャッシュメモリ 1個、(1)のパイプラインを使用した方式
- b) キャッシュメモリ 1個、(2)のパイプラインを使用した方式
- c) キャッシュメモリ 1個、(3)のパイプラインを使用した方式
- d) キャッシュメモリ 2個、(2)のパイプラインを使用した方式
- e) キャッシュメモリ 2個、(3)のパイプラインを使用した方式

上記5種のハードウェア構成の中で、ハードウェア構成e)が、性能的に、最も優れているのは、明白である。しかし、そのために投入するハードウェア量、制御の複雑さの量を考慮すると、性能の差を明白にする必要がある。

各方式の比較 C H Iにおける代表的なユニフィケーション処理用機械語命令の実行ステップ数を比較した結果を第1表に示す。○印は、他の構成より優れていることを示している。△印は、逆に、他の構成より劣っていることを示している。nは、ユニファイ処理の対象である2個のデータのタイプにより、0または、1以上の値となる。

(a) (b) (c)の結果から、キャッシュメモリが1個の場合のパイプライン構成と性能の関係が分かる。まず、(a)だけ性能が悪い機械語命令が12個存在することから、(a)の構成は好ましくない。(a)が(b) (c)に比べて劣る理由は、(b) (c)が、キャッシュメモリ内で命令先取りアドレスの演算を行なうのに対して、(a)では、プロセッサ部の演算器で行なうからである。したがって、命令先取りアドレス演算機能をキャッシュメモリ内に備える必要がある。

(b)と(c)の比較より、オペランド・アドレス先行演算による改善効果が小さいことが分かる。したがって、キャッシュメモリ1個の場合、オペランド・アドレス先行演算を導入する必要はない。同様に、(d) (e)の結果から、キャッシュメモリ2個の場合のパイプライン構成と性能の関係が分かる。キャッシュメモリ1個の場合より、オペランド・アドレス先行演算の効果が大きいことが分かる。これは、オペランド・アドレス先行演算によりキャッシュメモリの未使用ステップが減少したことが原因である。

キャッシュメモリの個数による影響は、(b)と(d)、または、(c)と(e)を比較することで分かる。(b)と(d)の結果から、オペランド・アドレス先行演算を行なわない場合は、キャッシュメモリの未使用ステップが多いため、性能が極端

に改善されてはいないことが分かる。しかし、オペランド・アドレス先行演算を行なう(c)と(e)の比較では、個々の性能改善効果が加わった結果、性能の改善効果が大きくなっている。

以上から、オペランド・アドレスの先行演算と2個のキャッシュメモリは、単独では、大きな効果が期待できず、方式(e)の様に、オペランド・アドレスの先行演算と2個のキャッシュメモリを組み合わせる必要がある。したがって、ハードウェア構成としては、ハードウェア構成が比較的簡単な(b)の構成が、命令の先行処理が比較的多い(c)が適している。

まとめ 各命令先取り方式について、C H Iにおける機械語命令単位での実行ステップ数の比較を行なった。この結果、中途半端な命令先取りは、あまり効果がないことが分かった。今後、さらに高性能なマシンを実現するため、ベンチマーク・プログラムによる性能評価を行なっていく予定である。

参考文献

- [1] Nakazaki, et al.: "Design of a High-Speed Prolog Machine" Computer Architecture, 1985
- [2] 中崎 他: 高性能PROLOGマシン: HPM - 基本アーキテクチャと処理法 - 情報処理学会第30回全国大会
- [3] Warren, D.H.: "An Abstract Prolog Instruction Set" Tech. report 1309, AI Center, SRI International, 1983
- [4] 堀田 他: 高性能PROLOGマシン: HPM - ハードウェア構成 - 情報処理学会第30回全国大会

	a	b	c	d	e	
mul temporary	2	2	2	2	2	
mul permanent	2	2	2	2	①	
mul temporary	△	1	1	1	1	
mul permanent	3	3	⑤	3	⑤	
mul list	△	1	1	1	1	
mul structure	△	2	2	2	2	
add permanent	2	2	2	2	①	
add temporary	△	1	1	1	1	
add permanent	3	3	⑤	3	⑤	
add temporary	△	1	1	1	1	
add list	△	1	1	1	1	
add structure	△	3	3	3	3	
unify permanent	3	3	3	3	3	read mode
unify temporary	3	3	3	②	②	write mode
unify permanent	2	2	2	③	③	read mode
unify permanent	△	3	3	3	3	write mode
unify temporary	△	1	1	1	1	read mode
unify value	3	3	3	②	②	write mode
unify list	△	1	1	1	1	read mode
unify structure	△	2	2	2	2	write mode
	△	3	3	3	3	listの場合
	△	3	3	3	3	structureの場合

