

TM-0078

情報処理学会第29回全国大会発表論文集
(37件)

ICOT ならびに再委託先メーカ

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

逐次型推論マシンのファームウェア

—— マイクロインタプリタの概要と特徴 ——

3-1

山本 明 横田 実 西川 宏 堀 和男 内田 俊一

(新世代コンピュータ技術開発機構)

中島 克人(三菱電機)

三井 正樹(神電機)

1. はじめに

パーソナル逐次型推論マシン(PSI)は第5世代コンピュータシステム研究開発の一環として開発されているソフトウェア開発用のパーソナルマシンである。

ψではKL0と呼ばれるPrologライクな論理型言語を機械語としており、このKL0をマイクロプログラムで記述されたインタプリタ(マイクロインタプリタ)と専用のハードウェアにより直接解釈実行する。本論文ではマイクロインタプリタの特徴と概要を報告する。

2. KL0

ψ上で実行されるプログラムはOSを含めて全て論理型言語で記述されるようになっており、OSを効率的に実行するためKL0にはPrologライクな機能の他に低レベルのハードウェアを操作出来る機能や拡張された制御機能が含まれている。KL0で記述されたクローズ表現は、コンパイラによりクローズヘッド部、ヘッドの引数部及び複数のゴールからなるボディゴール部から構成されるワード列のテーブル(内部形式)に展開される。マイクロインタプリタはこのテーブルを順次解釈しながら処理を行なう。

3. マイクロインタプリタの特徴

ψは論理型言語KL0及びその上に構築されたOSを効率的に実行するための各種の方式を採用している。

a. スタック方式とストラクチャ・シェアリング

KL0プログラム実行の基本的な動作はクローズの呼び出し、呼び出し元へのリターン及びユニフィケーションに失敗した際のバックトラック等の実行順序に関するものと、ユニフィケーションのようにデータ操作に関するものである。これらの動作を実現するため、マイクロインタプリタは4本のスタックを使用して実行環境の管理を行なっている。4本のスタックはそれぞれローカルスタック、グローバルスタック、トレイルスタック、コントロールスタックと呼ばれる。ローカルスタック、グローバルスタックは変数セルの格納用として、トレイルスタックは変数セルの無効化用情報の格納エリアとして、又コントロールスタックはクローズの実行順序を制御する情報の格納用として使用される。

ユニフィケーションで操作される構造体のデータ表現の形式として構造を共有するストラクチャ・シェアリング方

式と構造を新しくコピーするストラクチャ・コピー方式とがありそれぞれ長所と短所もつが、ψでは仮想記憶をサポートしていないこと、ストラクチャの共有によるアクセスのオーバーヘッドの方がコピーによるオーバーヘッドよりも少ないとの予想からストラクチャ・シェアリングの方式を採用している。

b. テイル・リカーション・オブティミゼーション

ユニフィケーションは呼び出し側ゴールの引数と呼び出された側のクローズヘッドの引数との間で行われる。しかし、これは呼び出し側での処理と呼び出された側での処理の2つのフェーズに分解して考えることが出来る。マイクロインタプリタはユニフィケーションに先立って呼び出し側ボディゴールの引数を評価し、呼び出される側の変数セル領域にその値をコピーした後、このコピーされたセルの値と呼び出された側のクローズヘッドの引数との間でユニフィケーションを行なう。このようにユニフィケーションを呼び出し側ゴールの処理と呼び出されたクローズヘッドの処理とに分解することにより、クローズヘッドのユニフィケーションが失敗し他の選択肢のユニフィケーションを行なう際にも以前に生成した呼び出し側の引数セルの値をそのまま使用でき、処理の高速化を計ることが出来る。

又、クローズの最後のボディゴールの呼び出しでは、そのクローズにもはや他の選択肢が存在せず決定的に処理が終了する場合そのクローズの変数セル領域は保存する必要がない。したがって、マイクロインタプリタはこのようなゴールの処理では呼び出し側の引数のコピーを完了した時点で呼び出し側クローズの変数セル領域を解放する。この方式はテイル・リカーション・オブティミゼーションと呼ばれ、不要な変数セル領域によるスタックの消費を極力抑えることが出来る。

c. 拡張制御機能

論理型言語Prologをプログラミング言語として実用的に使用出来るものになっている機能としてカットがある。カットはプログラムの実行時に自クローズの他の選択肢を刈り取ることににより、不要な枝の探索を行なわないようにする機能であるが、KL0にはさらに自クローズの先祖までもカットできる強力な遡隔カットの機能が備えられている。

その他、KL0にはユニフィケーションに失敗してバックトラックした際に指定された処理(on-backtrack)を行なう機能、ユニフィケーションによって特定の変数に値がバ

インドされた際、その時点であらかじめ指定しておいた処理を開始する機能 (bind-hook) が備えられている。

マイクロインタプリタはこれらの機能を実現するための制御情報をダイナミックに生成する。

d. OSサポート機能

OSの重要な機能のうちハードウェアにディペンデしたものととして割込み処理、プロセススイッチ、実メモリの割付け等がある。割込み処理では割込みの検出、割込み原因の固定エリアへのセーブ及び割込みハンドラへのプロセススイッチが必要であり、プロセススイッチ処理では実行中のプロセスの実行環境のセーブ及び新しいプロセスの実行環境のロードを行なう必要がある。これらの効率はシステム全体の性能に大きく影響する。したがって、 ψ における割込み処理、プロセススイッチはすべてマイクロインタプリタによって処理される。また、メモリ割付けについても実行中の各プログラムに対してマイクロインタプリタが動的に実メモリの割付け及び不要となったメモリの回収を行なう。マイクロインタプリタによりこれらの処理をサポートすることにより処理速度の向上だけでなく、OSのプログラム記述が容易になる。

4. マイクロインタプリタのモジュール構成

マイクロインタプリタはその機能により次の4モジュールから構成されている。(図1)

a. 核部

KL0 の内部形式と4本のスタックを使用して実行順序の制御及び管理を行なうと共にユーザ定義述語のユニフィケーション処理を行なう。内部形式中に相込述語のコードが検出されると相込述語部へ制御が移される。拡張制御機能実現のための実行順序の制御もこのモジュールで行なわれる。

b. 相込述語部

ψ には約160種の機械語として相込まれた述語(相込述語)が用意されており、核部により相込述語を実行する必要があることが検出されるとこのモジュール中の対応した相込述語処理ルーチンへ制御が移される。このモジュールでの各相込述語の処理が失敗または終了すると制御は核部へ戻される。

c. OSインタフェース部

核部又は相込述語部が割込みを検出するとこのモジュールに制御が渡される。このモジュールでは、割込み原因の解析、割込み原因の固定エリアへのセーブ及び割込みハンドラへのプロセススイッチを行なう。又、実メモリの自動割付けもこのモジュールで行なわれ、もし割付けの実メモリがなくなった場合にはガベージコレクタを呼び出す。

d. 初期化部

ψ を立ち上げるには実メモリやCPU内の各種のハードウ

ェアテーブルの初期化を行なう必要があり、これらの処理はこのモジュールによって行われる。

初期化の後OSのプログラムの基本部が実メモリにローディングされ、システムの立ち上げが完了する。

5. 開発サポート

ψ のマイクロインタプリタを早期に開発するためのツールとして、ほぼC言語風に記述出来るマイクロプログラムアセンブラとマイクロプログラムシミュレータが開発され、マイクロプログラムは入出力命令等のハードウェアにディペンデした部分を除いて全てマイクロプログラムシミュレータ上で開発された。

6. おわりに

マイクロインタプリタの開発はほぼ完了し、現在OSの開発が行なわれている。又これと平行してマイクロインタプリタの評価を行なっているが、性能はほぼ当初の目標を達成できる見込みである。今後さらに、個々の評価を行なうと共にシステム全体としての評価も行なう予定である。

最後に、日頃有益な助言をいただくICOTメンバ及び協力いただいたメーカの方々に深く感謝する。

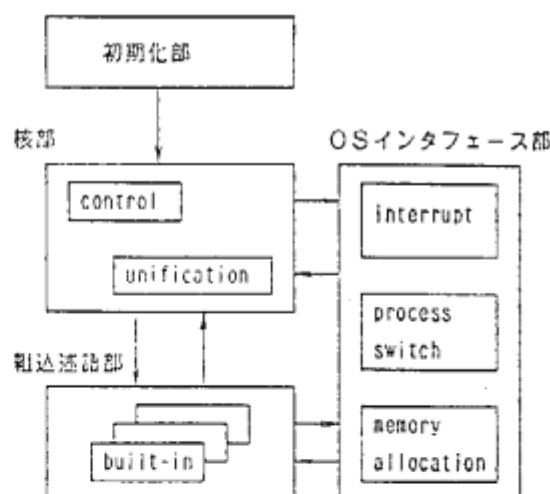


図1 マイクロインタプリタモジュール構成

参考文献

1. Warren, D., H.D. [An improved Prolog Implementation which optimises Tail Recursion] Proc. of the Logic programming workshop, Hungary (July, 1980)
2. Chikayama, T. et al. [Fifth Generation Kernel language Version-0] Proc. of the Logic programming conference, Japan (March, 1983)
3. Yokota, M. et al. [The Design and Implementation of a personal Sequential Inference Machine:PSI] New Generation Computing Vol.1 NO.2 '83 ohmsha, Ltd.

逐次型推論マシンのファームウェア

—— 基本制御部 ——

中島克人
(三菱電機)

山本 明
(新世代コンピュータ技術開発機構)

横田 実

1. はじめに

逐次型推論マシン (PSI) では KLO と呼ばれる Prolog ライクな論理型言語を直接解釈・実行するマイクロインタプリタ方式を採用している。本論文では述語呼出しやバックトラックなどの実行順序を制御するマイクロインタプリタ基本制御部の概要と、システム性能の向上のための処理方式について報告する。

2. 基本制御部の概要

基本制御部では 4 本のスタックを使用して、実行環境の管理を行なっている。

a. コントロール・スタック

述語呼出しの際の環境を、コントロール・フレームと呼ぶ 10 語を単位として格納する。コントロール・フレームは述語の呼出し元へのリターン時や、バックトラック時の環境の復帰に使用される。

b. ローカル・スタック

ヘッド引数および、ボディ引数にのみ現れる変数のためのセル (ローカル・フレーム) を述語呼出し単位に格納する。ローカル・フレームはスクラッチパッド・メモリ上に設けた 2 面のバッファに交互に生成され、保存が必要なフレームのみ、スタックに格納される。

処理対象のクローズが決定的に終了する場合には、コントロールおよびローカル・フレームはたたみ込まれる (テイル・リカーション・オブティミゼーション: TRO)。つまり、すでにそのクローズのフレームが存在すればそれを放棄し、まだ存在しなければ格納しない。

c. グローバル・スタック

構造体内変数のためのセルや、拡張制御機能用の情報など、クローズの実行が決定的に終了しても解放できない情報を格納する。

d. トレール・スタック

バックトラック時に値の束縛を解放 (UNDO) すべき変数セルのアドレスを積む。また、拡張制御機能用の情報の内、UNDO すべきものに関しても、そのアドレスとデータの 2 語を対として格納する。トレール・スタックの先頭部分は、やはりスクラッチパッド・メモリ上に設けたバッファに保持される。

基本制御部の概略フローを図 1 に示す。

図には示していないが、基本制御部では、述語呼出しの切れ目における割込チェックや、拡張制御機能に述語呼出しのレベルの深さの管理なども行なっている。

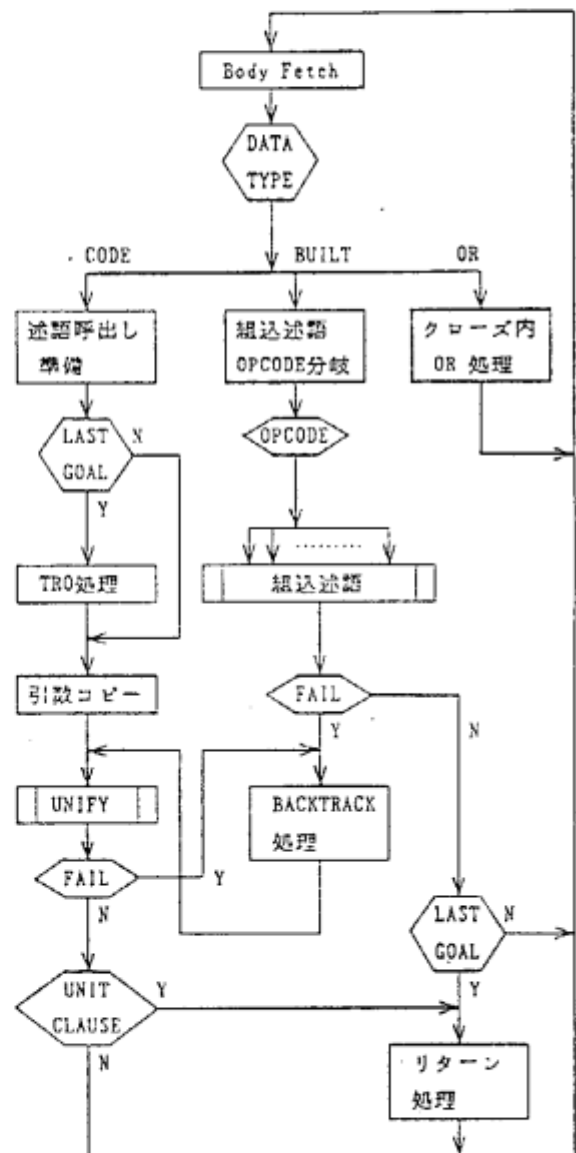


図1 マイクロインタプリタ基本制御部概略フロー

3. 引数コピー方式

KL0 の内部コードにおいては、クローズ内の変数は、そのクローズの呼出し元のフレーム内オフセットという意味の変数番号で表現されている（図2）。インタプリタでは、そのクローズのボディゴール部からの逆送呼出しにあたって、まずボディ引数の評価を行ない、ローカル・フレームを生成する。つまり、ボディ引数が変数の場合は、その変数番号をオフセットとしてそのクローズの呼出し元フレームをアクセスし、新しいフレームにコピーする。ボディ引数が定数の場合は、その値自身をコピーする。

ユニフィケーションは、ボディ引数の情報に基づきコピーされたローカル・フレームと、呼出された側のクローズヘッド引数との比較という事で行われる（図2）。

この方式は、ローカル・フレームへの格納・読出しが若干のオーバーヘッドになるが、

- (1) ユニフィケーションが成功する場合には、このローカル・フレームはいずれ呼出し側クローズのボディゴール引数の評価時に、呼出し元フレームとして使用されるため、引数評価の時間は同じである。
- (2) ユニフィケーションが失敗し、他のクローズが呼び出される場合には、コピー済みのフレームが再利用される。
- (3) 呼出しがクローズの最後のゴールであり、かつ、処理が決定的な場合には、必要な引数の値はすでに新しいフレームにコピーされているため、呼出し元のフレームの保存が不要となり、対応するスタック領域の解放が可能

となる。よってコピー方式の実現を目的とした専用のハードウェアにより、バッファ・アクセスのオーバーヘッドは小さくなっている。

4. クローズ・インデックス

KL0 では呼出し側の任意の引数によるクローズ・インデックスを可能としており、第1引数でのインデックスの後に、更に第2引数でインデックスするなど、多重のインデックスも可能である。KL0 コンパイラは、インデックスに用いる引数の位置、ハッシュによる分枝数および分枝テーブルをコードとして生成することにより、インタプリタへの指示を与える。インタプリタでは、クローズ・コードの位置に挿入されたこのコードを解釈し、指示された呼出し側引数を評価し、ハッシュ値を計算し、分枝テーブルによりクローズを選択する。

現在、ハッシュ関数は1種類であるが、コードの指示により、複数のハッシュ関数を選択する方式も容易に実現できる。

多重のクローズ・インデックスは、大量のデータ・ベース・クローズの選択などには非常に有効な手段となるであろう。

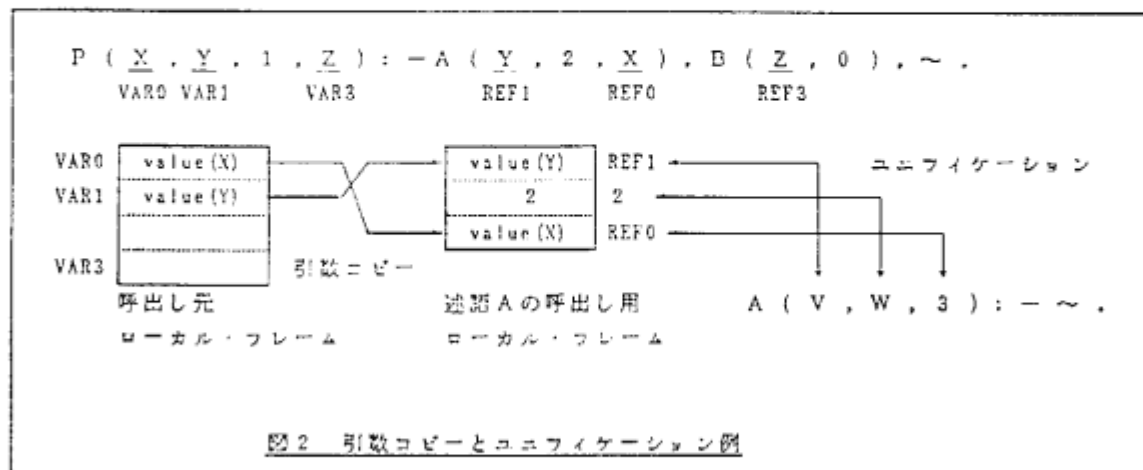
5. おわりに

マイクロインタプリタの基本制御部は約800語/64bitのマイクロ命令で構成されている。今後、処理速度向上のために、各種の評価および改良を行なっていく予定である。

最後に、本インタプリタの作成にあたり、適切な御指導・御助言をいただいた、ICOTメンバおよび、関連メーカの方々に感謝する。

参考文献

1. Warren, D. H. B. "An improved Prolog Implementation which optimises Tail Recursion" Proc. of the Logic programming workshop (July, 1980)



逐次型推論マシンのファームウェア — ユニフィケーション部 —

1E-3

三井 正樹 機田 実 西川 忠 藤村 茂
(沖電気) (新世代コンピュータ技術開発機構) (シャープ)

1. はじめに

パーソナル逐次型推論マシン (PSI) は K L 0 と呼ぶ Prolog 的な論理型言語を機械語としている。従って、そのファームウェアは K L 0 を直接解釈実行するインタプリタとして機能する。

本稿では、マイクロインタプリタ板部のうち、ユニフィケーション部について報告する。

2. 概要

ユニフィケーション部は K L 0 のユーザ定義述語に関して、呼出し側ボディゴール引数と呼ばれた側クロズヘッド引数との間の比較あるいは代入といったユニフィケーション機能を担っている。

呼出し側の引数は引数コピー方式の採用により既に基本制御部に於て評価され、呼ばれた側の変数セル領域の先頭にコピーされているので、変数セル領域のポインタ経由で逐次的に取出せることが約束されている。呼ばれた側の引数はコード中に存在し、命令ポインタ経由で取出せる。引数間のユニフィケーションの順序は本質的には任意であるが、ここでは前述の2つのポインタをインクリメントしながら第1引数より逐次的に行なっている。

ユニフィケーションの手順を図-1にフローチャートで示す。

全引数のユニフィケーションが成功すると基本制御部の成功リターン処理に制御を戻す。又、失敗した場合はその時点でユニフィケーションを打ち切り基本制御部のバックトラック処理に制御を戻す。

3. ユニフィケーションルール

ユニフィケーションの成功・失敗は取出した引数のデータタイプに応じたユニフィケーションルールの適用によって決定する。ユニフィケーションルールの説明と一覧表(表-1)を以下に示す。

(1) アトミックデータ

タグ部のデータタイプとデータ部の値の比較を行ない一致すればユニフィケーションは成功、一致しなければ失敗とする。

(2) 構造体データ

構造体データには本体がヒープ中に格納されている場合とグローバルスタック中にその本体セルが作り出されている場合の2種類がある。サイドエフェクトを伴うヒープ中の構造体は物理的に同一であ

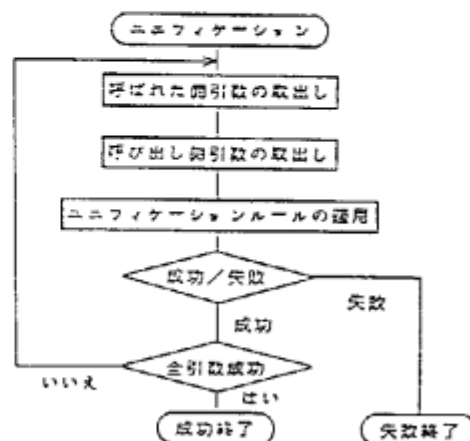


図-1 ユニフィケーションの手順

表-1 ユニフィケーションのルール

引数 データタイプ	変数	変数	変数	変数	変数	変数	変数	変数
変数	成功	成功	成功	成功	成功	成功	成功	成功
変数	成功	比較 失敗	失敗	失敗	失敗	失敗	失敗	失敗
変数	成功	失敗	失敗	失敗	失敗	失敗	失敗	失敗
変数	成功	失敗	失敗	失敗	失敗	失敗	失敗	失敗
変数	成功	失敗	失敗	失敗	失敗	失敗	失敗	失敗
変数	成功	失敗	失敗	失敗	失敗	失敗	失敗	失敗
変数	成功	失敗	失敗	失敗	失敗	失敗	失敗	失敗
変数	成功	失敗	失敗	失敗	失敗	失敗	失敗	失敗

* 比較の結果 一致すれば成功

* スタック中のベクタでは各要素について、ユニフィケーションのルールを適用する。

るとき成功とする。undo処理の対象となるスタック中の構造体は要素同士の比較を行なって成功・失敗を決定する。各構造体についてのユニフィケーションルールを示す。

(a) ヒープベクタ

両方のヒープベクタの要素数と本体アドレスの一致を調べる。

(b) ストリング

両方のストリングの要素数とタイプ、オフセット及び本体アドレスの一致を調べる。

(c) コード

両方のコードの本体アドレスの一致を調べる。

(d) プロテクティドタイプ

両方のタイプ記述子のアドレスの一致を調べる。

(e) スタックベクタ

両方のスタックベクタの要素数の一致を調べ、個々の要素同士と比較又は代入を行なう。

(3) アンバウンド変数セル

(a) 共にアンバウンド

アドレスの大きい方から小さい方にポインタを張って両方のデータが論理的に同一であることを示す。ポインタが代入される変数セルにフックがかけられていたらその内容を他の変数セルに移す。両方共フックがかけられているときは更にその内容にマージを施し、アドレスの小さい方にその結果を格納する。

(b) 片方のみアンバウンド

アンバウンドであるセル中に他方のセル内容をコピーする。フックがかけられていたらその内容は、内の特定レジスタとの間でマージを行なう。

4. 構成と実装性

ユニフィケーション部のルーチンは図-1の手順を図-2に示すモジュール構成で実現したマイクロプログラムである。ルーチンの大半はユニフィケーション部の適用部で表-1の縦軸を分類するグループと、横軸を分類し規定された処理を行なうグループの2つに大きく分けることができる。前者のグループはヘッド側引数オペランドに対して

- ・トップレベルのオペランドの分類
- ・短縮形オペランドの分類
- ・構造体要素オペランドの分類

の3種のモジュールとそれに付随するデレファレンスのモジュールから成っている。後者のグループは分類されたヘッド側引数オペランドのデータタイプに対応して

- ・アトミック、コード、プロテクティドタイプ
- ・未定義変数
- ・バインドフック変数
- ・ヒープベクタ
- ・ストリング
- ・スタックベクタ

といったデータタイプ別モジュールに分けられる。

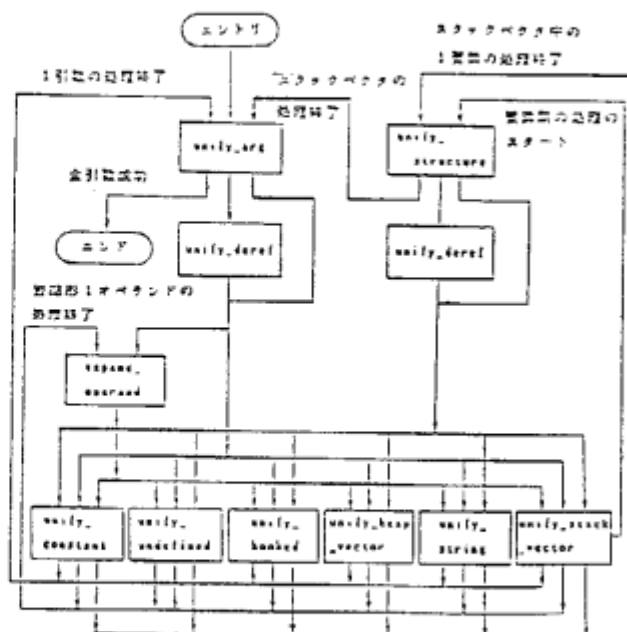


図-2 モジュール構成

後者のグループのモジュールは前者のグループで共用され、ボディゴール側引数オペランドを分類して表-1に対応する処理を行う。共用の方法は、ハードウェア資源のjr(jump register)に成功・失敗時の戻り先情報を格納しておき処理終了後jrの間接ジャンプオーダを発行するという形で実現している。

データタイプの分類はタグディスパッチ機能を有効に利用して高速化を図っている。前者のグループはヘッド側引数オペランドのデータタイプをタグディスパッチ機能で分類し後者のグループへ分岐する。後者のグループはボディゴール側引数オペランドとのタグの比較を陽に行なわず、タグディスパッチを利用した自動分類によってヘッド側引数オペランドとつぎあわせ、成功・失敗の処理を行なっている。

タグディスパッチは2種類の分岐条件指定を使用している。それぞれ5面、8面の受け口があり、あわせて38回のタグディスパッチオーダを発行しユニフィケーション部の高速化に役立たせている。

5. おわりに

マイクロインテグリティの開発はユニフィケーション部も含めてほぼ完了し評価の段階に入っている。今後はシステム全体の評価による見直しを進めていく予定である。

参考文献

- Yokota, M., et al. "The Design and Implementation of a Personal Sequential Inference Machine:PSI," New Generation Computing Vol.1 No.2 '83 ohusha, Ltd.

逐次型推論マシンのファームウェア —データ・タイプとその操作—

小松 英二 三井 正樹 (沖電気)
横田 実 山本 明 西川 宏 宮崎 敏彦
(新世代コンピュータ技術開発機構)

1. はじめに

パーソナル逐次型推論マシン ϕ は、第五世代コンピュータ・プロジェクトにおけるソフトウェア開発用パイロット・モデルである。 ϕ にはKL0と呼ばれるPrologライクな言語をベースとした機械語があり、約160種類の組込述語をサポートしている。データ操作系組込述語はこのうちの約50種類を占めている。 ϕ におけるデータ操作は、ヒープ上のデータ操作に特徴を有している。本編では、 ϕ ファームウェアにおけるデータ操作部として、ヒープ上のデータ操作に触れながら、データ操作系組込述語について報告する。以下、特に断らない限りスタックとはローカル・スタックまたはグローバル・スタックを指すとする。

2. データ・タイプ

ϕ には将来の拡張用も含めて約60種類のデータ・タイプがある。以下に特徴のあるデータ・タイプを幾つかあげる。

a. ヒープ・ベクター、スタック・ベクター

一次元の配列であり、その本体の置いている場所がヒープかスタックかにより、ヒープ・ベクターまたはスタック・ベクターと呼ぶ。

b. コード

KL0のプログラムのクロズを返すために用いられる。構造はヒープ・ベクターと同じであるが、誤動作を防ぐ為、別のデータ・タイプにしてある。

c. モレキュール

コード中のヒープ・ベクターは、変数の値がスタック中にある。これと、オブジェクト生成述語により、動的に作られたヒープ・ベクターを区別するために、読み出し時にモレキュールを作り、スタック・ベクターと同じ扱いにする。

d. ロケーション

このデータ・タイプは「visibleなポイント」と考えられ、構造は一要素のベクターと同じである。ロケーションの指す要素はdangling referenceを防ぐ為ヒープ上のデータ・タイプに限っている。図1、2にモレキュール及びロケーションの構造を示す。

e. プロテクテッド・タイプ

プロテクション・キーと呼ばれる保護キーと保護されるデータを持った構造体である。図3に構造を示す。

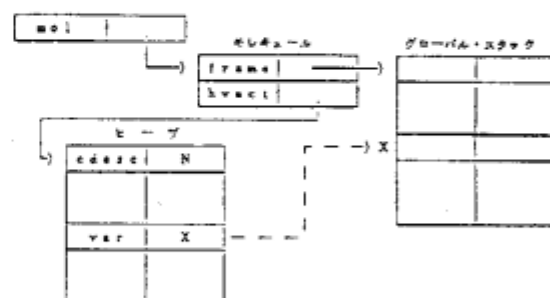


図1 モレキュールの構造

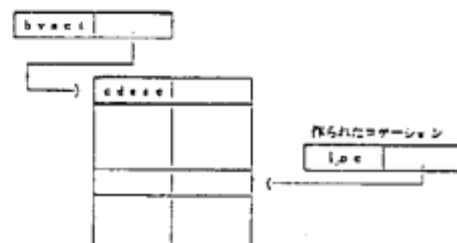


図2 ロケーションの構造

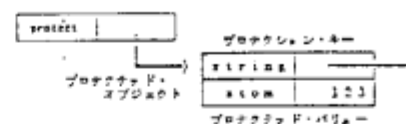


図3 プロテクテッド・タイプの構造

3. non stack objectとstack object

データは、ヒープ又はスタックに置くが、スタックは動的に伸縮するためにヒープからスタックへのポインタは、ロケーションのときのようにdangling referenceになる可能性がある。このため ϕ では、ヒープ上に置くことのできるデータ・タイプを制限することにより、ヒープからスタックへのポインタを防止している。

a. non stack object

ヒープ上に置くことのできるデータ・タイプをnon stack objectとよぶ。このグループには、未定義語のように将来ポインタを張る可能性のあるデータ・タイプも除いている。上記のヒープ・ベクター、コード、ロケーション及びロケーションの指すデータはnon stack objectである。

b. stack object

スタック上に現れるデータ・タイプをstack objectと呼ぶ。スタック・ベクタ、モジュール、はstack objectである。

c. その他

インタプリタ制御用データ・タイプの中には、トレイル・スタックにしか現れないものがある。このようなデータ・タイプは、non stack objectでもstack objectでもなく、このデータ・タイプがヒープ又はトレイル以外のスタックに現れたときはエラーとなる。

ψのデータ・タイプは上記のa)~c)のどれかに属する。

4. データ操作

ヒープ上のデータとスタック上のデータではデータの容換えに関して、データ操作における扱いが全く異なる。ヒープ上のデータは常にサイドエフェクトにより上書きが行われ、バックトラックによってundoされない。一方、スタック上のデータは、ユニフィケーションによってのみ書き換えられ、バックトラック時にはundoする。ヒープへの上書きはnon stack objectに限る。

5. データ操作系組込述語

データ操作系組込述語は、次のような3つのグループからなる。

a. 属性チェック組込述語

データのタイプ・チェック及びその他の属性の取りだしに用いられる。タイプ以外の属性としては、データのバリエーション、構造体データの長さ、ストリング要素のタイプ（ビット、バイト、2バイト、ワード）、プロテクテッド・タイプのプロテクション・キー及び値などがある。

述語の例: atom(X)、value(X, Value)

code(X, Length, Xarg)、string(X, Length, Type)

b. オブジェクト生成組込述語

アトム、ヒープ・ベクタ、スタック・ベクタ、ストリング、プロテクテッド・タイプ、コードを生成する述語を用意している。構造体データの本体は作るデータ・タイプによって、ヒープまたはスタックのトップに生成し、データは第一引数にセットする。

述語の例: new_heap_vector(X, Length)

new_protected_object(X, Key, Value)

c. 構造体データ操作組込述語

構造体データ操作述語は、構造体アクセス、サブ構造体アクセス、ロケーション・アクセスに分けられる。さらに、データのバインディングの方法が、ユニフィケーションによるものとサイドエフェクトによるものを用意している。

o 構造体アクセス

構造体の要素を取り出したり、上書きを行ったりする述語である。サイドエフェクトにより書き込みが行われる構造体はヒープ上に本体を持つ構造体になり、また書き込む要素はnon stack objectである。

述語の例: vector_element(Vector, Position, Element)

set_vector_element(Vector, Position, Element)

o サブ構造体アクセス

構造体のサブ構造体を取り出したり、他の構造体を上書きしたりする述語である。サイドエフェクトによりそのデータを書き込むベクタは要素がnon stack objectでなければいけない。サブ構造体は記述子だけを新たに生成し、要素は元の構造体と共用する。図3にサブ構造体の構造を示す。

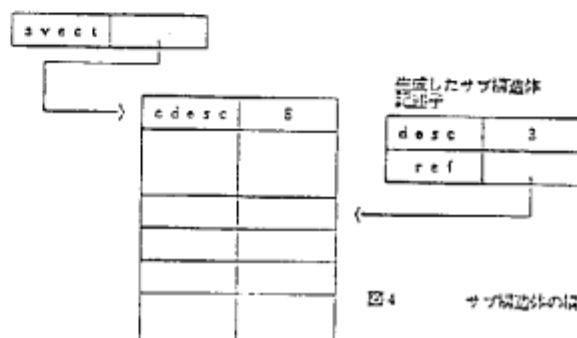


図4 サブ構造体の構造

述語の例: sub_vector(Vector, Position, Length, Subvector)

set_sub_vector(Vector, Position, Length, Subvector)

o ロケーション・アクセス

ロケーションの生成や、ロケーションの指しているデータの取り出し、及びデータ・セットを行なう述語を用意している。ロケーションにセットするデータはnon stack objectに限る。

述語の例: vector_element_location(Vector, Position, Element)

location_element(Location, Element)

set_location_element(Location, Element)

6. おわりに

ψでは、データ操作に関して、内部的にはヒープ等に関連したデータ・タイプのチェックを徹底的に行い信頼性を増す一方、K L O 言語レベルでは、ロケーションなどの導入により多様な機能を提供している。また、本編でのべた以外に、データ操作系全体での整合性や特権述語の設定により、ユーザ・プログラムであるコードの破壊防止なども考慮している。

逐次型推論マシンのファームウェア

—— 割込み処理とプロセス管理 ——

1E-5

池田 守宏 横田 実 山本 明 瀧 和男 西川 宏
(三菱電機) (新世代コンピュータ技術開発機構)

1. はじめに

パーソナル型逐次推論マシン ϕ (PSI) では、プログラムの実行の単位をプロセスと呼び、OSのプロセス管理ルーチンの管理下で複数プロセスの実行環境を作っている。このようなマルチプロセスの管理には、複雑で細かなハードウェア資源の操作が必要であるが、それによる処理速度の低下及び誤操作を避けるため多くの部分をファームウェア化し、OSとのインタフェースを簡潔にすることに努めている。このプロセス管理のファームウェアは、割込み、プロセス切替処理を中心としたOSインタフェース部と、ソフトウェアインタフェースのための組込み制御部の2つの部分から構成されている。

本論文ではこれらファームウェアに関し、論理型機械言語KL0実行マシンとして特徴的な処理を中心にその概要を報告する。

2. プロセス

a. プロセスの実行環境

ϕ では最大63個のプロセスを設けることができる。各プロセスはそれぞれ独立に機械語KL0のプログラムを実行する。このためにPCB、EPCBと呼ばれる実行制御用の情報をそれぞれのプロセスが個別に持っている。

PCB (process control block)はソフトウェアが参照、操作するステータス情報でプロセスの割込みレベル等より成り4語で構成される。EPCB (extended process control block)はファームウェア、マイクロインタプリタが参照、操作するための制御情報である。これはマイクロアドレススタックの内容、汎用マイクロレジスタ、命令レジスタ、スタック操作ポインタ等より成り最大60語で構成される。以上の制御情報の他各プロセスは、それぞれKL0実行時に必要な4本のスタックエリア (グローバル、ローカル、コントロール、トレイル) 用に主記憶が割当てられており、エリアごとのATP (area top pointer)がこれらスタックのトップを管理している。

b. プロセス切替

プロセス切替は上記プロセス実行環境を全てファームウェアにより切替る処理である。

プロセス実行時PCB、EPCB、ATPは、ハードウェアレジスタ上にロードされ、KL0実行に従ってファームウェアにより参照、更新されている。ここでプロセス切替の要求が起るとプロセス切替の処理ルーチンがサブルー

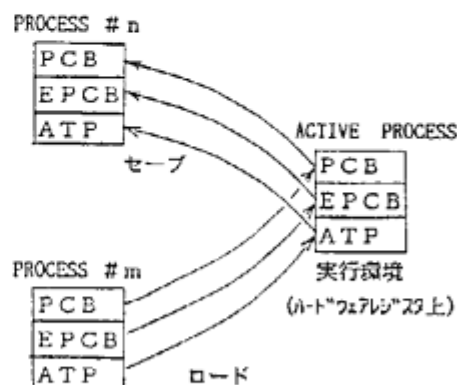


図 1 プロセス切替の状態

チンコールされる。これによりPCB、EPCB、ATPはプロセスごとの退避領域に格納される。その後新たに起動されるプロセスのそれら情報が、退避領域からロードされ新しい実行環境が作られる。このプロセス切替の状態を図1に示す。

3. 割込み処理

プロセス切替の要因として割込み及びトラップがある。割込みはプログラムと非同期に発生しPCB中のマスクレベルの設定により抑止出来るもの、トラップはプログラムとは同期して発生し抑止出来ないものである。それぞれの要因を表1に示す。ファームウェア上割込み、トラップは特に区別無く扱い、以後単に割込みと称する。

表 1 割込み、トラップ要因

名称	要因
カ-26 I5- 1 トラップ	ハートウェア 検出メモリアドレス I5-
カ-26 I5- 2 トラップ	ウェアウェア 検出I5-
GC コールトラップ	ガ-1-3コレクション要求
I57 リミットチェック トラップ	指定サイズを越えたメモリアドレス
トレース トラップ	KL0 ステップ 実行用
オーバーフロー割込み	エラー割込み
温度、電圧異常割込み	CPU 内部温度、電源異常
入出力高レベル割込み	入出力装置 (高レベル)
タイマ割込み	インタ-ハ-タイマ-
入出力低レベル割込み	入出力装置 (低レベル)
メモリアドレス訂正割込み	主記憶16ビットI5- 訂正

a. 割込みの検出

割込みの検出は原則としてKLO実行処理の切目、即ち、述語呼び出しの切目でハードウェアの割込みフラグを判定することにより行う。この割込み検出のタイミングを基本タイミングと呼び、ユーザ定義述語の場合は、ユニフィケーションの成功又は失敗の直後、組込み述語の場合は実行終了時点である。ここでのプロセス切替では、EPCBの退避情報が最も少く従って高速におこなえる。KLOプログラムと割込み基本タイミングの関係を図2に示す。

しかし許容時間の短い割込み又は緊急度の高い割込みに対しては、上記方式だけでは対処できずそれぞれ次のような処理を行う。許容時間の短い割込み（タイマ、電源異常割込み等）の場合、基本タイミングの他に実行時間の長い組込み述語中、及びユニフィケーション中に割込み検出点を設け、これにより一定間隔以内の割込みチェックを保証している。また要因発生時点で処理しないと以降の実行が継続できない緊急な割込み（物理メモリ不足によるGCOR等）については、要因発生時点で即刻処理ルーチンと呼び出す。これらの場合はEPCBの退避情報は基本タイミングに比べ多くなる。

b. 割込み処理

ファームウェア割込み処理ルーチンでは、割込みレジスタ上にセットされている割込み原因を読み取り、必要な情報を付加して割込み原因（割込みインデックス番号）に対応する「割込み原因コードテーブル」に設定する。さらに同じく割込み原因に対応する「割込みインデックステーブル」を読みだし、処理プログラムのプロセス番号を知りそのプロセスにプロセス切替を行う。割込みインデックステーブルはあらかじめソフトウェアにより設定されている。割込み応答プロセスのプログラムでは通常、割込み原因コードテーブルから必要な、原因情報を得る。

4. プロセス管理用組込み述語

プロセス管理用組込み述語は、ソフトウェアとのインタフェースのために設けられた一群の組込み述語である。

a. プロセス生成用組込み述語

新たにプロセスの実行環境を作るために initialize_process 及び set_process_control_block 組込み述語がある。initialize_process はEPCBの初期設定を行う組込み述語で3引数で起動され、第1引数でプロセス番号を指定し、第2引数でプログラムのスタートコード（ヒープ上当該プレディケートへのポインタ）、第3引数で使用する4本のスタックエリアのうちの先頭のエリア番号を指定する。あらかじめメモリ管理系の組込み述語を使用してスタックエリアを生成した後、この組込み述語を実行することによりEPCBの初期値が作られ格納される。

P: -Q, b1, b2, R.

↑ ↑ ↑

Q: -S, T.

↑

Q.

↑

P, Q, R, S, T; ユーザ定義述語

b1, b2; 組込み述語

↑; 割込み基本タイミング

図 2 KLOプログラムと割込み基本タイミング

set_process_control_block 組込み述語は4語構成のベクタ形式で表すPCBデータとプロセス番号の2引数で起動され、PCBデータを対応PCBに格納する。これによりPCBの初期設定又は値の更新がおこなえる。

b. プロセス切替、トラップ組込み述語

生成したプロセスへの実行の切替又は割込みプロセスからのリターン等、任意のプロセスへの切替には change_process 組込み述語が使われる。change_processでは引数として与えられたプロセス番号へ、無条件でプロセス切替を行う。

trap組込み述語は4語構成のベクタ形式で表す、割込み原因コードデータと、割込みインデックス番号の2引数で起動され、これにより任意のインデックス番号の割込みをソフトウェアで発生できる。

c. プロセス管理用テーブル、レジスタ類アクセス組込み述語

プロセス管理のためソフトウェアとのインタフェースとなるテーブル、レジスタ類をアクセスするための組込み述語で、割込みインデックステーブル、割込み原因コードテーブル、システムレジスタをそれぞれ読み、書くための組込み述語が用意されている。

5. おわりに

割込みトラップ処理、プロセス管理の他 ϕ では、OS記述用にメモリ管理等多くのファームウェアモジュールが用意されている。現在これらを使用してOSの開発が進められており、プログラムの記述が簡単になる等の効果がみられている。今後OS評価フェイズを持って性能面等の評価をおこなって行く予定である。

最後に本ファームウェアの開発に当り適切な御指導を戴いた、ICOTメンバーの方々及びコーディング、テストに御協力戴いた 三菱総合研究所、小沢政、岡沢敏に深く感謝する。

逐次型推論マシンのファームウェア

— マイクロプログラムシミュレータ —

守山 貞 (沖電気)

三井 正樹 (沖電気)

横田 実 (新世代コンピュータ技術開発機構)

1. はじめに

パーソナル逐次型推論マシン (PSI) は第5世代コンピュータシステム研究開発の一環として開発されているソフトウェア開発用のパーソナルマシンである。

ψマイクロプログラムシミュレータはファームウェアの開発をサポートするソフトウェアの一つとして、マイクロ命令の実行動作を論理的にソフトウェアでシミュレートすることにより、ψハードウェアの開発と並行してファームウェアの開発及び検証を行うことを目的として作成したものである。本論文ではマイクロプログラムシミュレータの機能及び構成等について報告する。

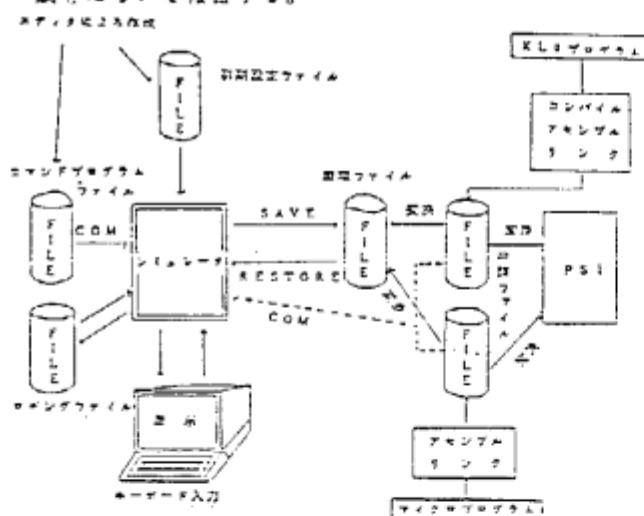


図-1 実行環境

2. 特徴

本マイクロプログラムシミュレータは、移植性、互換性を考慮に入れ、PASCALを用いて記述したものであり、約340個のモジュールに分かれ、約6000ステップより成る。

また、機能的には、割込み、アドレス変換等もサポートしており、ほとんど実機と同等であり、DEC2060上で、1秒間 (CPU TIME) に約400マイクロステップ実行可能である。

さらに、本シミュレータはデバッグを容易に行えるように、マイクロアドレスのトレース機能、トレース中トレース数が256に達したら実行を停止するリミット機能、マイクロ命令のステップ実行機能、

全てのハードウェア資源に対して設定可能なイベント条件を用いることにより、その条件が成り立てばトレースを開始するディレイドトレース機能または実行を停止するブレイク機能、そしてロギング機能を有する。そのほか、シミュレータは画面のスクロール範囲を自動的にコントロールし、自動表示、トレーススタック表示を行う機能、マクロコマンドの実行機能も有している。従って、これらの機能を利用することにより効率のよいデバッグが行える。

(図-1) にシミュレータの実行環境を示す。

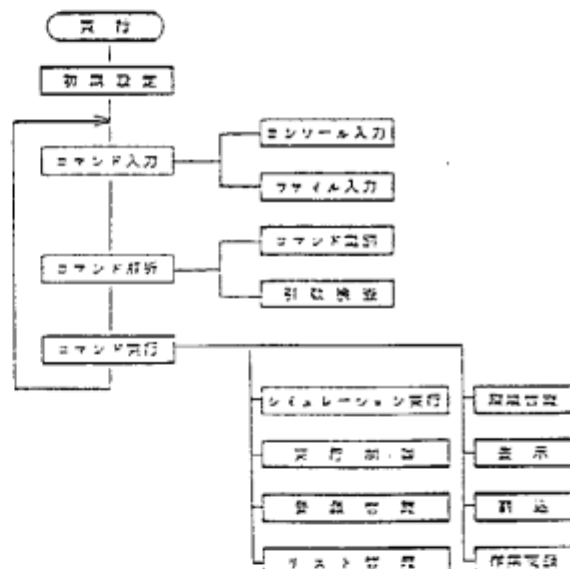


図-2 基本構成

3. プログラム構成

本シミュレータは、(図-2) に示す様に、大別して(a)初期設定部、(b)コマンド入力部、(c)コマンド解析部、(d)コマンド実行部から成り、シミュレーションを実行する。