第1表 実行ステップ数の比較

2G-6

高性能PROLOGマシン(CHI)
における最適化技法小長谷 明彦 中崎 良成 福田 伸一 梅村 護
日本電気株式会社 C&Cシステム研究所

1. はじめに

近年、新しいPrologのコンパイル方式としてWarrenの命令セット[1]が注目を集めている。本稿では同方式をベースとした高性能PrologマシンCHI[2]における最適化技法をそのコンパイル特性と合わせて報告する。

2. Warrenの命令セット

Warrenの命令セットは仮想的なPrologマシンの命令を定義したものであり、ユニフィケーション命令(`put`, `get`, `unify`)、ゴール呼び出し命令(`call`, `execute`, ...)、インデクシング命令(`try`, `retry`, `trush`, ...)より構成される。CHIの命令セットもこれに基づいている。Prologのクローズはこれらの命令の列に展開され、以下の最適化が可能である。

- (1) インデクシング
- (2) 終端再帰呼び出しの最適化
- (3) レジスタに割り当てられた変数のアクセス命令の最適化

3. インデクシング

インデクシングはクローズの特定引数に着目してクローズをあらかじめ分類しておき、冗長なバックトラッキングを防止する最適化手法である。この分類基準としてCHIでは現在以下のものを用いている。

- (1) データ型(変数、リスト、構造体、アトム・データ)による分類
- (2) アトム・データの値による分類
- (3) 構造体のファンクタの値による分類

データ型による分類に関する調査結果を表1、表2に示す。調査対象はCHI用に作成されたプログラム(資源管理プログラム、インタプリタ、コンパイラ等)である。平均的な手続きのクローズ数は2.7であった。表1より、アトム、構造体を扱うクローズは平均的に長く、変数とリストを扱うクローズは短い。これより、上記の分類は概ね適切であることがわかる。ただし、表2をみると、全てのインデクシング引数に変数であるクローズが非常に多い。この傾向はベクトル等リテラルとして現れないデータ構造を多用する資源管理プログラムにおいて特に顕著となる。このような手続きをよく調べると、

$$p(X) :- var(X), 1, \dots$$

$$p(X) :- integer(X), 1, \dots$$

のように定数で受けて、タイプチェックを行なっている例が非常に多い。したがって、変数の型を含めた分類を行なうとより効果的なインデクシングが実現できると予想される。

表3はアトム・データのインデクシングに関して、インデックス・テーブルを逐次的に検索する場合と、ハッシュして検索する場合の実行時間を比較したものである(VAX780上でCで記述されたCHI機械語エミュレータを使用)。表より明らかなように、ハッシュ検索と逐次検索の差は少なくデータベースの先頭の方のクローズの検索はかえって逐次検索の方が早い。これはハッシュ値の計算や衝突の処理にオーバーヘッドがかかったためと考えられる。クロスポイントはデータにもよるが10~20であり、CHIのようにアトム番号をハッシュ値とせず、アトム名から計算したハッシュ値をアトム番号とは別に持つようなシステムではクローズ数が少ない場合は逐次検索を用いた方が効果的である。また、コンパイル時間は逐次検索の方がハッシュ検索よりも2~2.5倍早い。

4. 終端再帰呼び出しの最適化

Prologの最後のゴールの呼び出しは制御を戻す必要がないためリブルーチン呼び出し(`call`)ではなく、ジャンプ命令(`execute`)で実現することができる。このためには、ゴール引数の回避が必要であり、Warren命令セットではゴール呼び出しの前にゴール引数を引数レジスタにセットし(`put`命令)、呼ばれた側では引数レジスタの値とユニフィケーションする(`get`命令)ことによりこれを実現している。引数レジスタを用いることにより以下の利点が生じる。(1) ユニットクローズ、本体部にゴールが1つしかないクローズ(以下トランジェント・クローズと呼ぶ)では変数領域をスタック上に確保する必要がない。(2) 引数レジスタを用いたアクセスの最適化(次節)ができる。このうち、性能的には前者の効果が大きい。実際、Warrenの命令セットを用いると他のベンチマークプログラムに比べ、`append`プログラムが非常に速い[3]が、これは、引数型によるインデクシングの効果とともにトランジェント・クローズであるため制御情報を格納するためのメモリアクセスが全く不要となる効果が大きい。