(a) 初期設定部

シミュレータは起動されると、ハードウェア資源に相当する領域とシミュレータ内部資源をクリアし、ユーザーにより作成されたファイルを読み込んでコマンド名、タグのモニタリングの設定等を行いシミュレーションに備える。

(b) コマンド入力部

シミュレータへの入力、通常のキーボードからの入力の他ファイルからの入力があり、マクロコマ

ンドを実現している。

(c) コマンド検知部

コマンドが与えられるとシミュレータは初期設定部で設定したどのコマンドと一致するかを判別し、一致すればそのコマンドに適した引数が揃っているかを検査する。

(d) コマンド実行部

コマンド実行部は以下の8つの機能よりなる。

(i) シミュレーション実行機能

マイクロプログラムの実行をシミュレートするシミュレータ本来の機能で、連続実行、継続実行等が可能である。

(ii) 実行制御機能

シミュレーション実行を制御する機能で、モードを設定することにより行われる。

(iii) 内容表示機能

φのハードウェア資源及びシミュレータの内部資源に対して内容表示、内容変更をする機能である。

(iv) リスト管理機能

実行制御モードの条件リスト、自動表示の対象を示すハードウェア資源リストを管理する機能である。

(v) 表示機能

実行制御モードの状態を常時表示すると共に、自動表示として指定されたハードウェア資源の内容を実行停止毎に画面上部から表示する機能である。

(vi) マクロコマンド実行機能

コマンドプログラムファイルに定義されているマクロコマンドを実行する機能である。

(vii) 環境変数管理機能

シミュレーションの一時停止、再実行等に必要な実行環境の退避、再設定を行う機能でバイナリファイルを用いて行う。

(viii) 動作速度支援機能

10進数と16進数の変換をサポートする機能である。

表-1に以上の機能を実現しているコマンドの一覧表を示す。

4. シミュレーション方式

本シミュレータは、ハードウェア資源（レジスタ、レジスタファイル等）をすべてソフト的な変数や配列として持ち、ソース1、ソース2等のバスも変数として持っている。それらの資源をmarr（マイクロアドレスレジスタ）の指すマイクロ命令に従って操作しシミュレーションを実現している。marrの

指すマイクロ命令は64ビットで構成されており、シミュレータはマイクロ命令の上位32ビット、下位32ビットを2つの整数として受け取り、各フィールド毎のビットパターンを値として切り出し、そのフィールドに応じた処理を行う。

また、特殊な処理として、割込み処理がある。シミュレータは、割込み名をコマンドとして受け取りその引数のON/OFFの指定により、割込み・トラップ用の論理的に連動する3つのレジスタを自動的に操作する。マイクロ命令のシミュレーション実行中に割当てられていないメモリへのアクセス等の割込みが起こると、シミュレータはコマンド指定による割込みと同じことを自動的に行う。そして、外部デバイス等からのシミュレータ内部では起こりえない割込みをシミュレートする場合は前述のブレイク機能やシミュレータ自身の割込み（^Y）を使ってマイクロ命令の実行を一時停止させ、割込み名を指定して割込みが起こった状態とした後、マイクロ命令を継続実行させる等の方法がある。

機 能	コマンド名	注 記
シミュレーション実行	GO TO	任意のアドレスからのマイクロプログラムの実行
	START	mainの指すマイクロ命令の実行
実行制御	STOP	ストップ コードの指定、解除
	BREAK	ブレイク コードの指定、解除
	TRACE	トレース コードの指定、解除
	STEP TRACE	ステップトレース コードの指定、解除
	LIMIT	リミット コードの指定、解除
表示機能	LIST	リスト コードの指定、解除
	DISPLAY	資源の内容表示
リスト管理	CHANGE	資源の内容変更
	SET	イベント表示リスト、及び自動表示対象ハードウェア資源リストの登録
	RESET	イベント表示リスト、及び自動表示対象ハードウェア資源リストの取消
マクロコマンド実行	COM. COMNU	コマンドプログラムファイルからコマンド入力、及び実行
	COM LOG	
	PAUSE	マクロコマンド実行の一時停止
	CONTINUE	マクロコマンド実行の継続
環境変数	SAVE	シミュレーション実行環境の保存
	RESTORE	シミュレーション実行環境の復元、再設定
基 本	RESET	資源の初期化
	PREVIOUS	一つ前の自動表示の内容を再表示
作業支援	DETOH	16進 → 10進
	HTOH	10進 → 16進

表-1 コマンド一覧

5. おわりに

φのファームウェアはすべて実機完成前にシミュレータを用いてデバッグが行われたため真機デバッグを短期間で完了できた。更に基本的な性能評価も本シミュレータを用いて行われている。

今後ともこうしたシミュレータはファームウェア開発には不可欠のものであろうが、今回その有効性を確認した。

最後に日頃有益な助言をいただくICOTの方々及びノーマルの方々に深く感謝する。

逐次型推論マシンのハードウェア

—— 入出力システムと割込制御方式 ——

1E-7

三石彰純 京 敬人 眞取東朋
(三菱電機)西川 宏
(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システム研究開発プロジェクトにおいてソフトウェア開発用ツールとして開発された逐次型推論マシン(PSI)は、マウス、ビットマップディスプレイ、LANインタフェース、ディスクなどを備えた入出力装置を備えている。これらの入出力装置はIEEE-796に準拠した入出力バスに接続されており、その制御部であるKLO(Kernel Language 4.0版)の入出力組込言語を用いて制御される。またここでは優先順位を付いたペフター割込方式の割込制御を採用しており、割込によってプロセスの切替えを行っている。本報告では入出力システムと割込の制御方式について述べる。

2. 入出力システム

この入出力装置は図1. に示すように全て入出力バスに接続され、入出力制御部を介してCPU

から制御される。入出力バスは16メガバイトのアドレス空間を有しており、この空間は主記憶のメモリ空間とは完全に独立している。入出力バス空間には512キロバイトの入出力バスメモリ、7.5メガバイトのビットマップメモリ(ビットマップディスプレイの制御部に内蔵)が実装され、64キロバイト分の空間が制御装置への起動用ある

いは割込のリセット用の制御装置識別アドレスとして使用されている。残りの空間は現在未使用。

CPUが入出力バスに接続されている制御装置との間でデータ転送を行なう場合、CPUは入出力組込言語、例えばput.byte, put.vector 等を用いて

全ての入出力装置に共通なHMCB(Hardware Module Control Block)及びTMSG(Transmission Message)と呼ぶ二種の情報を入出力バスメモリ上に生成し、入出力バス空間上の制御装置アドレスを用いて起動をかける。入出力制御装置はHMCB及びTMSGに従って動作し、動作完了後CPUに割込を送る。CPUは入出力完了の割込を受信後HMCB, TMSGの状態情報を確認することで一連の動作を完了する。なおデータ転送は入出力装置と入出力バス

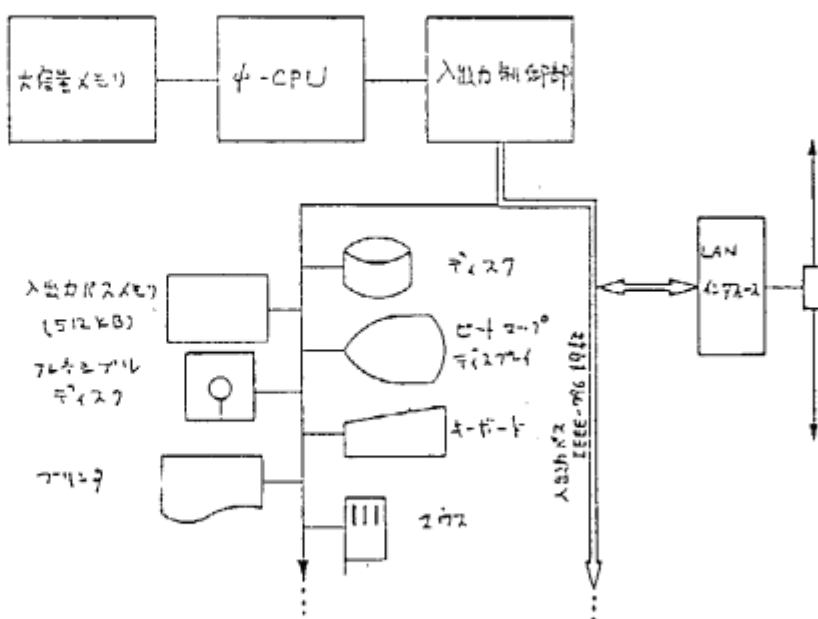


図1. この入出力システム構成

メモリの間で行われ、入出力バスメモリとCPUとの間の取り扱いは別途入出力制御装置を用いて行われ、その入出力制御装置は以下の7機能が用意されている。

(入力)

(出力)

get-byte	put-byte
get-double-byte	put-double-byte
get-double-byte-swapped	put-double-byte-swapped
get-vector	put-vector
get-vector-swapped	put-vector-swapped
get-string	put-string
get-string-swapped	put-string-swapped
get-with-tag	put-with-tag
(バスの初期化)	bus-reset

3. 割込制御

4は割込をプロセススイッチで実現しており、イベントによりプロセスを切替える機構として割込とトリップを併せている。割込は入出力動作やような外部的原因によりプログラムの実行とは非同期的に発生し、トリップは不当なページへのアクセスのような内部的原因によりプログラムの実行に同期して発生するものである。4の割込方式は優先順位が付いたベクタ割込方式であり、優先順位をレベル持っている。

7は優先順位はなく、マスクもできない。割込制御回路は図2に示すようにフリップフロップ群、2組のエンコーダ、4本のレジスタがあり、本回路でトリップも扱っている。割込またはトリップの要求が発生すると、フリップフロップ群内のフリップフロ

ッポがセットされる。割込要求の優先順位がIMSKR (Interrupt Mask Register) のビットに与えられた優先順位より高い場合、要求は受け付けられる。トリップ要求の場合はIMSKRの内容に関係なく受け付けられる。割込/トリップが受け付けられると、CPUはエンコードされた割込要因コードを用いて割込インデックステーブルを引き、対応する割込ハンドルのプロセスを起動すると共に、割込原因コードテーブルをセットする。割込インデックステーブルは主記憶中に存在する64ワードの表であり、各ワードには割込ハンドルのプロセス番号が入っている。割込原因コードテーブルも主記憶中の表であり、1エントリ4ワードで構成されている。各エントリには割込原因の詳細コードと付随するパラメータ及び被割込プロセスの番号が格納される。

4. おわりに

本報告では4の入出力システムと割込制御についてその制御方式を述べた。現在は入出力システムの機能強化を検討している。

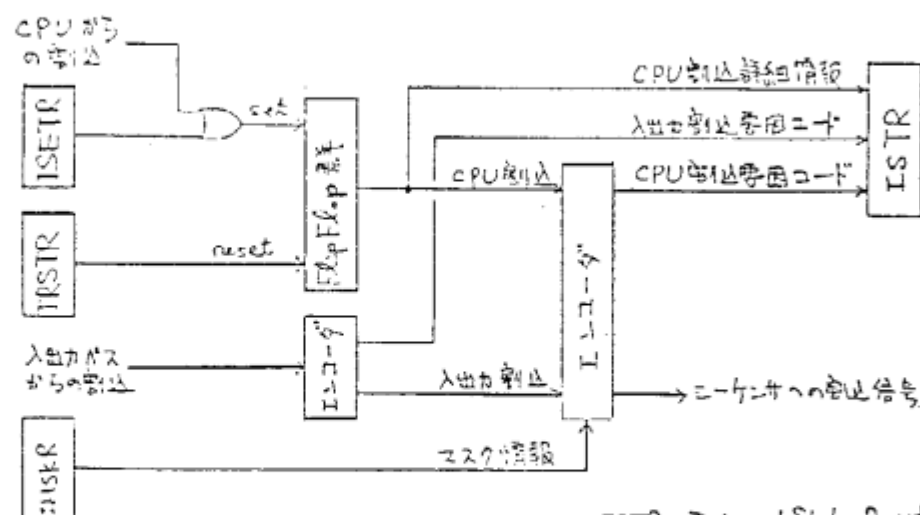


図2. 割込制御回路の構成

ISTR: Interrupt Status Register
 ISETR: Interrupt Set Register
 IRSTR: Interrupt Reset Register
 IMSKR: Interrupt Mask Register

逐次型推論マシンのハードウェア

12-2

—— システム制御とRASについて ——

大上 貴英 吉市 昌一

森 和男

(三菱電機)

(新世代コンピュータ技術開発機構)

1. はじめに

逐次型推論マシン(PSI)は第5世代コンピュータ・プロジェクトで開発が進められているソフトウェア開発用のパーソナル・マシン〔1〕であり、PROLOGに似た論理型言語KL0を直接実行する。ここでは、このψについてそのシステム制御とRASについて報告する。

2. システム制御

2.1. システム制御概要

ψのシステム制御においては、ファームウェアで実現されるマイクロインタプリタやKL0で記述されるプログラムの開発を支援する機能の実現を目標に設計されており、従来のミニコンの上位クラスのものに相当する機能になっている。これは、ψがパーソナル・マシンとはいえ、将来これを使って各種の大規模なソフトウェアの開発が行われることになっており、これを支援できるように考慮されているためである。

ψのシステム制御はコンソール・プロセッサ(CSP)を中心に行われる。CSPは、システムの立ち上げ、実行のスタート/ストップ、各種実行モードの設定、エラー処理、ファームウェアやソフトウェアのデバッグ支援、CPUやメモリの各種レジスタのリード/ライトなどを行うマイクロプロセッサをベースにしたプロセッサである。CSPの構成と機能については他の報告〔2〕で述べられているので本報告では主にCPUとのインタフェースについて述べる。

CSPとCPUとは8ビットのインタフェース・データ・バス、5ビットのアドレス、及びその他の制御信号によって接続され、各種のデータがインタフェース・データ・バスに接続されたレジスタに書き込まれ、また読み出される。図1にインタフェース・データ・バスとこれに接続されたレジスタを、また、表1にレジスタの一覧表を示す。レジ

スタの読み出し/書き込みは特別に1ビットずつ行うもの以外は8ビットを単位に行われ、この単位に5ビットのアドレスが付けられている。

2.2. 制御レジスタ

CSCNTRは8ビットのレジスタでCPUの制御/状態の各種情報を保持する。図2にCSCNTRの8ビットの意味を示す。CMAは、コンソールからCSMIRにセットされたマイクロ命令をアクティブにする制御ビット、BHEはDBGFフィールドでブレークを指定したマイクロ命令を実行した時、停止するかどうかを制御するビット、EFMはエラーの発生に無関係に命令を実行することを指定するビット、FERRは強制的にエラーを発生することを指定するビット、KBTはKL0のステップ実行のためのトラップを発生することを指定するビットである。また、STOPR、STOPB、STOPCはCPUが停止したときの要因を示すビットで、それぞれ、STOPRへの書き込みを指定したマイクロ命令の実行により停止したことを示すビット、DBGFフィールドでブレークを指定したマイクロ命令の実行により停止したことを示すビット、そして、エラーの発生により停止したことを示すビットである。尚、この3ビットは読み出すことはできるが書き込むことはできない。

CSCNTRの8ビットのうち、KBTは、KL0のプログラムのデバッグに使用されるもので、これを'1'にしておくと、述語呼出しの度に割込みが入る。これによって述語呼出しの流れをトレースしたり、実行を停止してチェックを行うことができ、非常に強力なデバッグ機能となっている。

また、CSCNTRの他に、CPUに対して動作を指定する1ビットの仮想的なレジスタが3つ(RESET、RUN、STEP)あり、RUN以外は書き込みだけが可能である。RESETはCPUとI/Oをリセットし、RUNはマイクロプログラムを連続的に

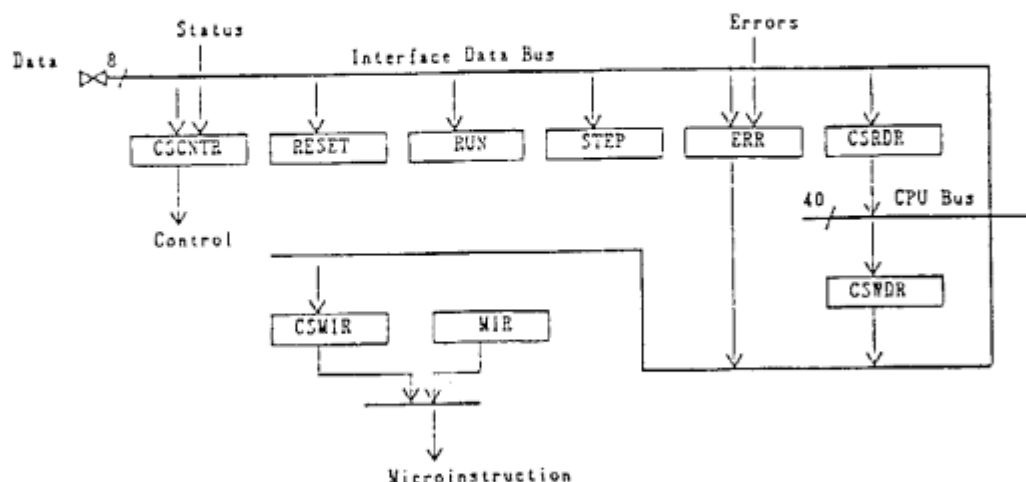


図1. CSPインタフェース

Register	# Bits	Read	Write
CCONTR	8	o	o x
RESET	1	x	o
RUN	1	o	o
STEP	1	x	o
CSNDR	40	x	o
CSNDR	40	o	x
ERR	32	o	o
CSMIR	64	x	o

* Except STOPP, STPT, and STPE Bits

表1. インタフェース・レジスタ

CMA	: Console Mode Active
BHE	: Break Halt Enable
EFM	: Error Free Mode
KBT	: KLO Break Trap
FERR	: Force Error
STOPS	: Stop due to STOPP Write
STOPE	: Stop due to Break
STOPE	: Stop due to Error

図2. CCONTRの制御/状態ビット

実行させることを指定し、STEPはマイクロ命令を1命令だけ実行することを指定する。

マイクロインタプリタの開発においては、STEPを使用して、1ステップずつマイクロ命令を実行させてデバッグを行うことができる。また、チェック・ポイントとなるマイクロ命令でブレークを指定しておき、CCONTRのERR ビットを'1'にしておくと、このチェック・ポイントを通る度に停止させることもできる。この時、レジスタの内容などをチェックした後にRUNに'1'を書き込めば、再び実行が開始される。

2. 3. その他のインタフェース・レジスタ

ERR はCPU およびメモリで発生するエラーの要因を保持する32ビットのエラー・レジスタである。CSNDR とCOND E はCPU の内部バス (32ビット・データ+8ビット・タグ) に接続された共に40ビットのレジスタで、CSNDRはCPU あるいはメモリ内のレジスタにCPU の内部バスを介してセットするデータを保持し、CSNDR はこのCPU の内部バスに出力されるデータを常時取り込んでいる。従って、CPU やメモリ内のレジスタの読み出しや書き込みにはこのCSNDR およびCOND E が用いられる。

CSMIR はコンソールからマイクロ命令をセットして実行させるための64ビットのレジスタで、これを使用すればマイクロ命令でアクセス可能なハードウェア資源を全てコンソールからアクセスすることができる。マイクロインタプリタやKLO プログラムの実行環境に、変更したりするのに有効である。

3. RAS

3. 1. RAS 装置

ψでは信頼性向上のために、ほとんどすべてのRAM と主なレジスタにはデータ8ビットに対して1ビットのパリティ

ビットが付加されており、データが転送されてレジスタにセットされたときにはパリティ・チェックが行われる。また、パリティ・ビットの付加単位である8ビットが他の8ビットとマージされたり、分解されるような論理においてはパリティのプレディクションが行われている。さらに、主記憶については40ビットのデータ・タグについて27ビットのSECDED (Single Error Correction and Double Error Detection) のECC (Error Correction Code) が付加されており、読み出し時に1ビットの誤りならば自動的に訂正し、2ビット以上の誤りならばその旨報告するようにになっている。各所で検出されたエラーは32ビットのエラー・レジスタERR に集められ、1ビットでもエラーがあるとCCONTRのSTPTビットが'1'になってCPU の動作が停止する。ERR はエラーによる停止後のエラー解析で使われる。

エラーが発生してCPU が停止した場合には、CSP がエラー発生時の状態 (ERR の値など) を読み、フレキシブル・ディスクにログ・アウツするため、後にこのディスクからエラー発生状態を知り、原因を究明することが可能になっている。

3. 2. メモリ系のエラー

アドレス変換においては、32ビットの論理アドレスが24ビットの物理アドレスに2段のテーブル変換を経て変換されるが、このとき、PMM (Page Map Memory) に保持されたページ・アドレスにはinvalid ビットが付加されており、このページを使ってアドレス変換を行った場合には割込みが発生するようになっている。これは不正ページ参照と評価される割込みで、プログラムが誤って割り付けされていないページをアクセスしようとした時に発生する。これはシステムのエラーで回復不能の重大な障害と見なされる。

アドレス変換におけるこのような不正ページ参照が発生した場合や、論理アドレスのパリティ・エラーなどその他のメモリ関係の論理にエラーが発生した場合には論理アドレスと物理アドレスを保持するアドレス・レジスタを連結し、メモリ関係のエラー解析を容易にしている。

4. まとめ

以上述べたようなシステム制御とRAS が逐次型推論マシンで実現されており、十分とはいえないまでも、ハードウェア、ファームウェア、ソフトウェアの開発においてかなり有効であった。これにより、実用的な推論マシンへの手掛かりが得られたと思われる。しかし、推論マシンは新しいタイプのマシンであり、従来の機能だけでは不十分で、論理型プログラムの高いレベルでの開発支援やデバッグ支援を十分に行うためにψでの経験を基により高度な機能の研究が必要であろう。

参考文献

- (1) 藤 他, 「パーソナル逐次型推論マシンPSI のハードウェア設計」, Proc. Logic Programming Conf., 1984, Session 8.1.
- (2) 中島 他, 「逐次型推論マシンのハードウェア—コンソール・プロセッサの構成と機能—」, 本全国大会講演論文集。

逐次型推論マシンのハードウェア

—— コンソール・プロセッサの構成と機能 ——

1E-9

中島 浩 三辺隆司
(三菱電機)渡 和男
(新世代コンピュータ技術開発機構)

1. はじめに

逐次型推論マシン(PSI)には、コンソール・プロセッサ(CSP)というマイクロプロセッサが内蔵されている。CSPは、システムの立上げ、異常処理、ファームウェア/ソフトウェアのデバッグ支援などの、 ϕ の開発や運用のための強力な支援機能を有している。

2. ハードウェア構成

Fig.1にCSPのハードウェア構成を示す。 ϕ のCPUには、Table 1に示すコンソール・インタフェース・レジスタ(CSIFR)があり、CSPはこれらのレジスタを介して、CPUの起動、停止、エラー情報の収集などを行うことができる。また、CSMIRにCPU内部のレジスタ等をアクセスするためのマイクロ命令を格納し、シングル・ステップの実行を行わせることにより、 ϕ の任意のハードウェア資源をアクセスすることができる。

なお、CSPには低速のI/O機器(マウス、キーボード、etc.)が接続されており、これらの機器のコントローラとしての機能も有している。

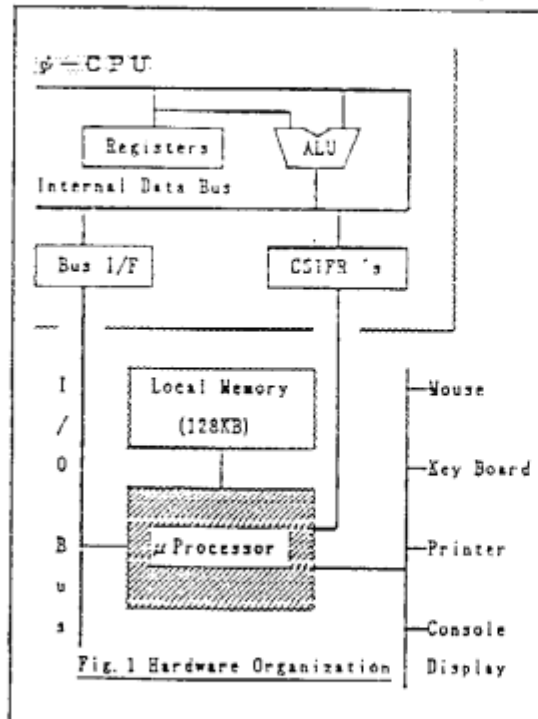


Fig. 1 Hardware Organization

Table 1 Console Interface Registers

name	function
CSCNTR	Run/Halt, Execution Mode, etc.
ERR	Hardware Error Status
CSWDR	Put Data to Data Bus
CSWDR	Get Data from Data Bus
CSMIR	μ Instruction for Resource Access

3. コンソール・システム

CSP上にインプリメントされたコンソール・システムでは、システムの立上げ、異常処理、ファイル管理、ファームウェア/ソフトウェアのデバッグ支援を行う。これらの機能は、73種に及ぶコンソール・コマンドにより会話的に実行される。以下、各機能について説明する。

3.1 システムの立上げ

Fig.2に示すように、BOOTコマンドによって、一連の立上げ手順が連続的に実行される。これらの手順は、コンソール・パネルのスイッチ設定により、電源投入と同時に自動的に実行することもできる。また、開発の段階に応じて、手順の途中までを行ったり、個々の手順のみを行ったりすることもできる。

```

$ BOOT
$$ CLEAR      Register/Memory Clear
$$ INITIALIZE  Initialize Set Up
$$ IWPL       Initial  $\mu$  Program Load
$$ IPL        Initial Program Load
$$ START      System Start
  
```

Fig. 2 System Start Up Sequence

3.2 異常処理

ハードウェアの重大な障害や、マイクロインタプリタやOSのバグによる実行不能の障害が発生した場合、 ϕ のCPUは直ちに停止する。コンソール・システムはCPUの停止状態を識別し、状況に応じてコンソール・ディスプレイに障害の内容を表示

する。また、同じ内容をファイルにログし、筆者の解析の便宜を図る。

3.3 ファイル管理

コンソール・システムでは、ファームウェア/ソフトウェアのロード・モジュールを始めとする種々のファイルを管理している。これらのファイルはフレキシブル・ディスク又は固定ディスクに格納される。コンソール・コマンドにより、開発中のプログラムのロード/セーブや、ディレクトリの表示、ファイルのコピーや削除などの、ファイルの保守も行うことができる。

3.4 デバッグ支援

ファームウェアのデバッグ支援のために、レジスタの表示/変更、連続/ステップ実行、ブレーク・ポイントの設定や解除などのコマンドが用意されている。これらの基本的なコマンドの他に、より効率的なデバッグを行うためのコマンドが用意されている。その一つがFig.3に示す、停止時のレジスタ表示である。これは、あらかじめ定義されたレジスタ（複数の定義も可）の内容を、CPUが停止するたびに自動的に表示する機能であり、コマンド入力の手間を大きく軽減することができる。また、Work File(WF)の特定のアドレスを、そのアドレスに割付けられた制御用レジスタのニモニツク(Logical Name)で参照する機能や、データ・タイプを識別するタグの値をニモニツク・コードで指定する機能などもある。

```
> SET BREAK 78B      [X]Define Break Point
> SET BREAK 13AC      [X]Define Break Point
> SET DISPLAY 0 PMAR  [X]Define Register to
> GO 1F3              [X]Display
XI-BKXPNT.BREAK POINT
PMAR = 78B           [X]Auto-Display
> SET DISPLAY 0 CLRR  [X]Current Local Base
> GO                  [X]Register (WF 16)
XI-BKXPNT.BREAK POINT
PMAR = 13AC          [X]Auto-Display
CLRR = REF(21) 0D00E53C
[REF:reference]
```

Fig.3 Firmware Debug Support

ソフトウェアのデバッグ支援のためには、機械語(KL0)レベルでのステップ実行やブレーク・ポイントの他に、主記憶や制御用レジスタの内容があらかじめ定義された値になったときに停止する、データ・ブレーク機能や、実行した述語のアドレスなどの情報をトレースする機能がある(Fig.4)。これらの機能は、KL0のマイクロインタプリタと協力して、以下に示す手順で行われる。

- (1)コマンドに応じたデバッグ支援機能の内容を、特定のアドレスに割付けられた、ファームウェア制御用レジスタ(PCB)にセットする(例えば、データ・ブレーク機能を示すコード、アドレス及びデータをセット)。
- (2)KL0ブレークのフラグをオンにし、述語呼出しのたびに、割込が発生するようにする。
- (3)マイクロインタプリタは、述語呼出しを行う際にKL0ブレークの割込を検知し、デバッグ支援用のマイクロルーチンへ分岐する。
- (4)デバッグ支援用ルーチンでは、FCRの内容を解析し、その内容に応じて、デバッグ情報の表示や実行停止などの処理を行う(例えば、データ・ブレークが発生したアドレスとデータを表示して停止する)。

```
> STEP /XL0          [X]LO Step Exec
*** ..... IRR = 000127B8 ***
> SET DBREAK 0:13A6B  [X]Define XL0 Break
> GO                  [X]Point
*** ..... IRR = 00013A6B ***
> SET DBREAK 0:20 INT 7 [X]Define Data Break
> GO
*** ..... MW(00000020) = 04 00000007
```

Fig.4 Software Debug Support

4. おわりに

コンソール・システムによる強力なデバッグ支援は、このファームウェア/ソフトウェアの開発の効率化に大きく貢献した。現在は、システムの評価を行うためのデータ収集機能の追加を行っており、性能評価や改良などに利用される予定である。

データベース・マシンの機能設計について

シフ-1

宮崎収兄, 青田健男, 柴山茂樹, 後田治夫, 村上国男

(財)新世代コンピュータ技術開発機構

1 はじめに

データベース・マシン(DBM)の開発にはその基本処理方式やハードウェア構成だけでなくDBMが全体として実現すべき機能の検討も必要である。

DBM機能の検討はホストとDBMとの間の機能および負荷の最適配分の問題として捉えられる。

2 DBMの使用形態

DBMはホストとの関係によつて2つに大別できる。

(1) バックエンド型DBM

(2) サーバ型DBM

バックエンド型DBMは図1に示すように1または少数のコンピュータのバックエンドとして用いられる。その導入の目的はホストの性能向上と処理負荷軽減にある。他方サーバ型DBMは一般に図2のようなネットワーク環境で使用され多数のコンピュータにデータベース機能を提供する。この場合の目的はデータベースの共有化・大容量化や性能向上にある。

DBMの設計にあたってはできるだけ機能をマシン側に移しホストの負荷を軽減すべきだとの意見もある。しかしサーバ型DBMでは機能をマシンに集中するとボトルネックとなる恐れもあり負荷バランスを考えた設計が必要であろう。どちらで実現しても処理量があり変化する並列処理による性能向上が期待できないものはホスト側機能とした方がよい場合もある。

DBMのトラヒック・パターンはそのアプリケーションによって変わるのでそれを考慮した機能の設計が必要である。例えば定型業務の場合 predefined request機能が必須であるし, Prolog

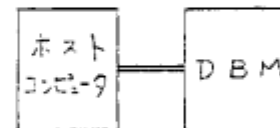


図1 バックエンド型DBM

(LAN等)

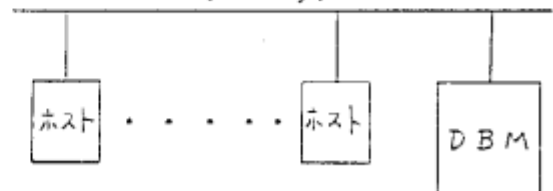


図2 サーバ型DBM

のfactを格納するにばりレ・ショナル・モデルが最適である。

3 DBMの機能

DBMの機能を論じるにはそれぞれ3層に分類すると良い。

(1) データベース・アクセス機能

(2) データベース制御機能

(3) システム管理機能

3.1 データベース・アクセス機能

ホストより指定された検索・更新などを実行するDBMの基本機能である。従って本機能の性能向上はさかんに検討されている。検討すべき機能として以下のものがある。

(1) DBアクセスのレベル (フィジカル内部スキーマ, スキーマ, サブスキーマ)

(2) データ・モデル

(3) DBアクセス言語 (コマンド)

(4) インタラクションの単位 (カプル/集合)

(5) データ独立柱の保持 (フォーマットおよびコード, 型, 単位等の変換)

(6) DBアクセス機能の実現度

・関係完備性, 条件式表現,

- ・ 算術演算, aggregate 演算
- ・ 出力のソート・ユニーク処理
- ・ 集合比較
- ・ least fixed point operation

(7) predefined request

3.2 データベース制御機能

データベース制御機能はデータベースアクセスを制御する際のデータ保護などの機能である。DBMS として求められる機能は必須であるが何などのようを実現すべきかは、マシンのアクセス機能ほどには検討されていない。また、実現される場合でも通常の DBMS の技術の応用にとどま、ている場合も多い。

- (1) 複数トランザクションのコンカレンシ制御
- (2) データリカバリ (ロールバック, バックアップ, ロールフォワード)
- (3) セキュリティ管理 (アクセス権, 暗号など)
- (4) インテグリティ管理
- (5) ディクショナリ管理 (DCL プロセスの提供, コンカレンシ制御, データのコンシステンシ維持)
- (6) 複数 DB, 共有 DB, 個人 DB
- (7) DB administratively のサポート
 - ・ バック・ロードニング
 - ・ DB 再構成, フトレ・ミ管理
- (8) データベース設計支援

3.3 システム制御機能

システム制御機能は派用システム DBMS とは全く OS の機能である。バックエンドに DBMS はある程度ホストに分離できることもできるがサーバ型 DBMS は重要な機能となる。

- (1) コンソール・インタフェース
- (2) システム状態制御
- (3) ホストからの状態監視, 実行中制御指示
- (4) ユーザ登録, 管理
- (5) 課金処理
- (6) システム制御パラメータの変更
- (7) ログレ・ブ等の保存

(8) 障害検出と回復, 系の再構成

- (9) 故障診断
- (10) 機能増設, 機能追加
- (11) 簡便なデータ収集

などが検討課題となる。

4. DBMS 機能の設計

近年「機能的に完全な DBMS」の研究がさかんになってきた。これは基礎研究だけでなく実用化にも重点がおかれてきたことを示唆している。しかし「機能的に完全な DBMS」がどのようなものかについてあまり議論されていない。この作業仮説としてきちで述べた全項目を高度に実現したマシンを「機能的に完全な DBMS」とみなし、アクセス処理と実行する各種基本方式のもとでその実現方法を検討することは基本方式の比較や DBMS 機能の決定を行うための有効な手段であろう。

データベース処理におけるデータベース制御処理のための比率はかなり高い。DBMS におけるデータベース制御機能の高性能化が検討されるようになってきたがアクセス処理の基本方式やハードウェア構成との関係はまだ充分解明されていない。DBMS の機能の設計においてはこれらの点の検討が重要である。また以下の点も考慮せねばならない。

- (1) DBMS の用途・アプリケーションによる機能の必要性
- (2) DBMS の使用形態, ホストの種別
- (3) 全機能の DBMS およびホストでの実現方法による性能とコスト比較
- (4) 将来の機能追加・拡張

DBMS のように新しいシステムでは未解決の問題もあり (4) は重要である。

5. まとめ

DBMS 機能の設計における検討課題と留意点を考察した。機能検討にあたってはホストを含めた全体系の中での機能分担の視点が重要である。

Unification in A Knowledge Base Machine

3F-2

Haruo Yokota, Shigeki Shibayama, Nobuyoshi Miyazaki,
Takeo Kakuta, Kunio Murakami
Institute for New Generation Computer Technology (ICOT)

1. Introduction

It is important for the knowledge information processing system to manage large amounts knowledge efficiently. When we implement such a system using a logic programming language based on Prolog, collections of Horn-clauses can be viewed as a kind of knowledge. If the volume of knowledge exceeds the capacity of main memory, the current version of Prolog processor cannot handle the load. We must consider how to deal with knowledge by using the secondary memory effectively.

Horn-clauses are classified into three types: unit clauses without variables, unit clauses with variables, and non-unit clauses [1]. It is known that a unit clause without variables bears a one-to-one correspondence to a tuple in a relational database. We plan to construct a knowledge base machine as an enhanced relational database machine. We are now developing relational database machine Delta which manipulates streams between the hierarchical memory system and the relational database engine. The knowledge base machine then handles Horn-clauses by manipulating the streams.

We have presented a draft for putting three types of Horn-clauses into the knowledge base in a tuple style [1]. In order to retrieve all types of clauses, we proposed using the unification command as a knowledge base operation; it is treated as an enhanced relational algebra command. In this paper, we further investigate that approach.

2. Scheme for Storing Horn-clause

We group clauses by name and arity (number of arguments) of their positive literal. Each group forms a table. The name of the positive literal in the grouped clauses becomes the name of the table. A tuple in the table corresponds to a clause and has a number of attributes. Clause arguments are converted to attributes on a one-to-one basis. Two special attributes are provided for the bodies of clauses and for the most general unifiers generated by the unification commands. If a clause has no body, the corresponding attribute is empty.

Example 1: We use the DEC-10 Prolog syntax here. If there are four clauses, such as:

```
p(abc,12).
p(efg,22).
p(hij,X) :- q([klm,X],Y),r(Y).
p([X,Y],11) :- q([X],stt(klm,Z,Z)), Y \= Z.
```

the following table, named p, is formed.

arg1	arg2	body of clause	most general unifier
abc	12	-	-
efg	22	-	-
hij	X	$q([klm,X],Y),r(Y)$	-
[X,Y]	11	$q([X],stt(klm,Z,Z)), Y \neq Z$	-

When a query $?-p(A,12)$ is made, the knowledge base machine searches for unifiable tuples and generates the most general unifiers for the clauses found, as follows:

arg1	arg2	body of clause	most general unifier
abc	12	-	{A/abc}
hij	X	$q([klm,X],Y),r(Y)$	{A/hij,12/X}

From the latter table, one answer, $A=abc$, is obtained and a new query, $?-q([klm,12],Y),r(Y)$, is created. []

3. Argument Forms

An argument of a Prolog-clause is formed by either an atom, a variable, or a structure. We distinguish each type using tags. An atom is merely a character string with a tag. A variable can be represented using a sequential number with a tag. Since the current version of Prolog does not support global variables, a variable is used to relate the argument to other arguments in the clause.

The Prolog processor does not place structures in the same location at which arguments are stored. There are pointers to indicate the location at which the structures are actually stored. It is for this reason that structures are treated dynamically in Prolog. For instance, a variable can be unified with any kind of structure.

However, when unifiable clauses are searched for, instantiation is not necessary. It is only needed to verify whether the variables are unifiable with the structure. Moreover, if we use streams to retrieve clauses, pointers are not suitable. Thus, in the knowledge base machine, structures may be placed at the same location at which the arguments are stored. Structures are represented by functor name, argument count, and a list of the arguments. Notice that variable-length records must be supported for storing structures.

Example 2: Arguments appearing in Example 1 are represented as follows:

```
atom:
abc      => [a|abc]
22       => [a|22]
variable:
X        => [v| ]
structure:
[klm,X]  => [s| [2|a|klm|s| [2|v|1|a| ] ] ]
[N,Y]    => [s| [2|v|1|s| [2|v|2|a| ] ] ]
sit(klm,Z,Z) => [s| [3|a|klm|v|1|v|1| ] ]
```

Here, [a,b] stands for (a,(b,[])). []

4. Operational Flow

When a knowledge base query, which is a Prolog goal sequence, is made, one goal is selected from the sequence and unifiable clauses are searched for in the table having the same name as the goal. A new goal sequence can be created from each unifiable clause and the most general unifier for that clause.

If a unifiable clause has no body, the rest of the sequence to which the most general unifier has been applied becomes a new goal sequence. Otherwise, i.e., if the unifiable clause has a body, the most general unifier is applied to both the body and the rest of the sequence. A concatenation of these two instances becomes one of the new goal sequences. Emptiness of the new goal sequence indicates that one of the answers to the query has been obtained.

Concurrently obtaining an entire new goal sequences for the query implies the use of OR-parallelism. However, further new goal sequences must be generated concurrently for each obtained goal sequence. The number of goals that must be handled concurrently increases rapidly. This phenomenon is called OR-parallelism explosion.

There are three types of optimization for sequentially generating new goal sequences. First, we allow processes to be invoked in the inference

mechanism, especially recursive call processes, which access only main memory. Consequently, we partition knowledge into two types, according to its volume and contents. Small quantities of iterative knowledge are put into the inference mechanism. Large quantities of knowledge are stored in the knowledge base mechanism. To take this approach, some interface between the inference mechanism and knowledge base mechanism is needed. In [2], we propose a deferred evaluation method for interfacing Prolog with relational databases. We can improve this method to interface these two mechanisms without massive modifications.

Second, we can group processes that access the same table. A head literal in the body of the unifiable clause indicates a table that will be accessed in the next stage. We can collect clauses that have the same head literal in each body by going through the clause-body attribute of the unifiable clauses table being searched.

Last, we can control the order of clause selection according to the literal count in the clause body. The length of the body often indicates the complexity of the resolution process when the process does not contain a recursive call. The count of the negative literals can be used as a kind of evaluation function for priority control. A unifiable clause that has a shorter body is invoked prior to those having longer bodies.

5. Conclusion

We have presented a method that stores any type of Horn-clauses in the knowledge base. We have shown that atoms, variables and structures can be handled by the knowledge base. It is difficult to implement OR-parallelism in the knowledge base mechanism. For sequential execution, we have proposed three types of optimization.

Acknowledgement

The authors thank the members of ICOT Working Group 1 (KBM subgroup) for their valuable discussion.

References

1. H. Yokota, et al., An Investigation for Building A Knowledge Base Machine, *Proceeding of the 27th IPSJ Conference*, 2K-7 (Nagoya, October, 1983). [in Japanese, English version to appear as ICOT Technical Memorandum]
2. H. Yokota, et al., An Enhanced Inference Mechanism for Generating Relational Algebra Queries, *Proceeding of the 3rd ACM SIGACT-SIGMOD Symposium on Principles of Database Systems*, pp.229-238 (Waterloo, April, 1984).

Relational Database Processing on an Attribute-Based Schema

3 F - 3

Shigeki Shibayama, Takeo Kakuta, Nobuyoshi Miyazaki
Haruo Yokota, Kunio Murakami

Research Center
Institute for New Generation Computer Technology (ICOT)

1. Introduction

ICOT's Relational database machine (Delta) has adopted an attribute-based database schema for storing relations in secondary storage devices (moving-head-disks). The reasons we selected the attribute-based schema are: a) predefined access paths are not available, b) unnecessary accesses to disk should be minimized, and c) tuples are expected to be accessed evenly. In this paper, we will investigate the properties associated with this schema.

2. Attribute-based Schema and Tuple-based Schema

The differences between attribute-based schemas and tuple-based schemas are as follows:

- (1) Tuple reconstruction time: attribute-based schemas must reconstruct corresponding attributes for output or tuple-oriented operations such as union with duplicate elimination.
- (2) External storage access cost (both seek + latency time and data transfer time): tuple-based schemas must access extra (unnecessary) attributes in tuples not used in the query.
- (3) Clustering of tuples/attributes: there are a number of methods of storing tuples/attributes in clusters and provide access to them. This is closely related to how symmetrically the access methods can handle the attributes of a relation.
- (4) Space efficiency: an attribute-based schema usually requires tuple-identifiers (tids) to indicate tuple relationships among attributes. This requires additional storage.
- (5) Temporary relation representation: in a tuple-based schema, usually, temporary relations must actually be made. With an attribute-based schema, temporary relations exist in concept only; however, efficient internal representation of temporary relations is possible.

3. An Attribute-Based Schema with Two-Level Clustering

In Delta, we have chosen an attribute-based schema with two-level clustering [1]. An example of this clustering is shown in Fig. 1. This schema is devised to match the requirements for a database machine used within a logic programming research environment. In particular, the symmetrical attribute access property is thought to be essential. In Prolog, a representative logic programming language, facts (unit clauses with no variables) directly correspond to tuples in a relational database. Prolog facts do not explicitly incorporate the concept of keys. We believe that unequal processing times for manipulating various attributes is an undesirable feature for

database machine users. Therefore, we decided to handle all the attributes symmetrically. The appropriateness of this specific schema remains to be evaluated.

4. Hardware and Software Requirements for Attribute-Based Schema

There are several hardware and software requirements for adopting an attribute-based schema in a database machine.

- 1) Fast join capability: generally, the process of performing relational database operations with an attribute-based schema consists of, a) producing a tid list denoting result relations, and b) generating result relations from the tid list. As the attributes used in the query are stored separately with their corresponding tids, tid-restrictions to the tid list are necessary for all the attributes. This is the same, from a processing point of view, as equi-joins between attributes. If equi-joins are performed slowly, overall performance will not be satisfactory.
- 2) Fast tuple/attribute transposition capability: tid-restriction is the first stage of generating result relations. After the tid-restrictions are complete, attributes (tid-restricted and sorted with the tids) must be transposed with tuples.
- 3) Cluster-formation capability: this is specific to the two-level clustering schema of Delta. After the transposition of attributes (in the form of attribute buffers), each attribute must go through the clustering process. This is done by, a) sorting each attribute with the attribute value, b) dividing the value-sorted buffers into a certain number of smaller sections (zones) according to value ranges, c) sorting each zone with the tid-values, and d) dividing each zone into a number of smaller sections (pages) according to the tid value ranges and creating a directory for the clusters.

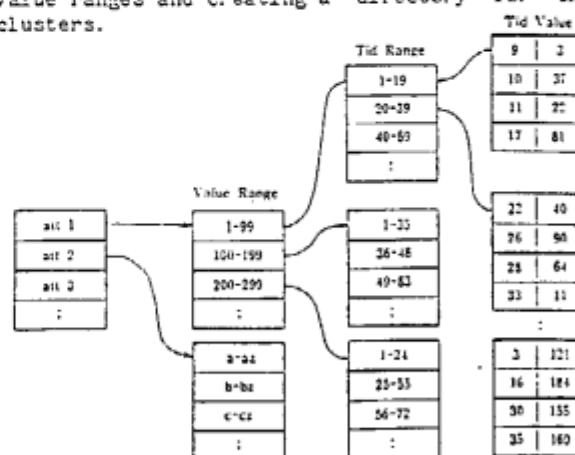


Fig. 1 Delta's Attribute-Based Schema

8. An Example of Relational Database Command Interpretation using the Attribute-Based Schema

Consider the following SQL-like query:

[Sample Relations]

```
icot(name, age, company, laboratory, group)
company(company_name, location)
```

[Sample query]

```
select name, laboratory, group
from icot
where company = next
select company_name
from company
where location = [yokohama]
```

The following lists the processes performed in response to this query. Shared item numbers indicates source-result relationships. (Increasing primes indicate successive source-result relationships.)