表4はCHIインタプリタのクローズのタイプを

調べたものである。カット等の組込述語を通常のゴールと同じように扱うと終端再帰呼び出しの最適化の効果はあまり期待できないが、組込述語をインライン展開するとトランジェント・クローズが非常に多くなることがわかる。組込述語のインライン展開は変数分類を複雑にするが[3]、Prologの場合は単なるインライン展開の効果以上の高速化が期待できる。

5. 一時変数の最適化

鎖領域をローカルスタック上にとらずに済む変数は一時変数と呼ばれる。この鎖領域を引数レジスタ上にとるとユニフィケーション命令が省略できる場合があり、これは一時変数の最適化と呼ばれる。

単純な例としては、

$p(X, Y) :- q(X, a).$

のXのようにヘッドゴールと同一位置にある変数やヘッドゴールに現われるボイド変数、

$p(, a).$

があるが、レジスタのライフタイムを調べることで、より、より高度な最適化が可能である。例えば、

$append([X|Y], Z, [X|YZ]) :- append(Y, Z, YZ).$

では、Zに対するユニフィケーション命令(get命令、put命令)だけでなく、本体部のY、YZの引数レジスタへのセット命令(put命令)も省略できる[2]。

一時変数は任意のゴールに存在するが最適化可能な一時変数はヘッドゴールと本体部の第一ゴールのトップレベルに現われる最初の出現に限られる。現在CHIのコンパイラで行なっている一時変数の最適化条件を次に示す。

1) ヘッドゴールでのget命令の省略

①ヘッドゴールのみに出現する変数

例 $p(X, Y, [Y]).$

②本体部での出現より後ろにある変数

例 $p(Y, a, X) :- q(X, X, Y).$

③同一位置に本体部の出現がある変数

例 $p(a, X, b) :- q(X, X, X).$

2) 本体部でのput命令の省略

①ヘッドゴールでの出現と同じか前にある変数

例 $p(X, Y, X) :- q(X, X).$

②ヘッドゴールでの同一位置にボイドがある変数

例 $p(X,) :- q(b, X).$

③ヘッドゴールの引数より後に出現する変数

例 $p(X, Y) :- q(a, b, X).$

一時変数の最適化は様々なパターンがあるが、省略できる命令はレジスタ転送命令のため、実行性能に対する効果はあまり大きくなく、appendの例でも10%以下でしかない。

6. まとめ

Prologにおける最適化の基本はいかに冗長なメモリアクセスを減らせるかにあり、インデクシング、終端再帰呼び出しの最適化の利点もバックトラッキング情報や制御情報を格納するフレームの生成を省略できる点にある。最適化に関する今後の課題としては変数の型を含めたインデクシングや、本稿では述べなかったが、冗長なカットの除去やクローズ内ORの効率的な実現がある。ただし、これらを実現するにはソースプログラムの解析だけでは不十分な点があり、より効果的な最適化を実現するためには変数の型宣言や引数の入出力指定等の導入が必要と思われる。

最後に、本稿を記述するにあたって種々のデータを収集してくれた神戸日本電気ソフトウェア株式会社の阪野氏に感謝いたします。

参考文献

- [1] Warren: "An Abstract Prolog Instruction Set" TR309, AI Center, SRI International, 1983
- [2] Nakazaki et al.: "Design of a High-Speed Prolog Machine" In Int'l. Symp. on Comp. Archi. 1985
- [3] Roy P. V.: "A Prolog Compiler for the PLM" UCRCSD 84/203, Berkeley 1984

表1 データ型によるインデクシング後の型別平均クローズ数(単純平均クローズ数2.7)

変数	1.7
構造体	12.3
リスト	1.1
アトムック	6.7

表2 クローズのインデックス引数の手続き内での出現パターン

全て変数	74%
全て非変数項	12%
変数/非変数項	1%
非変数項/変数	12%
それ以外	1%

表3 データベース検索問題(220クローズ)におけるインデクシング方式の実行速度比較

	ハッシュ検索	逐次検索	非決定的検索
先頭	0.9	0.5	0.9
中間	1.3	1.5	43.5
最後	1.0	2.6	84.3

表4 インタプリタのクローズ型の比率

組込述語展開	無し	有り
ユニット	5%	5%
トランジェント	18%	59%
その他	77%	36%