- (1) prepare comp-location attribute
- (1)' tid(c):comp-location (= [yokohama])
- (1)' tid(c):comp-location (sort by tid(c))
- (2) prepare comp-name (= tid(c) of (1)')
- (2)' tid(c):comp-name (= tid(c) of (1)')
- (3) prepare icot-company (= comp-name of (2)')
- (2)', (3) → (3)' new-tid:tid(c):tid(i)
- (natural-join by comp-name and icot-company)
- (3)' new-tid:tid(c):tid(i) (sort by tid(i))
- (4) icot-name (= tid(i) of (3)')
- (4)' new-tid:icot-name (tid-restriction by tid(i))
- (5) icot-lab (= tid(i) of (3)')
- (5)' new-tid:icot-lab (tid-restriction)
- (6) icot-group (= tid(i) of (3)')
- (6)' new-tid:icot-group (tid-restriction)
- (4)', (5)', (6)' → (7) new-tid:icot-name (sort by new-tid)
- (5)', (6)' → (7) new-tid:icot-lab (sort by new-tid)
- (6)', (7)' → (7) new-tid:icot-group (sort by new-tid)
- (4)', (5)', (6)' → (7) icot-name:lab:group
- (transposing to tuples)

here, tid(c), tid(i) and new-tid denote tid of company relation, icot relation and generated new tid, respectively. For example, new-tid:icot-name denotes a buffer with a data structure of new-tid and icot-name attribute.

As can be seen from this example, attribute-based schemas carry out essential database query operations using concise representations. Queries that frequently handle large temporary relations are performed efficiently with this schema. The most time-consuming factor of this schema is the actual formation of result relations. This requires, a) readout of output attributes from a disk device (with disk cache), b) tid-restriction of the output attributes and, c) transposition of attributes with tuples. In Delta, a large semiconductor disk (about 120MB) is incorporated for disk cache use to improve the hit ratio and for buffer use to ensure that almost all the intermediate results in a query remain in high-speed storage. Frequent operations with this schema, such as tid-restriction and join, are performed using special-purpose hardware called the Relational Database Engine (RDBE) [2].

The relative performance is compared to that of result tuple selectivity in Fig. 2. Estimation model is a selection put in SQL as:

```
select A1, A2, ..., An
from A
where A1 (cmp) C(criterion)
```

The following conditions are assumed:

- (1) a specific page can be identified by either its tid or value
- (2) tid-restricted attributes are read from disk using a secondary index
- (3) average seek and latency time are required for each page access
- (4) tid-restriction requires buffer memory read and write
- (5) transposition requires 10 buffer read/write operations
- (6) there is constant overhead time (accumulation of overhead associated with each operation)

The estimate shows that, a) a full scan of tid-restricted attributes is required for a selectivity range larger than 1 %, b) tuple reconstruction time affects performance for selectivity ranges larger than 10 %, c) tid-restriction time is short compared with disk access time, and d) overall processing time increases gradually. Evaluation of the join operation will show similar results, because most of the processing time consists of tid-restricted attributes access and transposition time.

6. Conclusion

Attribute-based schemas can exploit efficient intermediate representations of temporary relations. The processing cost lies in read-out of tid-restricted attributes. Tuple reconstruction time is not very large in low selectivity ranges.

[References]

- [1] Shibayama, S., et al., "A Relational Database Machine with Large Semiconductor Disk and Hardware Relational Algebra Processor," New Generation Computing, Omsha-Springer, Vol.2, No.2, pp.131 - 135.
- [2] Oka, T., et al., "Development of a Relational Database Engine (1) to (6)," Proc. 29th IPSJ Conference, 4F-5 to 4F-10.

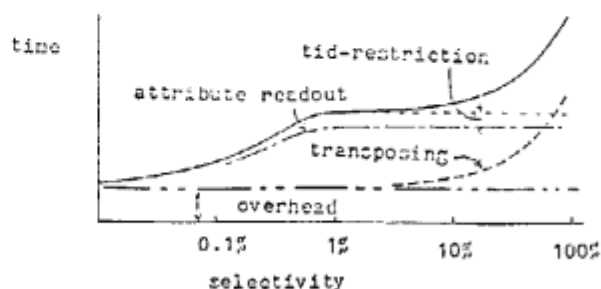


Fig. 2 Performance Estimation

RDBM Delta のリカバリ方式

3F-4

角田健男, 宮崎収兄, 柴山茂樹, 横田治夫, 村上国男

(財) 新世代コンピュータ技術開発機構

1. はじめに

ICOT で開発中の関係データベースマシン (RDBM) Delta は、ホストマシンと多数台のパーソナル逐次型推論マシン (PSI) と LAN を介して接続され、本プロジェクトの中期研究開発課題の一つである知識ベースマシン開発のための基礎的実験を行うこと、および PSI 上のユーザプログラムからの外部データベースアクセス要求を効率良く処理することを主目標としている。外部データベースには英和・和英辞書などの大容量共用データや各ユーザ専用のデータベースなどが格納され、検索・更新処理の対象となるため各種障害発生時のデータベースのリカバリ処理を考慮し実現する必要がある。本稿では、Delta のリカバリ方式の概要について報告する。

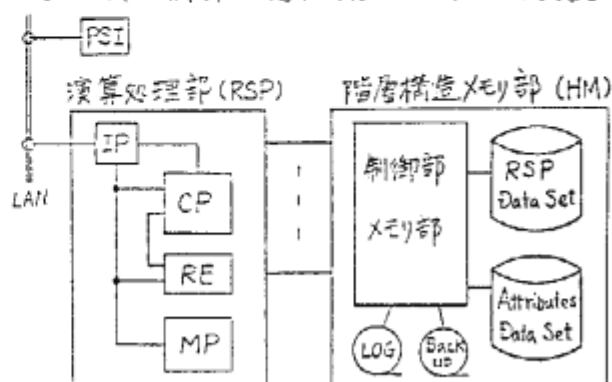
2. Delta のデータ格納

Delta では、リレーションを構成するタプルをアトリビュートに分割し、アトリビュート毎に格納する縦型格納法⁽¹⁾を採用した。一方、リレーション管理情報には、ホストが参照するディクショナリと Delta の Control Processor (CP) が Delta コマンド解析・実行やリカバリ処理に使用するディレクトリがある。階層構造メモリ部 (HM) のメモリ管理上、アトリビュート情報とディクショナリ情報はアトリビュート・データセットとして、またディレクトリ情報は RSP データセットとして取扱う。

Delta システム構成を図 1 に示す。

Delta のリカバリ処理は、Delta コマンドの実行と制御する CP と Delta システムの状態管理、オペレータインタフェースを処理する Maintenance Processor

(MP) との制御下で HM のデータセットを、障害発生時には実行中のトランザクション開始前の状態に復旧させることである。



IP: Interface Processor, RE: RDBM Engine
RSP: RDBM Supervisory and Processing subsystem
HM: Hierarchical Memory subsystem

図1. Deltaのシステム構成

3. Delta のリカバリ処理基本方針

- (1) 更新処理を含むトランザクションの開始後、トランザクション凍結前まではホストからのアボート指示または Delta 内で検出した Dead Lock やリソース不足等のもとでトランザクション開始時点の状態を戻し、実行中の更新処理を全て取消す。
- (2) Delta システムダウン時: システム再起上げ後、オペレータによるシステム状態チェックを行い、コミット/ロールバック指示に基づいてリカバリ処理を行う。
- (3) 媒体系の障害時: 人間による原因調査、装置構成変更等の後、バックアップ情報とログ情報によりロールフォワードし、次にロールバックを行う。
- (4) 電源障害時: 無停電電源装置と使用し電源ダウン時のデータベース破壊を防止すると共に処理再開できる環境を確立後システムダウンさせる。その後の処理は(2)と同様である。

ニコンカレンシイ制

Deltaでは、2フェーズ・ロック方式[1]に基づくコンカレンシイ制御をCPで行い、最大64個までのコマンド・トランザクションを同時動作させる。

- (1) ロックの単位：リレーション
- (2) ロックの期間：トランザクションの開始から凍結または終了まで。
- (3) ロックの種類
 - ・ READ ロック (シェアラブル)
 - ・ WRITE ロック (エクスクルーシブ)
- (4) ロックの相互関係

ロック 取得	READ	WRITE
READ	○	X
WRITE	X	X

○：許可
X：禁止

5. Deltaのリカバリ方式

Deltaのデータベース・リカバリ処理方法は、以下の2ケースによる。

- ① オンライン・ロールバック
- ② オフライン・ロールバック

(1) リカバリ用サブコマンド

CPで管理するディレクトリには、リレーションおよびアトリビュートスキーマ情報があり、実行中のトランザクションID、ロック中のリレーションIDとリカバリ用は上記スキーマ情報のシャドウ情報をトランザクション毎に持つ。一方、HMでは各アトリビュートID毎にアトリビュート情報を管理し、リカバリ系はアトリビュート情報のシャドウを持つ。

CPは、コマンド・トランザクションの同時実行のためトランザクションID毎にDeltaコマンドの解析の後各サブコマンドに分解生成し、HMとREを起動し、その実行を管理する。

CPからHMへ送出される復元、更新系サブコマンドには、トランザクションIDやアトリビュートIDなどのパラメータが含まれている。

リカバリ系サブコマンドには次の2つがあり、これもCPからHMに送出される。

- (a) Commit-attributes サブコマンド
ホサブコマンドで指定のアトリビュートID群とトランザクションIDが更新処理実行中のトランザクションIDとアトリビュートID群と一致するとき、リカバリ対象と判定しそのアトリビュート情報をコミットする。

- (b) Rollback-attributes サブコマンド
パラメータ判定によりリカバリ対象であるときは、指定アトリビュートとロールバックする。

(2) オンライン・ロールバック

更新処理中に上記リカバリ系サブコマンドを発行したときのディレクトリとアトリビュート情報の状態遷移を下表に示す。

HM サブコマンド	ディレクトリ		アトリビュート情報		備考
	ICエントリ	シャドウ	IDエントリ	シャドウ	
更新系 ↓ Commit-att.	(更新) ↓ 新エントリ	旧エントリ ↓ 廃棄	(更新) ↓ 新エントリ	旧エントリ ↓ 廃棄	
更新系 ↓ Rollback-att.	(更新) ↓ 廃棄	旧エントリ ↓ 旧エントリ	(更新) ↓ 廃棄	旧エントリ ↓ 旧エントリ	トランザクション アボート (オンライン)

(3) オフライン・ロールバック

3の(2)～(4)項に記す通りリカバリ処理を行う。

6. おわりに

Deltaのリカバリ処理を実現するための障害の分類と処理手順について更に詳細検討を進め、Deltaのリカバリ機能を充実させてゆく予定である。

今後Deltaの開発に尽力いただいている(株)東芝および(株)日立製作所の各位に感謝致します。

参考文献

- (1) 宮崎他, データベースマシンにおけるデータ格納法に関する一考察, 27 情報大全 (1983.10)
- (2) Eswaran et al., The Notions of Consistency and Predicate Locks in a Database System, Com. of the ACM, No.11, Vol.19, 1976

SIMPOSのオペレーティング・システム概要

飯部 隆 辻順一郎 内田俊一 横井俊夫
((財) ICOT)

1. はじめに

SIMPOS (Sequential Inference Machine Programming and Operating System) は逐次型推論マシンPSIのためのプログラミング/オペレーティング・システムである。PSIはICOTで開発されたソフトウェア開発用パイロット・モデルであり、SIMPOSはそのために新規に設計、開発されている。本稿では、SIMPOSのオペレーティング・システムの構築概念とシステム構成を述べる。

PSIのノリセリとメイン・メモリは、後言語第0版(KLO)と呼ばれるPrologベースの機械語を効率良く実行するように設計されている。また、入出力装置として、ハード・ディスク、フレキシブル・ディスク、ビット・マップ・ディスプレイ、オブティカル・マウス、およびローカル・エリア・ネットワーク(LAN)インタフェースなどを装備している。

SIMPOSの主目標は、PSIの特質や装置を効果的に使用して、PSIを良好な論理型プログラミング・ワークステーションとすることである。また、その設計思想は機能の豊富さよりも簡潔性と拡張性を重視している。これは、パイロット・モデルとして今後の種々の試みを容易に実現できるようにするためである。

オペレーティング・システムはプログラミング・システムの下にあって、マルチ・プロセッシング、マルチ・ウィンドウ、ディスク・ファイル、およびネットワーク通信などをサポートしている。プログラミング・システムは、テキストプログラム・エディタ、デバッガ/インタプリタ、ライブラリ/コンパイラなどを含み、それらとユーザとの会話はコーディネータにより管理される。

2. 構築概念

SIMPOSではPrologにモジュール性と副作用を導入する基本的手段としてオブジェクト指向的アプローチを選んだ。オブジェクトは、外部的にはそのオブジェクトに許される(つまり、オブジェクトが受けつける)操作の集りにより定義され、内部的には操作を記述するクローズおよび開のオブジェクトを保持するスロットで定義される。同じ定義をもつオブジェクトはクラスとして分類され、それらのテンプレートはクラス定義で与えられ、各オブジェクトはこのテンプレートから生成される。SIM

POSのシステム記述言語であるESPでは、クラス定義は以下のシンタックスで与えられる。

```
class クラス名
    [ マクロ・バンク宣言 ]
has
    [ 継承定義 : ]
    [ クラス・スロット定義 : ]
    [ クラス・クローズ定義 : ]
[ instance
    [ インスタンス・スロット定義 : ]
    [ インスタンス・クローズ定義 : ] ]
[ local
    [ ローカル・クローズ定義 : ] ]
end.
```

クラスはいくつかのクラスを多重継承して定義できる。そのクラスは外部的に、自ら定義した操作に加えて、継承したクラスの操作をすべてもつことになる。内部的には、これらのクラスにある同じ操作のクローズがORされ、スロットが集められる。また、デモン述語をクラスおよびインスタンス操作に定義できる。これはプライマリ述語の前か後にANDされ、暗黙的に呼出される。これらの機構はオペレーティング・システムの記述に有効である。

3. システム構成

オペレーティング・システムの機能は、最小限の機能をもつ、限られた数の基本クラスと、これらを継承した付加機能をもつクラスとして提供され、アプリケーション・プログラムはこれらクラスのオブジェクトを介してその機能を利用する。オペレーティング・システムでは、これらのシステム・クラスを入出力メディア・システム、実行管理(スーパーバイザ)、および資源管理(カーネル)の3層に分割する。

3.1 資源管理

カーネルはprocessor やareaで表現されるハードウェア資源を管理し、diskやkeyboardなどで表現される物理入出力装置を制御する。なお、プロセッサやメイン・メモリはファームウェアで基本的に管理されている。

また、カーネルは IPL (Initial Program Loader) や スタンド・アロン・デバッガであるシステム・トレーサを含む。

3.2 実行管理

スーパーバイザは SIMPOS における実行モデルをサポートする。そのためのクラスとして、process (実行主体)、program (処理手続)、stream (プロセス間通信)、pool (オブジェクトの収納)、world (実行環境) などを含む。このモデルでは、プロセスがワールドの下でプログラムを実行することでオブジェクトを操作し、また必要ならば、ストリームを介して他のプロセスと通信する。

なお、プログラムはクラス as-program を継承したユーザ定義のクラスとして与えられる。

3.3 入出力メディア・システム

入出力メディア・システムは、window (ユーザと会話的に通信する)、file (ディスクにデータを格納、検索する)、cable (ネットワークを介して他のマシンと通信する) などのクラスで論理入出力機能をサポートする。

(1) ウィンドウ・サブシステム

ウィンドウ・サブシステムは、ビット・マップ・ディスプレイ (スクリーン)、キーボード、およびマウスを制御して、複数の会話セッションをユーザに提供し、並行していくつかの仕事をできるようにする。この機能は、マルチ・ウィンドウと呼ばれ、ワークステーションでは必須の機能となりつつある。SIMPOS では、多手継承の機構を利用して、種々のタイプのウィンドウをユーザが容易に定義できるようにしている。このため、ウィンドウ・サブシステムはウィンドウ構成用クラスを提供している。

また、ウィンドウ・マネージャにより、ユーザはマウスを使ってスクリーン上のウィンドウを操作することができる。

(2) ファイル・サブシステム

ファイル・サブシステムは、ハード・ディスクおよびフレキシブル・ディスクを管理し、データの格納や検索をサポートする。SIMPOS で扱うファイルは、一般のデータを対象とするデータ・ファイルと、オブジェクトを対象とするインスタンス・ファイルとに分類される。

データ・ファイルには、バイナリ・ファイル、固定長レコードを格納するテーブル・ファイル、可変長レコードを格納するヒープ・ファイルがある。アクセス法として、ダイレクト・アクセス (ファイルに対し) およびシークエンシャル・アクセス (ファイル・タプに対し) がサポートされる。

インスタンス・ファイルには、保存したいオブジェクトを格納し、後でそのオブジェクトを取出して再生できる。例えば、ファイルやユーザといったオブジェクトがインスタンス・ファイルに格納される。さらに、ディレクトリ・ファイルを使って、これらのオブジェクトに名前を付けて管理できる。

また、ユーザはファイル・マネージャにより、ウィンドウを利用したファイルの操作ができる予定である。

(3) ネットワーク・サブシステム

ネットワーク・サブシステムは、LAN で接続されたマシン (PSI や DELTA) 間でのデータ通信を実現する。SIMPOS がサポートするネットワーク通信は3つのレベルに分類される。まず、マシン間通信のレベルは、異種マシン間の通信に使われ、データの入出力として扱われる。次に、プロセス間通信のレベルでは、異なる PSI 上のプロセスの間でチャネルを介したメッセージ通信を可能にする。最後に、リモート操作のレベルでは、他の PSI 上のオブジェクト (リモート・オブジェクト) を操作する機能を実現するためのサポートを提供する。

4. おわりに

SIMPOS は、新しいマシン上の新しいシステムであり、様々な試みがなされているが、限られた環境下で開発できたことについては、正しい選択であったと考えられる。また、利用上の評価はこれからユーザがアプリケーション・プログラムを開発し始めたときになされるであろうが、そこからのフィードバックにより SIMPOS をよりよいものにしていく予定である。

なお、今回の SIMPOS オペレーティング・システムの紹介では、以下の順でその各部分を説明する。

- ・オペレーティング・システム概要 (本稿)
- ・資源管理
- ・実行管理 — プロセスとストリーム
- ・実行管理 — プール
- ・実行管理 — ワールド
- ・ウィンドウ・サブシステム
- ・ネットワーク・サブシステム
- ・ファイル・サブシステム
- ・IPL 方式
- ・システム・トレーサ

参考文献

I. Hattori, et al., "An Operating System for Sequential Inference Machine PSI", ICOT TH-0061 (日本語版 TH-0065).

SIMPOSの資源管理

川上 孝仁
(三菱電機㈱)上田 尚純
(三菱電機㈱)堀江 建彦
(三菱電機㈱)設部 隆
(㈱ICOT)

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の一環としてICOTを中心に、知識情報処理ソフトウェアの開発用ツールとして、逐次型推論マシンφ (PSI) とそのプログラミング/オペレーティング・システム (SIMPOS) を開発中である。本論文では、SIMPOSのうち、その最下層に位置し、φのハードウェア資源 (プロセッサ、メモリ、入出力デバイス) を管理する資源管理 (カーネルとも呼ばれる) について報告する。

2. 資源管理の概要

資源管理の目的は、つぎの3つである。

- ハードウェア資源をオブジェクトとして記述し、SIMPOSの基本構成概念であるクラスの枠組にのせること。
- 実行管理プログラム (スーパーバイザ) およびファームウェアと協力して、63個までのプロセスが並行動作できるマルチ・プロセス環境を提供すること。
- デバイス仕様から独立した高い論理レベル・インタフェースの入出力実行機能を提供すること。

φのハードウェア資源としては、

- (i) プロセッサ、プロセス・コントロール・ブロック
- (ii) エリア、ページ・マップ・メモリ、物理ページ
- (iii) 各種入出力デバイス

があり、これらを管理するソフトウェアをそれぞれ、プロセッサ管理、メモリ管理、デバイス管理と呼ぶ (図1)。

3. プロセッサ管理

プロセッサ管理は、2つのクラスprocessor, context からなる。クラスprocessorは、φのCPUに対応したクラスであり、インスタンスは1個である。割込インデックス・テーブル、割込原因コード・テーブル、CPU内シ

テム・レジスタへのアクセス機能および、CPUを制御する各種特権組込述語の実行機能を実現する。

プロセス実行に必要なハードウェア資源として、つぎのものがある。

- (i) 1組のPCB (Process Control Block) / EPCB (Extended PCB)
- (ii) 4組のACB (Area Control Block)
- (iii) ページ・マップ・メモリ領域

これらを総称してハードウェア・コンテキストと呼ぶ (図2)。クラスcontextは、このハードウェア・コ

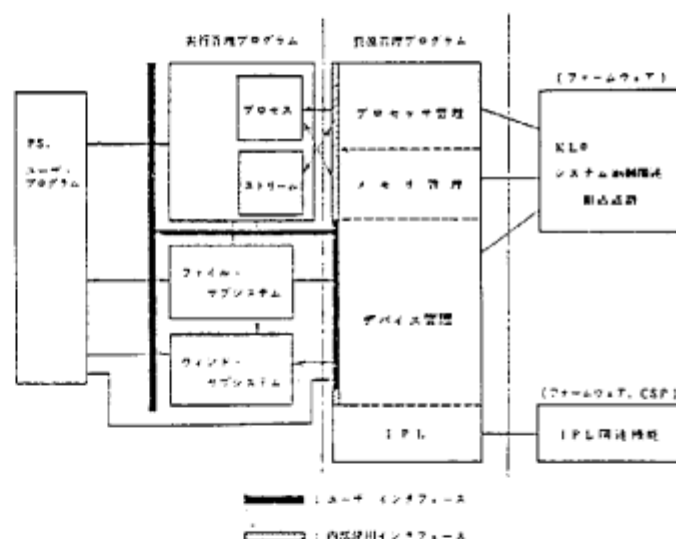


図1. 資源管理プログラムの機能分担

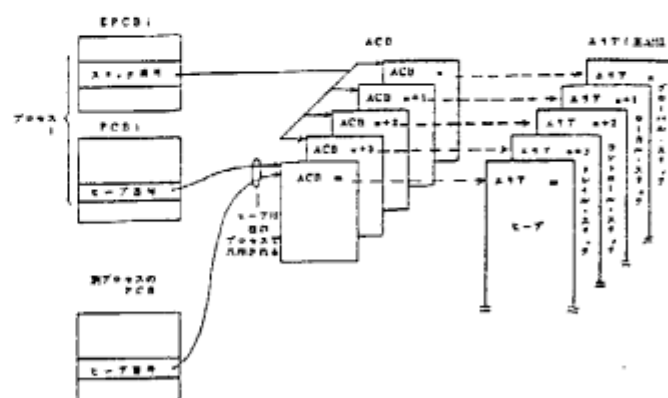


図2. ハードウェア・コンテキスト

ンテクトに、ソフトウェアでプロセスを管理するためのいくつかの情報を付加したものである。インスタンスは、ハードウェアのPCBの数と同じく63個、存在する。

4. メモリ管理

前述のように、プロセスを実行するとき、4個のスタックに対応した4個のエリア(ACB)と、それぞれのエリア毎に、そのページ数分の大きさのページ・マップ・メモリ領域を使用する。ページ・マップ・メモリとは、すべてのエリアの論理ページ・アドレスを物理ページ・アドレスに変換する変換テーブル(ハードウェア・レジスタ)である。プロセスの実行に伴って、ヒープ・エリアの伸長と、スタック・エリアの伸縮があるが、エリアを拡張するとき、ページ・マップ・メモリの確保と、物理ページの確保、および、物理ページ・アドレスのページ・マップ・メモリへの書き込みは、ファームウェアが行う。すなわち、ページ・マップ・メモリ管理の一部と、物理ページ管理は、ファームウェアが分担している。

ソフトウェアとしてのメモリ管理は2つのクラスarea, page-map からなる。ハードウェアのACBに対応するクラスareaには、インスタンスが256個あり、各インスタンスは、ヒープ、4種類のスタックのいずれかとして用いられる。クラスpage-mapはページ・マップ・メモリに対応し、ページ・マップ・メモリの領域管理を行う。

上層の実行管理プログラムがプロセスを生成するときには、プロセッサ管理、メモリ管理を用いて、つぎのように行う。

- (1) contextのひとつを確保する。
- (2) 4個のareaを確保する。
- (3) それぞれのareaに対し、適当な大きさのpage-mapを割付ける。
- (4) page-mapの割付けられた4個のareaをcontextにセットする。

5. デバイス管理

デバイス管理は、入出力デバイスに対応した下記のクラス、

- ・ disk-with-lrucache
- ・ flexible-disk
- ・ bitmap-display
- ・ keyboard

- ・ mouse
- ・ printer
- ・ console-CRT

制御プログラムのクラス、

- ・ device-handler

上層とデバイス管理との間で受渡しされる制御テーブル(入出力メッセージと呼ぶ)のクラス、

- ・ io-message

および、内部で使用するクラスからなる。

入出力デバイス毎に割込番号が決められており、割込が発生するとその番号に対応したプロセス(割込プロセスと呼ぶ)がハードウェアにより起動される。その割込プロセスで実行されるプログラムのクラスがdevice-handlerである。

disk-with-lrucacheは、ディスク・キャッシュ付きの固定ディスクのクラスである。マルチバス・メモリをセット・アソシティブ方式のキャッシュ・メモリとして使用することにより、ディスク入出力の高速化をはかった。

ビットマップ・ディスプレイは、ビットマップ・メモリの任意の大きさの矩形領域を単位として表示/転送処理を行う装置である。ビットマップ・ディスプレイに対応したクラスbitmap-displayでは、

- 矩形的ビットマップ・メモリの割り付け/解放
- ビットマップ・メモリ間の演算転送
- マーカの画面への表示

などの機能を実現している。

上層のサブシステムで入出力を行うには、つぎのようにする。

- (1) io-messageをひとつ確保する。
- (2) io-messageに必要な入出力パラメータ、データ・バッファなどをセットする。
- (3) 実行管理プログラムのチャネル機能を用いてio-messageをデバイス管理に送る。
- (4) チャネルを介してデバイス管理からio-messageを受取り、入出力動作の終了を知る。

6. おわりに

φのハードウェア資源であるプロセッサ、メモリ、入出力デバイスの管理が、SIMPOSの中でどのように実現されているかについて述べた。性能の評価と改良が今後の課題である。

S I M P O S の実行管理 —プロセスとストリーム

△三・三

島津秀雄、 吉田紀彦、 斎藤慎一、
(日本電気(株)) ((株)三菱総合研究所) (ヒューテック(株))
渡辺久晃、 服部隆
(沖電気工業(株)) ((財)ICOT)

1. はじめに

対象指向の言語でOSを構築する場合、プロセスの概念をうまく導入させる事は、重要な事柄である。本稿では、ICOTで開発した逐次型推論マシンPSI上のOSであるSIMPOSのOS核部に含まれるマルチプロセスの実現機構とプロセス間同期・通信機構について述べる。SIMPOSは対象指向の機能を含んだ論理型言語ESPで記述されており、本稿で述べるサブシステムも対象指向的に構築されている。

2. マルチプロセス機構

2.1 プロセス

SIMPOSは、対象指向システムとして構築されているため、システムのすべての基本構成要素はオブジェクトである。

オブジェクトで構築されたシステムが実際に計算機上で動作するためには、各オブジェクトにcontext(実行環境)を与える必要がある。

それを実現するのに、

- (1) 最も並列性を重視したものとしては、システム中に存在するオブジェクトの1つ1つに別のcontextを与えるactor理論のような方法がある。
- (2) 最も逐次的なものとしては、すべてのオブジェクトの実行を同一のcontext上で行なうという方法がある。

SIMPOSは逐次マシン上に構築されているので、オブジェクトの1つ1つにcontextを与えても、実際に並列に動作させることは不可能であるし、また、contextの切りかえに伴うオーバーヘッドが多大なものになってしまう危険性がある。一方、既存のOSにもあるプロセスやタスクといった疑似的な並列性を表現する仕組みは、ある程度以上の大きさのシステムを構築する上で非常に有効である。そこでSIMPOSでは、processという特別なクラスを導入し、クラスprocessのオブジェクトのみが独立のcontextを持ち得ることにした。従って、利用者があるプログラムをプロセスとして動作させたいときには、クラスprocessのオブジェクトを陽に生成し、生成されたプロセス・オブジェクトにプログラムを渡すという方式をとれば良い。こうすることで、普通のオブジェクトと並行

して動作するオブジェクト(プロセス・オブジェクト)とを利用者が陽に使われられるようにした。

2.2 プロセス管理

クラスprocessは、contextを物理的に切りかえるという下位のレベルの機能を使って、マルチプロセスの環境を利用者に提供することを行なう。クラスprocessは、プロセスの生成・起動・中断・再開・終了・開放を実現する述語を提供する。また、スケジューラの部分はすべてのプロセス・オブジェクトで共有するので、クラスprocessのクラススロットで保持している。

プロセスは、新たに生成されるときに、プロセスとして実行すべきプログラムを渡されるが、このプログラムは、クラスas_programを継承したクラスのインスタンス・オブジェクトである。

図1は、クラスprocessのクラス構成を示している。ハードウェアで認識するプロセス制御ブロックを保持するcontextと、そのプロセスの保持する資源をチェインするresourceは、プロセスのインスタンススロットで保持される。

3. プロセス間同期・通信機構

3.1 ストリーム

オブジェクト間のインタラクションは、一方が他方のオブジェクトに定義された述語を呼び出すことで実現される。プロセス・オブジェクトを除くすべてのオブジェクト間のインタラクションは、逐次的に実行されるので問題ない。しかし、プロセス・オブジェクト同士のインタラクションでは、プロセス間の同期・通信を矛盾なく実現するために、何らかの同期メカニズムを提供する必要がある。

SIMPOSでは、プロセス間の唯一の同期・通信機構として、クラスstreamを提供している。ストリームはオブジェクトを流すパイプである。ストリームの一方の端にオブジェクトを入れ、他端からオブジェクトを取り出すことが出来るという一種のFIFOキューである。もしあるプロセスが、オブジェクトを取り出そうとしたときに、ストリームが空であると、他のプロセスがそのストリームにオブジェクトを入れるまで待たされる。この機能に

より、プロセス間の同期が実現される。このようにプロセス・オブジェクト間の同期・通信には必ずストリームを介することが必要である。SIMPOSでは、クラスstreamを基本クラスとして提供し、その上に種々の上位機能をもつプロセス間同期・通信機構を構築している。

3.2 その他の同期・通信機構

ストリーム・サブシステムは、機能的に次の3つのレベルに分かれる。

- ストリーム：同一プロセッサ上のプロセス間の単方向通信をつかさどる同期機構を有したパイプ
- チャネル：特にメッセージを送受信するためのストリーム
- ポート：チャネルを、双方向通信可能になるよう拡張したもの

これらのクラスは、継承によって階層化されている。ストリーム・サブシステムでは、これらのクラスに加えて、さらにこれらに種々の機能（優先度制御付き・End_of_streamのあるもの・有限個バッファで構築されるもの・オブジェクトを取りださずに見ることができるもの・複数個のストリームを同時に待てるもの）を加えたストリーム類も提供している。それらの付加機能は付加機能単位に1つのクラスとして利用者に提供される。従って例えば、利用者が優先度制御付きで、かつ有限長のバッファで構築されているポートを欲するとき、クラスport、クラスas_priority_stream、クラスwith_bounded_bufferを多重継承するだけで、必要十分な新たなクラスを作成することが出来る。

4. おわりに

SIMPOSにおけるマルチプロセスの実現機構とプロセス間同期・通信機構について述べた。現在これらのサブシステムはPSI上で稼働中である。



図1 プロセスのクラス構成図

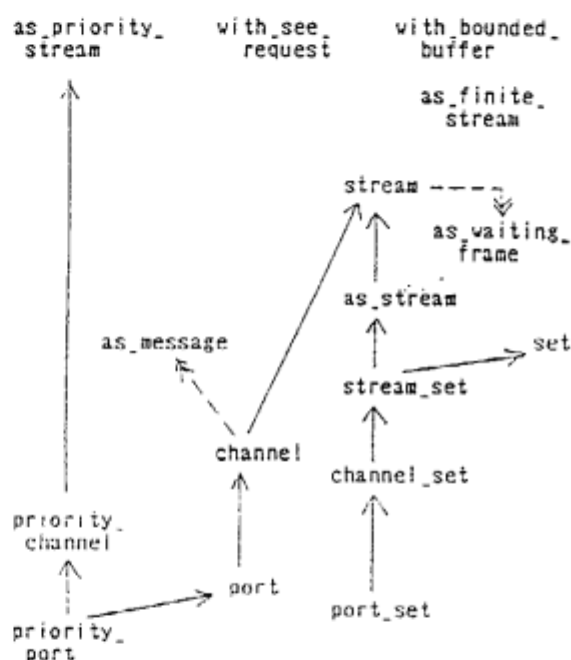


図2 ストリームのクラス構成図

```
class priority_port_with_bounded_buffer
nature
  as_priority_stream,
  with_bounded_buffer,
  port;
end.
```

図3 付加機能を持ったポートのクラス定義例

SIMPOSの実行管理 — プール

二 三

斎藤 慎一 渡辺 久典 島津 秀雄 吉田 紀彦 服部 隆
(ピーエンスシステム)(神電気工業株)(日本電気株)(三菱総合研究所)((財)ICOT)

1. はじめに

SIMPOSはオブジェクト指向的に構築されているため、実行主体となるプロセスが複数のオブジェクトを集合として保持する場合が多い。この時、必要なデータ構造を個別に用意したのでは大変だし、拡張性に乏しい。そこで、プール管理はSIMPOSの中で汎用的に使用されるオブジェクトの格納機能を抽象化し、これを継承関係を用いて様々なバリエーションを実現した。これにより、拡張が容易で統一されたアクセス法を持つデータ構造が実現された。

以下では、プールの設計思想と実現法について説明する。

2. プール

プールはオブジェクトを格納する容器である。プールに対しては以下のような操作を行なうことができる。

- ・プールにオブジェクトを格納する。
- ・プールからオブジェクトを取り出す。
- ・プール中のオブジェクトを削除する。
- ・プール中のオブジェクトの格納情報に関する問い合わせを行なう。

プール管理は、プールの格納機能をクラスとして抽象化し、これを継承させることにより様々な格納機能を持つプールを提供している。このようなプールの格納機能の抽象化は以下の視点で考えた。

- ・格納されるオブジェクト間の関係として順序があるかどうか。
- ・オブジェクトの格納位置をどのような形式で指定するか。
- ・同じオブジェクトの重複を許すかどうか。
- ・どのようなアクセス法が可能か。

(1) 順序

プールは、格納されるオブジェクト間の関係から、順序のあるプールと順序の無いプールに分類される。

順序のあるプールはオブジェクトを順序づけて格納するプールであり、格納されたオブジェクトはその格納位置によって識別される。順序のあるプールの基本となるものはリストとアレイである

(図2)。

i. リスト

要素数が可変のプールであり、いくつでもオブジェクトを格納することができる。

ii. アレイ

要素数が固定のプールであり、格納できるオブジェクトの数は有限である。

両者は一次元の順序を持つプールであるが、これらを多層に構成することによって多次元の順序を持つプール(バケット)を作ることにもできる。例えば、アレイを多層に組み合わせることによりアレイ・バケットができる(図3)。

(2) 格納位置

順序のあるプールにおける格納位置は、0から始まる昇順の整数で表わされる。この格納位置は固定的に決まっている場合(アレイ)と、相対的に決まる場合(リスト)とがある。

また、バケットにおける格納位置はパス・ポジションによって識別される。パス・ポジションは検索される順路に沿って複数の位置の並びを表現するものであり、次のような形式で表わす。

[position1, position2, ..., positionN]

(3) 重複

プールは、格納されるオブジェクトの重複を許すプールと重複を許さないプールとがある。順序の無いプールについては、重複を許すプールとしてバグ、重複を許さないプールとしてセットがある。セットは、オブジェクトを格納する時に重複するオブジェクトを排除することによって実現される。

3. アクセス方式

プールに対する操作を行なうためのアクセス方式としては次の3つがある。

- ・直接アクセス
- ・順次アクセス
- ・キーによるアクセス

それぞれのアクセス法においては、その操作のインタフェースが同じになるように設計した。これによって、プールの内部構造を意識せずにアクセスすることが可能となっている。

(1) 直接アクセス

直接アクセスは格納位置を指定してアクセスを行なう方式であり、隣接のあるプールの機能として実現される。アレイの場合は文字どおりそれぞれの位置へ直接アクセスできるが、リストの場合は内部的に順次検索が必要である。

(2) 順次アクセス

順次アクセスはタップという考え方を導入することによって実現している。タップとはプールの蛇口という概念を持つものであり、すべてのプールに対して設定可能である(図4)。

タップを用いることにより、プールの内部構造やオブジェクトの格納位置を意識せずに順次アクセスを行なうことが可能である。すなわち、リストでもアレイでも同様の操作によってタップが設定され、同様の操作でオブジェクトを取り出すことができる。

(3) キーによるアクセス

キーによるアクセスはインデックスを用いることで可能となる。インデックスはオブジェクトにキーを対応させて格納するプールであり、そのデータ構造や操作の方法は順序のあるプールと類似したものとして実現されている。

キーによるアクセスを行なう場合、内部での検索方法として、

- ・順次検索
- ・キーをハッシングする検索

があるが、両者の操作方法は同じものとして設計されている。

4. おわりに

以上述べたように、プールは様々な格納機能を抽象化することによって、インタフェースの共通化や系統的なバリエーションの展開が達成された。

しかし、抽象化と処理効率の間には背反的な事情がある。

- ・抽象化を重視すると処理効率が低下する。特にプールのようなOSの基本部分では無視できない点である。

- ・処理効率を重視すると拡張性が低下する。

現在、プールはこのような問題を合っており、今後改善の余地があると思われる。

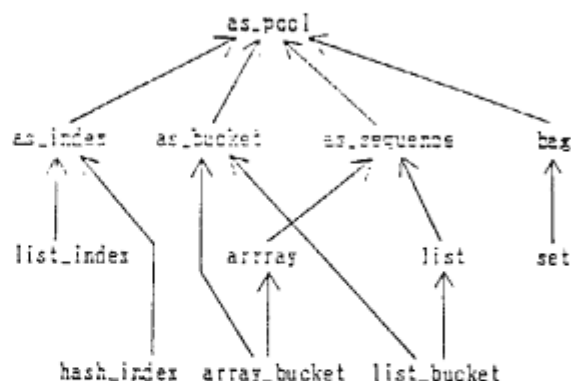


図1 プールのクラス継承図

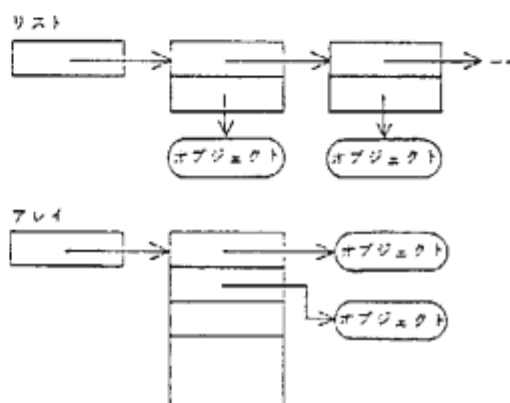


図2 リストとアレイ



図3 アレイ・バケット

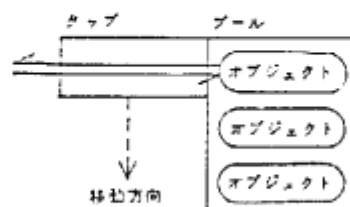


図4 タップ

SIMPOSの実行管理 — ワールド

4E-5

渡辺 久晃 (中電気工業 (株)) 島津 秀雄 (日本電気 (株)) 吉田 紀彦 ((株) 三菱総合研究所)
 齊藤 慎一 (ビーコンシステム (株)) 関部 隆 ((財) ICOT)

1. はじめに

SIMPOSは、逐次型推論マシン(PSI)のために開発されているプログラミング/オペレーティング・システムである。

ワールド管理は、SIMPOSの実行管理プログラムの1つであり、プロセスに実行環境を与えることを目的としている。

本稿では、その概念及び基本構成・機能を述べる。

2. ワールド管理の概念

ワールド管理では、主記憶上のオブジェクト(資源)を名前前で管理して、ユーザがプロセス実行中の任意の時点でオブジェクトにアクセスできる環境(実行環境)を提供する。オブジェクトを管理するにあたっては、ディレクトリ方式を採用している。ディレクトリは、名前前でオブジェクトを登録するテーブルであり、階層構造をとることができる。

システムには、システムに唯一のルートのディレクトリ(システム・ルート・ディレクトリ)を起点とするディレクトリの階層木から成るシステム・ディレクトリ・トリーがある。これは、ディスク上のオブジェクトへアクセスする際にインタフェースとなるディレクトリ(パーマネント・ディレクトリ)を含んでいる。システム・ディレクトリ・トリーは、各プロセスから共通にアクセスできる。

各プロセスは、システム・ディレクトリ・トリーの一部や、プロセス上でテンポラリに作られたディレクトリをリストで連結したものを実行環境として持つ。これは、ワールドと呼ばれる。

次にディレクトリ、パーマネント・ディレクトリ、ワールドのそれぞれについて基本構成・機能を述べる。

3. ディレクトリ

名前前でオブジェクトを登録する管理テーブルである。ディレクトリは、次のものから構成されている。

- (1) 一連のディレクトリ・エントリ
- (2) ディレクトリに対する付加情報

(1)の各々は、ディレクトリに登録されるオブジェクトの名前とオブジェクトから成る。

(2)は、ディレクトリの新規作成・参照・変更の日時、ディレクトリへアクセスするためのパス名、ディレクトリへのアクセス許可情報から成る。

ディレクトリは、システム・ルート・ディレクトリを起点に階層構造をとることができ、ディレクトリの階層構造上の位置付けは、(4)のパス名によって知ることができる。ディレクトリの提供する基本機能は、主記憶上のオブジェクトの登録、参照、削除等である。

パス名は、階層構造下にあるディレクトリやパーマネント・ディレクトリを辿ってオブジェクトにアクセスする場合、あるいは現在実行中のプロセスのもつワールドの中の部分ディレクトリを辿って、オブジェクトにアクセスする場合に利用され、名前と") " 記号を組合せたストリングによって表現される。例えば、図1のような階層構造下にあるオブジェクトdにアクセスするには、パス名" a 1) b 1) c 1 " が必要になる。

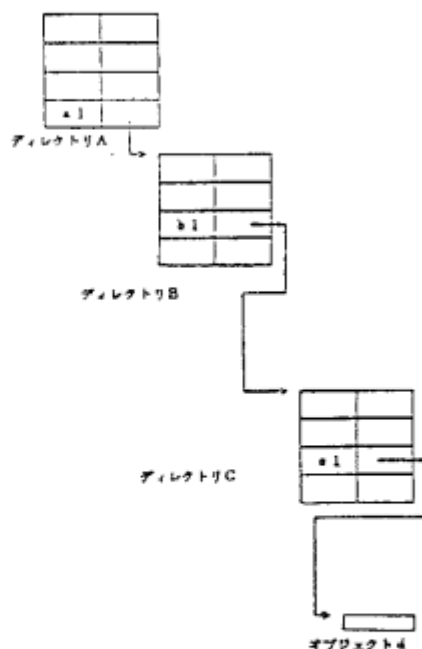


図1 ディレクトリの階層

4. パーマネント・ディレクトリ

パーマネント・ディレクトリは、ディレクトリ・ファイルを保持しているディレクトリである。ディレクトリ・ファイルは、名前でディスク上のオブジェクトを管理している管理テーブルである。パーマネント・ディレクトリは、主記憶上でディレクトリの下に作られ、主記憶上のオブジェクトをディスクへ格納したり、逆にディスクから主記憶上へ取り出す場合のインターフェースとして使われる(図2)。このとき、パーマネント・ディレクトリが保持するディレクトリ・ファイルが、キャッシュの役割を果たす。

一度パーマネント・ディレクトリを作れば、内部的には、それに対応するディレクトリ・ファイルがディスク上に1つ作られ、そのディレクトリ・ファイルがパーマネント・ディレクトリにセットされるので、ユーザは、直接ディレクトリ・ファイルへアクセスする必要がなくなる。ユーザが主記憶上のオブジェクトをディスクへ格納したり、ディスクから取り出したい時は、パーマネント・ディレクトリに対して命令を出せばよい。

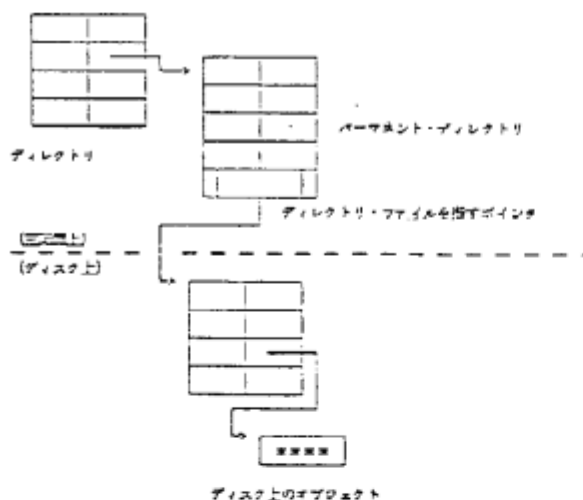


図2 パーマネント・ディレクトリ

5. ワールド

ワールドは、一連のディレクトリがリストで連結された構成になっている。ワールドは、それぞれのプロセスに対して設定できるようになっている。各々のディレクトリは、システムに唯一のルート・ディレクトリを起点に階層構造をなしているディレクトリの部分でも、また、プロセス実行中にユーザがテンポラリに定義したディレクトリのいずれであっ

てもよい(図3)。

基本機能は、次に示す通りである。

- (1) ワールドにディレクトリを付加する。
- (2) ワールドからディレクトリを取り除く。
- (3) ディレクトリの機能のサポートにより、ワールドにおける主記憶上のオブジェクトの検索・付加を行なう。

ユーザは、プロセス実行中に、作ったオブジェクトの格納・取り出しなどができる環境が欲しい時には、ワールドを作ればよい。システム・ルート・ディレクトリを起点とした階層構造下にあるディレクトリを実行プロセス中のワールドにセットして、そこにオブジェクトを登録すれば、プロセス実行中の必要な時点でそのオブジェクトへアクセスできる。再度、ルート・ディレクトリから遡ってオブジェクトへアクセスする必要はなくなる。

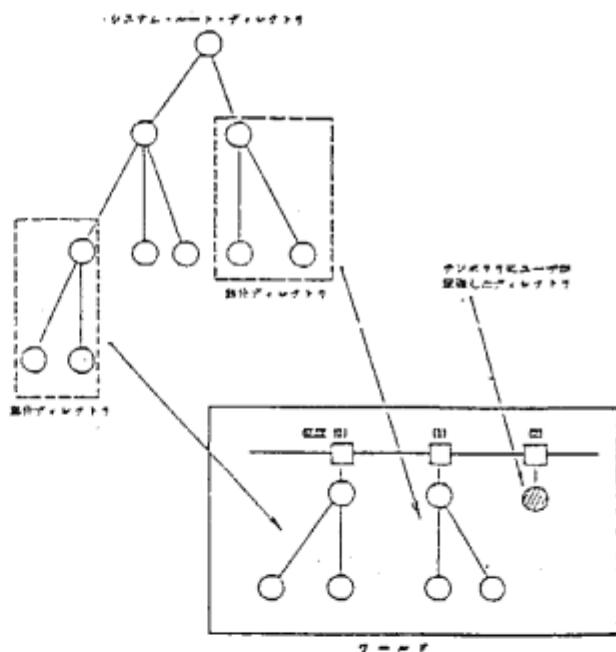


図3 ワールド

6. おわりに

以上のような基本構成と機能をもつ本システムのうち、現在、ディレクトリとワールドは、PSIを使って、実機デバッグ中である。今後は、ファイル・サブシステムを取り込んだのパーマネント・ディレクトリの実機デバッグ、ディレクトリに対するアクセス許可メカニズムの設計の仕様検討などが、重要な課題である。

43-6

SIMPOSのウィンドウ・サブシステム

飯間 豊
(沖電気工業(株))

中沢 修

板本 章二
(松下電器(株))辻 順一郎
(財)ICOT)

1. はじめに

SIMPOSは、PSI(逐次型推論マシン)上に開発中のプログラミング/オペレーティングシステムである。

ウィンドウ・サブシステムは、SIMPOSの中核部であるカーネル/スーパーバイザと、上位システムであるプログラミングシステムとの間に位置し、ビットマップディスプレイ、キーボード及びマウスを対象とする入出力管理プログラムである。これらの入出力装置に対する複数プロセスからの同時利用を可能とすることにより、PSI上のプログラミング環境を向上させることを目的としている。

本稿ではウィンドウ・サブシステムの概要及び特徴について述べる。

2. 概要

ウィンドウとは、ビットマップディスプレイ、キーボード及びマウス上に作られた論理的端末装置である。ウィンドウ・サブシステムはユーザプロセスからの要求に従って、互いに異なった特性を持つ複数のウィンドウを作成する。各プロセスは物理的装置を意識することなく入出力を実行出来る。

ユーザ側から見ると、ウィンドウはある特定のプロセスに接続されている端末である。ディスプレイスクリーン上に表示される複数のウィンドウを通じて、ユーザは同時に複数のプロセスと対話出来る。

ウィンドウはビットマップディスプレイ上に長方形領域として表示される。ウィンドウのサイズ、スクリーン上の位置は自由に設定可能である。同時に多数のウィンドウが存在する場合には、ウィンドウ群は、机上の紙の様に互いに重なり合っているかの様に表示される。

本システムの機能を以下に示す。

- ・オーバーラップマルチウィンドウシステム
- ・ウィンドウのサブウィンドウ分割可能
- ・マウスによるポジショニング
- ・文字及びビットグラフィック出力
- ・ウィンドウボーダ/ラベル等の表示
- ・メニュー等の多様なウィンドウのサポート
- ・ウィンドウのオーダメイドが容易

3. 構成

本システムの構成を図. 1に示す。ウィンドウ・マネージャは、ウィンドウ・マネージャプロセス下でバックグラウンド的に実行されるプログラムであり、主にウィンドウ・サブシステムとユーザプロセスあるいはデバイスハンド

ラプロセスとの間のプロセス間通信を司る。ウィンドウの実体は、ウィンドウクラスのインスタンスであるウィンドウオブジェクトである。ウィンドウに対する入出力等の要求は、ウィンドウ・オブジェクトによって処理される。

ウィンドウクラス群は、図形出力や文字入力といった単一機能について記述されている多数の部品クラスと、それらのうちのいくつかを選択的に継承することによりウィンドウとしてのまとまった機能を実現している完成品クラスとから構成されている。ウィンドウクラスの追加により、容易にユーザ定義ウィンドウの作成が可能である。

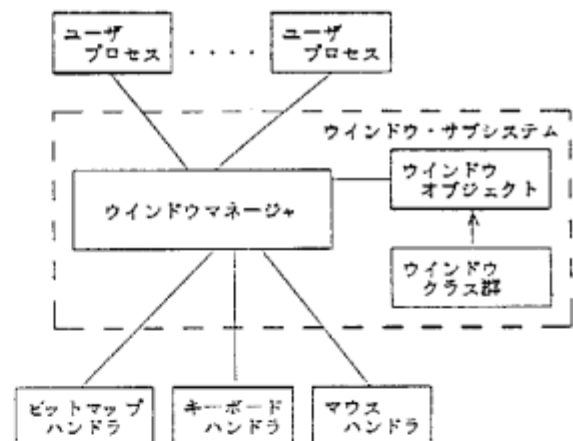


図. 1 システム構成

4. ウィンドウの階層化

高度なディスプレイ制御を行なうには、単にディスプレイスクリーンのマルチウィンドウ化のみでは不十分であり、ウィンドウ個々がさらに再帰的にマルチウィンドウ化出来る必要がある。本システムでは、ウィンドウに階層性を持たせることによって、再帰的マルチウィンドウを容易に実現している。

図. 2に示す様に、階層のルートには、スクリーンと呼ばれるディスプレイ画面全体を表すシステムウィンドウがある。ユーザ・プロセスからの要求によって作られるユーザウィンドウ(ウィンドウ1及び2)はすべてスクリーンの子ウィンドウとなる。また、ユーザウィンドウに子ウィンドウ(ウィンドウ3及び4)を作ることで、ユーザウィンドウがさらにマルチウィンドウ化される。図. 2の階層構成に対応する画面表示の一例を図. 3に示す。



図. 2 ウィンドウ階層

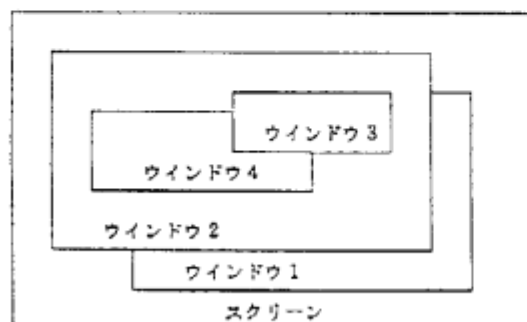


図. 3 表示画面

本システムには、集約的管理プログラムは存在しておらず、ウィンドウの管理はすべてウィンドウ階層内で行なわれる。即ち、ウィンドウは親ウィンドウによって管理される。また、親ウィンドウがさらにその親ウィンドウによって管理されているため、間接的にすべての祖先ウィンドウによって管理される。

5. 出力

オーバーラップマルチウィンドウシステムにおいては、他のウィンドウによって隠されているために、ユーザに見えないウィンドウが存在する。ウィンドウに文字等を出力する際には、それがユーザに見えるものであるか否かを考慮した処理が必要となる。

最も単純な出力方式は、完全に見えているウィンドウにのみ出力を許し、隠されているウィンドウへの出力を要求したプロセスは、そのウィンドウが見える状態となるまで待たされるというものである。しかしこの方式では、バックグラウンド的に実行されているプロセスからのメッセージをユーザに即座に伝えることが出来ないという欠点がある。他にも多くの方式が考えられるが、ユーザプロセスの多様性を考慮すると、いずれにも何らかの欠点があり、ただ1つの方式を導入することは困難である。

そこで本システムでは、以下に示す4種類の出力方式をサポートし、いずれの方式を用いるかはウィンドウ個々に設定することとした。

- ・WAIT 当該ウィンドウが見える状態となるまで、ユーザプロセスを待たせる
- ・NOTIFY 警告メッセージを表示する
- ・OUTPUT 隠されたまま、出力を行なう
- ・SHOW 当該ウィンドウを強制的に見える状態とする

6. 入力

ウィンドウは入力デバイスとして、キーボードとマウスを持つ。これら2つのデバイスは、かなり異なった性質を有しており、それぞれ別個の方式で制御されている。

文字入力デバイスであるキーボードは、複数プロセスからの同時共有が本質的に不可能なデバイスであり、キーボード入力の送り先ウィンドウ（セレクトッドウィンドウ）は、一意的に決定されている必要がある。このため、セレクトッドウィンドウは、主にユーザの指定により決まるものとする。

一方ポジショニングデバイスであるマウスは、ディスプレイスクリーン上を対象とするデバイスであり、スクリーンに表示されている複数のウィンドウからの共有が可能である。マウス入力の送り先ウィンドウは、その時点でマウスがどのウィンドウ上に在るか、即ちマウス位置を示すマウスマーカがどのウィンドウ上に在るかによって決定される。

7. メニューウィンドウ

マウスはメニュー選択のために最も頻繁に使われると考えられる。そこで、本システムではメニュー選択機能を備えた種々のウィンドウ（メニューウィンドウ）を提供している。メニューウィンドウは、項目の表示及びマウスによってそれらを選択する機能のみを持つウィンドウである。項目を選択することにより、ユーザによって定義されたコマンドがユーザプロセスに送り返される。また、項目の選択と同時にウィンドウイメージがスクリーン上から自動的に消えるメニューウィンドウも作成出来る。

8. おわりに

ウィンドウ・サブシステム単体における基本動作については、PSI上において検証済である。さらに、エディタ等の上位サブシステムと接続して、総合的検証を進めて行く。

本システムの様に、マンマシンインタフェースに係わるソフトウェアにおいては、実際の環境下での使用経験に基づく継続的な改良が必要である。今後、他サブシステムとのインタフェースをも含めて、処理効率の向上、機能拡充について検討を加える予定である。

SIMPOSのネットワーク・サブシステム

4E-7

高山 幸男
(沖電気工業 (株))

服部 隆
((財) ICOT)

1. はじめに

パーソナル逐次型推論マシン (PSI) は第5世代コンピュータシステム開発の一環として作られているソフトウェア開発用のワークステーションである。

PSIにはネットワーク機能があり、複数のPSIや関係データベースマシンDELTAと繋ぐことにより、ユーザに対してスタンドアロンのマシンでは得られない幅広いプログラミング環境を提供する。すなわち、他の人が作ったプログラムを共用したり、メールシステムを利用したり、あるいは、DELTAに格納されている膨大なデータベースを利用したりできる。ネットワーク機能は、ウィンドウ、マウスなどの高度なマンマシンインタフェース、ワークステーションならではの一定したレスポンスタイムとともに、PSIを高性能で使いやすいprolog処理系にするための機能として重要である。

本論文では、ネットワーク機能をサポートするために開発されたSIMPOSネットワーク・サブシステムの概要を報告する。

2. ハードウェア構成

SIMPOSネットワーク・サブシステムは、LIA (LAN Interface Adapter)、LIB (LAN Interface Board) という2つのデバイスを使ってネットワーク機能を実現する。

LIAは、LAN用の高性能デバイスであり、コマンドデータを送ることにより仮想回線の形成および解除、更に形成された仮想回線を使ったデータの送受信を行なうことができる。

このように、従来のOSに於いて通信機能の下位レイヤとして記述されていたソフトウェアがLIAに吸収されているので、OSで記述しなければならない部分が軽くなり、見通し良く設計できるという利点がある。半面、デバイスの機能が高めに固定されているので、OSをトップダウンに設計してゆく際の自由度が制限されてしまう難点がある。

LIBは、PSIとLIAとの間のインターフェースボードであり、PSI上のデータをLIAに送ったり、逆にLIAのデータをPSI上に書き込んだりする。

ネットワークの物理的構成を図-1に示す。

3. ネットワーク・サブシステムの機能と特徴

ネットワークOSとしてSIMPOSを見たとき、その機能は大きく3つのレベルに分れる。

1) マシン間通信モード

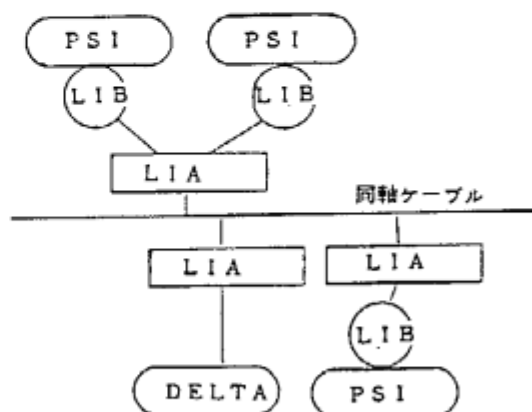


図-1 ネットワークの物理構成

LIAの機能を直接的に利用して、通信網を形成、消滅させたり、データを送受信したりする機能を提供する。これらの操作は、プラグと呼ばれるオブジェクトを通じて行えるようになっている。

また、通信網は、全二重のセッションによるメッシュ方式を想定し、仮想回線はケーブルと言うオブジェクトで表される。ケーブルは、ソケットと呼ばれる、位置を示すためのオブジェクトの間に張られる。プラグは、いずれかのソケットに“差し込まれて”おり、そのソケットから他のソケットにケーブルが張られると、そのケーブルに対してポインタを張る。ケーブルは、ひとつのソケットから複数本張ることができ、ユーザは一つのプラグからこれらを操作することができる。また、PSIやDELTAなどのマシンは、ノードと呼ばれるオブジェクトで表現され、マシンアドレスなどの情報が管理される。

ノードとソケットは、ディレクトリに管理されており、ユーザは必要に応じてパスネームを指定してそれらのオブジェクトポインタを手に入れることができる。ノード、ソケット、ケーブル、プラグの関係を図-2に示す。

2) プロセス間通信モード

[PSI-PSI通信のみ]

単一マシン内のプロセス間通信をネットワーク環境にまで拡張する機能を提供する。

SIMPOSにおけるプロセス間通信は、スーパーバイザで提供されているポート/チャネル方式である。このモードでは、他マシン内のポート/チャネルを自マシン内のそれと同様にアクセスできる。

3) リモートオブジェクト操作モード

[PSI-PSI通信のみ]

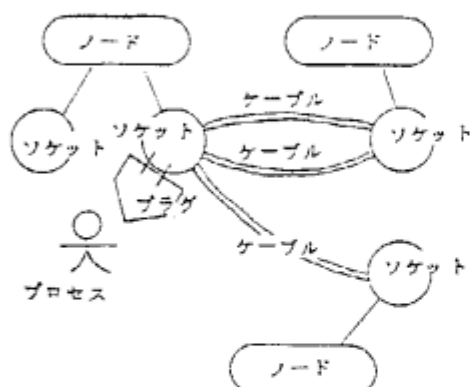


図-2 ノード、ソケット、ケーブル、プラグの関係

他マシン内のオブジェクトを、自マシンのそれと同様に操作できる機能を提供する。プロセス間通信モードとこのモードによって、OSがユーザに提供できるプログラミング環境が単一のPSIに留まらず、複数PSIにまで自然に拡張される。

ただし現在想定しているものは、2) のためのポート/チャンネル、ディレクトリ、ファイルのリモート操作である。

4. インプリメントの概要

1) マシン間通信モード

このモードでは、ユーザプロセス、ネットワークマネージャプロセス、LIBハンドラプロセスの3つのプロセスが互いにメッセージ通信をして操作を実現する。プラグは、ポートとしての機能を持っていて、ネットワークマネージャプロセスとのメッセージ通信に使われる。

ユーザプロセスが、プラグに操作をすると、その内容がプラグメッセージとして、ネットワークマネージャに送られる。ネットワークマネージャは、そのメッセージを解釈してLIA用のコマンドデータのオブジェクトを作り、LIBメッセージの形でLIBハンドラに送る。LIBハンドラはそのコマンドデータのオブジェクトをコードに変換してLIBに読み込ませる。

逆に、LIBからPSIに送られてきたコマンドデータは、LIBハンドラによってコマンドデータのオブジェクトに変換され、LIBメッセージの形でネットワークマネージャに送る。ネットワークマネージャは、コマンドデータオブジェクトからプラグメッセージを作り、プラグに返す。

ネットワークマネージャは、プラグメッセージとLIAのコマンドデータとの間の変換、および、プラグメッセージ、LIBメッセージの通信制御を行なう。

2) リモートオブジェクト操作およびプロセス間通信モード

他マシン内のオブジェクト(実オブジェクト)を操作する場合、そのオブジェクトの“代理オブジェクト”(リモ

ートオブジェクト)を自マシン内に生成する。同時に、他マシン内には“代理プロセス”が生成される。リモートオブジェクトは、実オブジェクトと同じユーザインタフェースをもっており、ユーザはそれが“代理オブジェクト”であることを意識せずに操作できる。

ユーザがリモートオブジェクトを操作すると、その内容がリモートオブジェクトメッセージの形で相手マシンの“代理プロセス”に送られる。“代理プロセス”は、そのメッセージに従って実オブジェクトに操作を実現する。

リモートオブジェクトメッセージは、マシン間通信モードのプラグのデータ送受信機能を使って相手側マシンに送られる。

プロセス間通信モードのインプリメントは、相手マシン内のポートやチャンネルの代理オブジェクトを自マシン内に生成することによって行なわれる。

SIMPOSのリモートオブジェクトは、CMUで開発されたACCENTで提案されたリモートポートの考え方を拡張したものである。

5. おわりに

現在、マシン間通信モードの開発はほぼ終了している。プロセス間通信モード、リモートオブジェクト操作モードは、我々の提案した基本方式をもとにスーパーバイザ開発グループなどにおいて検討中、ないし開発中である。

マシン間通信モードは、LIAの仕様に沿って開発されており、LIAにソケットの皮を被せたような構成になっている。それに対して、プロセス間通信モード、及びリモートオブジェクト操作モードは、デバイスの仕様からは独立しており、SIMPOSの論理構成に調和した設計を模索・実験する、といった方針で開発されている。

リモートオブジェクト操作モードなどのインプリメントを、よりすっきりしたかたちで行なえるようにするためには、オブジェクト・ポインティングメカニズムや手続き呼出しの方法など、SIMPOS全体の論理構成や、システム記述言語についての検討を更に深めてゆく必要があると思われる。また、ACCENT(前出)や、同じくネットワークをサポートし、オブジェクトオリエンテッド的な考え方で作られているCMUのHydraでは、プロテクションメカニズムがかなり追求されている。SIMPOSにおいてより強力なプロテクションメカニズムを考えた場合、ネットワークサブシステムをどう設計すればよいかという興味深い問題が起ってくるであろう。

参考文献

- Richard F. Rashid, et al. ACCENT: An communication oriented network operating system kernel. CMU (1981)
- R. A. Wulf, et al. Hydra/C.mmp: An experimental computer system. McGraw-Hill, New York (1981)

小松光雅 真野忠志 小長谷明彦 服部隆
(ビーコンシステム(株)) (ビーコンシステム(株)) (日本電気(株)) ((財) ICOT)

1. はじめに

SIMPPOSのファイル・サブシステムはPSIにおけるディスクへのデータ格納機能とアクセス機能を提供する。その最大の特徴は、ディスク上のファイルをオブジェクトとしてとらえる点にある。これにより、ファイル・サブシステムをSIMPPOSの対象指向の枠組の中に自然に組入れられる。本稿では、この考え方に基づく対象指向ファイル・システムの設計思想と実現法について述べる。

2. ファイル・オブジェクト

一般に、ファイル中に格納されるデータは半永久的な存在であり、プロセスとともに生成、または、消滅するメモリ中のデータとは本質的に性質を異にする。通常、ファイルは外界に存在する特別な記憶領域であり、このデータ格納領域とファイルへのアクセスの為に操作命令を分解してファイル・システムの構築を考えると、SIMPPOS全体がオブジェクトを基本単位とした対象指向システムとして構成されることについての一貫性が損なわれる。

SIMPPOSのファイル・サブシステムでは、このような問題点を解決する為に、ファイル・オブジェクトという考え方を導入している。ファイル・オブジェクトは、概念的には、格納領域としてのディスク・ファイルにファイル・アクセス機能を付加したものと考えられる。ファイルへの入出力は、このファイル・オブジェクトに備えられたアクセス機能を対象指向呼出しで起動することにより実現される。

また、ファイルをオブジェクトとしてとらえることの利点の1つとして、アクセス機能をクラスとして抽象化させることにより、クラスの継承機構を有効に利用して種々のファイル・オブジェクトを定義できることがあげられる(図1)。

このファイルのクラス化にあたって、格納できるデータに着目し、(1) ビット列、整数、または、文字列等のデータ、(2) オブジェクト、(3) 知識の3段階に分けて実現することにした。さらに、一般のファイルだけでなく、ディレクトリ・ファイルやファイルの物理的な格納情報を格納する場所(VTOC)もオブジェクトとして扱って、対象指向システムとしての統一性を計っている。

3. データ・ファイル

ファイル・サブシステムでは、最も基本的な格納機能として、データを格納するファイルを提供する。このデータ・ファイルに対しては、さらに、データの格納形式及びアクセス法の違いから細分化を行った。すなわち、データ格納形式の違いからは、ビット列の並びから構成されるバイナリ・ファイルと、いくつかの意味のあるデータの固まり(レコード)の並びから構成されるレコード・ファイルに分類される。このレコード・ファイルに対しては、可変長のレコードを扱うヒープ・ファイルと、固定長のレコードを扱うテーブル・ファイルに分類される。

次に、アクセス法の違いからは、前述の3つのファイルが直接アクセス法を持つものに対して、テーブル・ファイルについては、キー値によるインデックス・アクセス法(インデックス・ファイル)が考えられる。さらに、それぞれのファイルに対して順次アクセス法でデータの格納、または、取り出しを行なう機能としてファイル・タプが用意される。

4. インスタンス・ファイル

一般に、対象指向で構成されるシステムにおいて、メモリ中のオブジェクトをディスク・ファイルに格納できる機能を提供することは、対象指向システムにおける単一メモリの実現を可能とする。しかしながら、この機能の実現に当っては、メモリ上で複雑なネットワーク構造を取るオブジェクトを一次元的な領域(インスタンス・レコード)に変換する必要がある。この為の基本的な機能を体系化するには、まだ多くの検討課題を残している。たとえば、インスタンス・レコードを実際に作成する場合、各オブジェクト毎にインスタンス・レコードへの変換手続きを定義しておく必要があり、この変換機能をインスタンス・ファイルが標準的に用意することは、今の所、困難である。

本ファイル・サブシステムでは、レコード・ファイルに対応して、可変長のインスタンス・レコードを扱うインスタンス・ヒープ・ファイルと、固定長かつ同一フィールド様式のインスタンス・レコードを扱うインスタンス・テーブル・ファイル(同じクラスのオブジェクトの格納)の

提供を考えた。この中で、インスタンス・テーブル・ファイルの実際のインプリメント例として、ファイル・オブジェクトを格納する為のファイル・ファイルを実現している。ファイル・ファイルは、通常のファイル・システムにおけるボリューム管理テーブル(VTBC)に相当するものであり、ファイル・オブジェクトのボリューム内での格納情報を管理している。また、ファイル・ファイルはマウントされるボリュームに対して唯一のみ存在し、異なるクラスで定義されるファイル・オブジェクトを1つのインスタンス・テーブル・ファイルに収容する為に、フィールド情報の1つにファイル・クラスを識別する情報を持たせている。

さらに、インスタンス・ファイルのインスタンス・レコードは、格納されるオブジェクトの名前と対応付けられる(図2)。これを実現するものがディレクトリ・ファイルであり、オブジェクトの名前をキー値として、レコード・ポインタをデータとして保持している。レコード・ポインタはメモリ上のオブジェクト・ポインタに相当する。このディレクトリ・ファイルは階層構造を形成していて、オブジェクトは順路名で検索される。また、この順路名は、ワールド・サブシステムの提供するディレクトリ・システムの中に統一されるものである。

5. おわりに

SIMPOS設計の基本概念である対象指向に拘ったファイル・サブシステムの構成について述べた。本ファイル・サブシステムをクラス・システムで実際にインプリメントしてみても、新しいファイルの定義やファイル・オブジェクトへの操作としての新しい機能の追加が、クラスの継承機構を利用して比較的容易に行なえることが分った。しかし、反面、クラスの数が多くなることから発生する保守の問題や、継承したクラスのメソッドの中に再定義しなければならないものもあって、必ずしも単純に機能拡張が行なえない場合がある。

また、本ファイル・サブシステムは、格納機能とアクセス機能の中の基本的な機能についてインプリメントがなされているが、将来的には、データ・ベース・システムに統合される方向に向かいつつある。さらに、最終的には、格納データとして知識を格納する知識ベースの構築に対するサポートを目指すものであり、これについては、知識の表現方法等の未解決の問題をかかえて、知識ベース研究の大枠の中からのアプローチ課題が残されている。

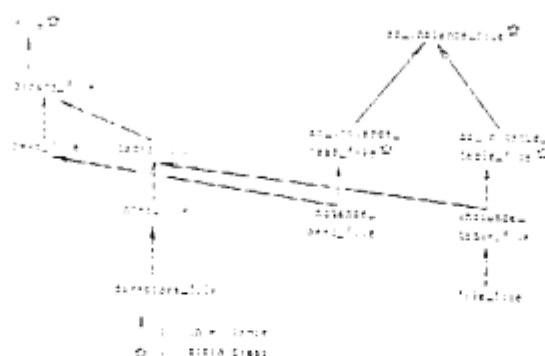


図2 インスタンス・ファイルの構造図

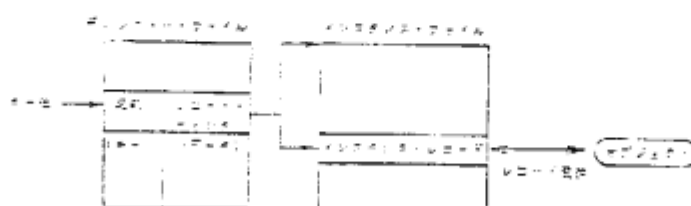


図3 ディレクトリ・ファイルのインプリメント図

SIMPOSのIPL方式

上田 尚純

(三菱電機(株))

東条 敏

(株)三菱総合研究所

黒川 利明

(財)ICOT

1. はじめに

個人使用向けのProlog高速実行マシンである逐次型推論マシン ϕ (PSI)上に、論理型プログラム開発支援用プログラミング/オペレーティング・システムSIMPOSを開発中である。 ϕ はPrologインタプリタや一部OS機能をファームウェアで実現している事、SIMPOSは対象指向の抽象化機構を持つ言語(ESPと称される)で記述されている事等の事情で、そのIPLは汎用計算機での通常のOSのIPLと一部異なる面がある。SIMPOSのIPL時のシステム立上げ手順について、その特徴を中心に述べる。

2. SIMPOSのIPLの問題点

SIMPOSのIPL手順を検討する際、次の点が問題となった。

- (1) ϕ ではメモリ管理、一部プロセス管理をファームウェアで行っており、IPLで最初に行われるプログラム(ブータ)の実行環境もやはりファームウェアで設定せねばならない。
- (2) ϕ の機械命令(KL0と称される)はテーブル形式でPrologを表現したものになっており、虫のバッチをワード単位でIPL時に施すのは困難である。
- (3) SIMPOSを構成するクラスは全てライブラリと称するSIMPOS内の一部機能で管理されるが、IPL中はこの機能がまだ作動していない。

これら問題点に対して、以下の方針でIPL手順を設計した。

- (1) IPL時、ブータが動作できるシングル・プロセス環境をファームウェアで設定する。マルチ・プロセス環境はブータが作動して作り上げる。
- (2) 虫のIPL時の修正はクラス単位の置換で行う。
- (3) ブータ等のIPLで使用するプログラムも対象指向言語で作成し、IPL中はIPL専用の簡易方式でクラス管理を行う。IPL終了後、これらクラスを正式にライブラリに登録する。

IPLで使用するクラスは、出来るだけSIMPOSの標準クラスとして作成したものを一部手直しして流用したが、SIMPOSのシステム記述言語であるESPが提供するクラス継承機能等を利用して容易にこの作業を行えた。例えば、割込を使用して動作する入出力ハンドラを一部変更して、割込を使用しないポーリング版の入出力ハンドラを作成した。

3. IPL手順

SIMPOSのIPL手順を図1に示す。 ϕ には低速デバイスの制御、本体部の診断等を行うマイクロ・プロセッサ内蔵のコンソール・プロセッサ(CSP)が装着されており、IPLもCSPにより開始される。CSPはハードウェアの初期設定やファームウェアのWCSへのロード等を行う。次いでファームウェアの初期設定ルーチンを実行

させてPrologインタプリタの動作環境の整備、シングル・プロセス環境の設定を行う。次にブータをディスクから読み込み主記憶にロードし、ブータを起動する。以後はソフトウェアの担当となり、ブータがSIMPOSのクラスをディスクから順次ロードしていく。

クラスはフロッピー・ディスク(FD)、または固定ディスク上の論理フロッピー・ディスク領域のファイル(ここでは以後クラス・ファイルと呼ぶ)で保持されている。1ファイルには1つないし複数個のクラスが含まれている。FDはIBM互換形式であり、固定ディスク上の論理フロッピー・ディスク領域も論理的に同じ形式となっている。CSPおよびブータはIBM互換形式FDを扱う機能を持っている。例えばブータはFD上に"SYSIPL"というファイル名で登録されており、CSPはFDのディレクトリをサーチしてこのファイルを見付けてロードする。ブータ

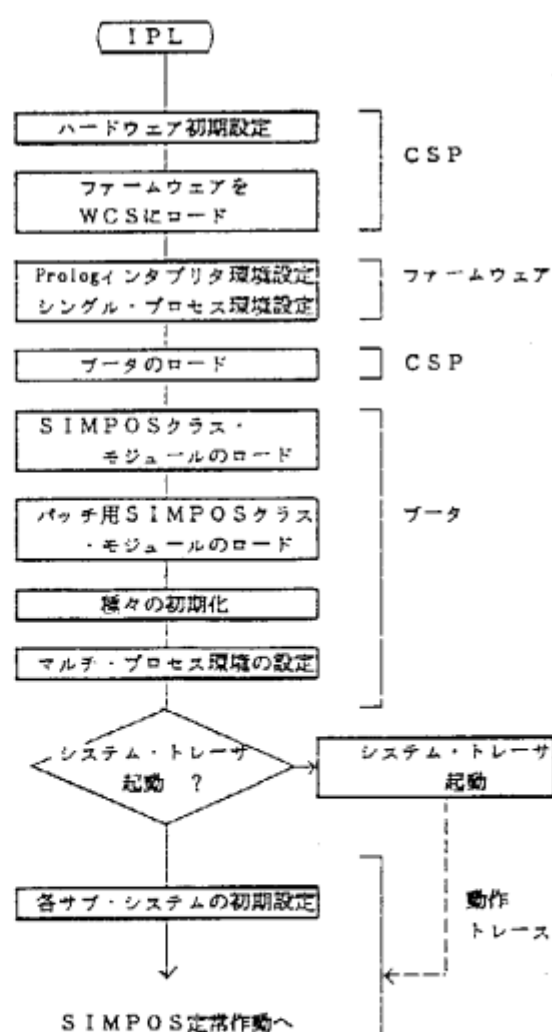


図1 SIMPOSのIPL手順

タのクラス・ファイルのロード手順は次のとおりである。

- (a) ロードすべきクラス・ファイル名を保持するファイル ("IPLLIST") の読み込み。
- (b) "IPLLIST" で指定されている各ファイルのロード。これによりコード部、ユーババビ、ウィンドウ・システム等のSIMPLOS中核部を構成するクラスが主記憶にロードされる。
- (c) パッチ用クラス・ファイル名を保持するファイル ("PATCHLIST") の読み込み。
- (d) "PATCHLIST" で指定されている各パッチ用クラス・ファイルのロード。

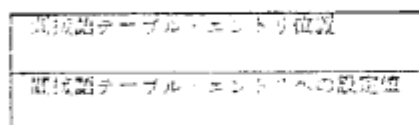
パッチ用クラス・ファイルは別のFDからも与えられるようになっており、搬送する機構により、通常のクラス・ファイルと同じ方法でロードすれば自動的に虫のある古いクラスと置換えられるようになって、また、正常なクラス・ファイルとパッチ用クラス・ファイルの形式は全く同じである。クラスのロードが完了すると、マルチ・プロセス環境の設定等のシステムの初期化を行う。

4. FD上のファイル形式と間接語テーブル

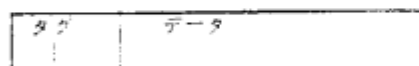
ブータによりロードされるクラス・ファイルのFD上のファイル形式を図2に示す。リロケータブル形式であり、ブータによりロード時にリロケートされる。コード本体部はKLOで表現されており、各ワードは8ビットのタグと32ビットのデータ部とより成る。タグの先頭2ビットは基本GC用であるが、IPL時はリロケーション用として使用している。間接語情報部は間接語テーブルにセットする値を指定するもので、対象となる間接語テーブル・エントリを指し示すワードとそのエントリに入れる値を保持す



(a) 上記のファイル形式



(b) 間接語ベアの形式



- 00XXXXXX : 絶対データ (リロケート不要)
- 10XXXXXX : 相対データ (リロケート要)
- 01XXXXXX : 間接語テーブル・エントリ位置

(c) コード部、間接語情報部のワード形式

図2. FD上のクラス・モジュールのファイル形式

間接語テーブル (IWT)

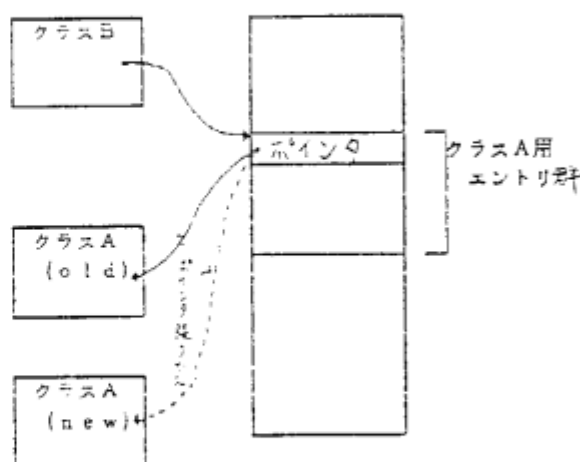


図3. 間接語テーブルによるクラス間リンクとそれによるパッチ方法

るワードのペアが並んだものである。

間接語テーブル (IWT) は主記憶上に置かれた、IPL中のクラス間の参照関係を管理するリンク用テーブルである。(図3) 各クラスが使用するIWTのエントリは、特別なデータベースで管理されており、一旦クラスにエントリ群が与えられると、そのクラスを再コンパイルしても同じエントリ群を使用する。あるクラス (A) が他のクラス (B) から参照される場合、他クラス (B) の参照ポインタは直接クラス (A) 内の被参照箇所を指すのではなく、その被参照箇所に対して確保されているIWTエントリを指すようコンパイルされる。IPL時にクラス (A) をロードする際にIWTエントリに被参照箇所へのポインタがセットされるが、これを指示するのが間接語情報部の間接語ペアである。KLOではこのような間接ポインタは特別なタグを持ち、インタプリタは自動的にポインタをたどってKLOを実行する。

クラス (A) に虫がいた場合、虫を修正して再コンパイルし、パッチ・モジュールとしてロードする。クラス (A) に対応するIWTエントリの内容が全て書き換えられ、新しくロードされたモジュールが古いモジュールに取って換わることになる。なお、古いモジュールが占めている領域はいずれGCで回収される。

5. おわりに

φはハードウェア、ソフトウェア共に高機能のため、IPL時の負荷もその分大きくなっているが、以上述べた方法により効率よいIPLを実現した。

(参考文献)

- SIMPLOSの概要 高木他 Logic Programming Conf. 1984. 3
- PSIのHW設計 藤他 同上

SIMPOSのシステム・トレーサ

佐藤裕幸，渡辺久晃，堀敦史，上田尚純，近山隆
 (三菱電機(株)) (沖電気工業(株)) ((株)三菱総合研究所) (三菱電機(株)) ((財)ICOT)

1. はじめに

我々は現在Prolog高速実行マシンである逐次型推論マシンφ (PSI) のオペレーティングシステム SIMPOSを開発中である。φはPrologを表形式で表現した高いレベルの機械命令体系 (KL0と称される) を持つ。SIMPOSは、Prologにオブジェクト操作機能を組込んだESPと称するシステム記述言語で統一して記述されている。このSIMPOSの開発支援システムとしては、CSP (コンソール・プロセッサ) 等があるが、これらはハードウェア・レベルのデバッグを目的として作られたものであり、symbolicなデバッグができず、ソフトウェア・レベル (Prolog) のデバッグには適していない。そのため、ソフトウェアで作られているESPプログラムの実行機構で、メソッド呼出し (対象指向のメッセージ・パッシング) 単位でのトレース機能が提供されている。しかし、マクロなメソッド・レベルのデバッグしか行えず、それより細かいPrologレベルのデバッグは行えない。そこで、KL0レベルの細かいトレースもでき、symbolicなデバッグが行えるシステム・トレーサを開発した。以下では、システム・トレーサのトレース機能とその実現法について、また、メモリのダンプ等トレース以外の機能についても報告する。

2. トレース機能

システム・トレーサのKL0トレースは、ヘッド・ユニフィケーションを基本としている。従って、DEC-

10Prolog等のゴール呼出しを基本としているボックスモデルのトレースとは異なっている。これはファームウェアのトレース・トラップ機能を利用しているからである。このトレース・トラップ機能には、ユーザ定義述語のヘッド・ユニフィケーション終了時と組込述語終了時にトラップ (プログラムの動きに同期して起こる割込の一種) を発生するKL0 step trap とある指定されたクローズのヘッド・ユニフィケーション終了時にトラップを発生するKL0 clause trap の二種類がある。これらを利用してシステム・トレーサではトレースを行うタイミングによって以下の5種類のモードを用意している。

■step

ユーザ定義述語のヘッド・ユニフィケーション終了時と組込述語終了時にトレースする。

■clause

ユーザ定義述語のヘッド・ユニフィケーション終了時のみトレースする。

■procedure

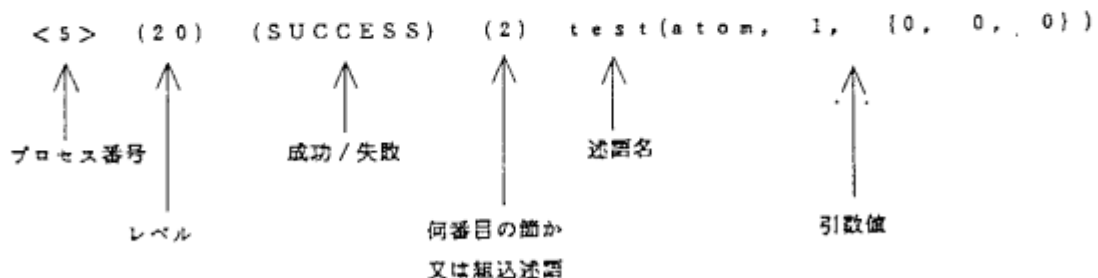
last clause以外のユニフィケーション成功時とlast clauseのユニフィケーション終了時にトレースする。

■spy

スパイポイントでのみトレースする。述語単位にスパイポイントを設定できる。

■off

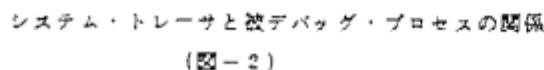
トレースは行わない。



トレース表示例 (図-1)

をコンソールに表示する。(図-1)

- トラップ時の実行コード・アドレス (IBR)
- トラップ時のローカル・スタックのベース・アドレス (CLBR)
- トラップ時のグローバル・スタックのベース・アドレス (GGBR)



現在、システム・トレーサはSIMP05のデバッグに利用されている。今後は、プリンターやファイル等へのログ機能、割込等のイベント・トレース機能、メモリのsymbolic表示等の機能拡張を要している。

関係データベースエンジンの開発 (その1)

△7

—ハードウェア構成—

岡 幸夫 松田 進
(株) 東芝 青森工場岩田 和秀 神谷 良雄
(株) 東芝 総合研究所

1. はじめに

通産省第五世代プロジェクトの一環として、関係データベースマシン Deltaのユニットのひとつである関係データベースエンジン(以下、RDBEと略記する)の開発を行った。

RDBEは、ソート処理を主体とした関係演算を専用ハードウェアで、また制御系の演算と全体の制御をソフトウェアで行うことにより、関係データベースの処理が効率よく行われるよう設計されている。本稿では、Deltaの概要と、RDBEのハードウェア構成の概要を報告する。

2. Deltaの概要

Deltaは、図1に示されているように、HMサブシステムとRSPサブシステムより構成され、ネットワークを介してホスト・マシンから送られてくるコマンドを実行する。RSPサブシステムはIP、CP、RDBE、MPの各ユニットから構成され、各ユニットが次のような機能を分担している。

IPは、LIAと、CP及びHMとのインタフェースをとるプロセッサ、CPは、RDBE、HMなどの主要なハードウェアのリソースを管理するプロセッサ、RDBEは、CPに制御されて、HMよりストリームで供給されるリレーションに対して関係演算を施すプロセッサ、MPは、Deltaの各プロセッサを監視するのを目的とするプロセッサである。

3. RDBEのハードウェア構成

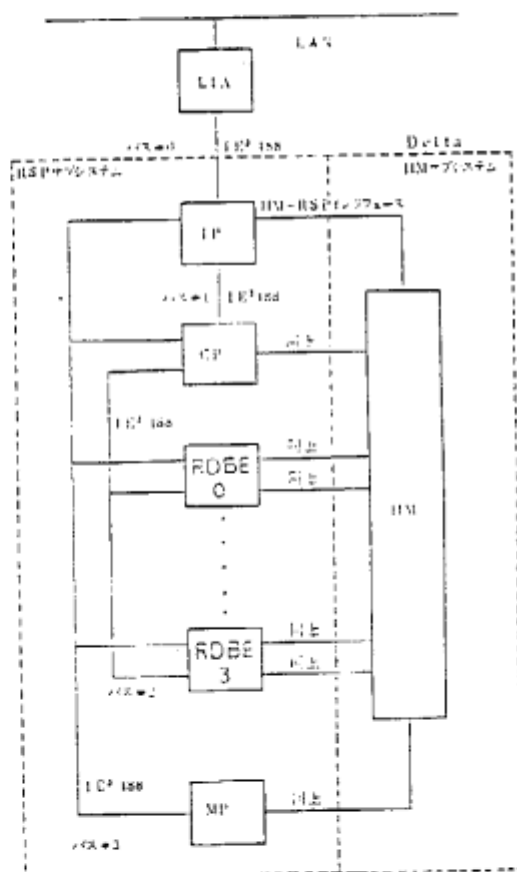
関係データベースの処理を効率よく行うためには、ソート処理や、Restrict、Join等の関係演算を高速に行う必要がある。このため、RDBEでは、ソート処理と、関係演算処理を速動して、入力データに開閉したパイプライン処理を行う専用ハードウェア(以下、エンジン・コアと略記する)を設計しこれをCPU及びその制御プログラムで制御する方式を採用した。

図2は、RDBEの概略構成図であり、点線で囲った部分がエンジン・コアである。エンジン・コアは更にINアライナ、ソータ、マージャの3つのモジュールから構成されている。INアライナは、入力データのフォーマット変換を行うハードウェア、ソータはソート処理専用のハードウェア、マージャは関係演算処理専用のハードウェアである。

ソータとマージャは2個のレコード間のKEYフィールドの比較を行い、比較結果が確定した時、該レコードの種類と比較結果に応じて、出力すべきレコードを確定する。したがって、ソータとマージャの各部では、KEYフィールドの比較が終了するまで、そのレコードをバッファリングする必要がある。しかし、パイプ

ライン処理の単位毎にレコードのバッファリングを行うことは、ハードウェアを大型化し、処理速度を低下させる。そこで、RDBEではエンジン・コアの入口にあるINアライナに1レコード分のバッファメモリを用意して、KEYフィールドをレコードの先頭に移動させるレコード・フォーマット変換を行うようにした。

(図3(a)、(b)を参照)



LAN : Local Area Network
LIA : LAN Interface Adaptor
RSP : Relational Supervisory Processing Subsystem
HM : Hierarchical Memory Subsystem
IP : Interface Processor
CP : Control Processor
RDBE : Relational Data Base Engine
MP : Maintenance Processor

図1 Deltaの構成図

ソータは、12個のソータ・セルと、ソータ・チェックより構成される。12個のソータ・セルは、2Wayパイプライン・マージソータ方式により、4Kレコードのソーティングを行う。また、ソータ・チェックは、1レコード分のバッファ・メモリを持ち、ソータの出力の連続した2レコードの大小比較を行うことにより、ソータ結果が正しいかを監視する。

マージは主に、2個のリレーション間の関係演算を行う専用ハードウェアであり、各リレーションを格納するために、2個のバッファ・メモリを用意している。マージは、ソータされたリレーション(A)を一方のバッファ・メモリに格納し、次にリレーション(B)がソータより転送されて来ると、他方のバッファ・メモリに格納しつつ、2つのリレーション間の関係演算を開始する。従って、マージは、エンジン・コアへの入力に同期した速度で関係演算を実行する。また、マージは、その出力部にリレーション(A)あるいはリレーション(B)に対応したレコードバッファを持ち、マージでの処理結果が出力される時にレコード・フォーマットの逆変換を行うことにより、エンジン・コアへ入力されたフォーマットに戻る。

本エンジン・コアでは、ソータ、マージを連続しているもので、マージに入力されるリレーションは、ソータによりあらかじめソーティングされている。マージは、これを前提として演算を行うので、関係演算処理の性能を著しく向上することができる。以上のように、エンジン・コアは、CPUの制御プログラムから、演算対象となるリレーションの総レコード数、レコード長、KEYフィールド長、KEYフィールドの位置、演算の種類、等の各種パラメータを受取り、各モジュールが連携して動作することにより、全体としてその機能を実現する。

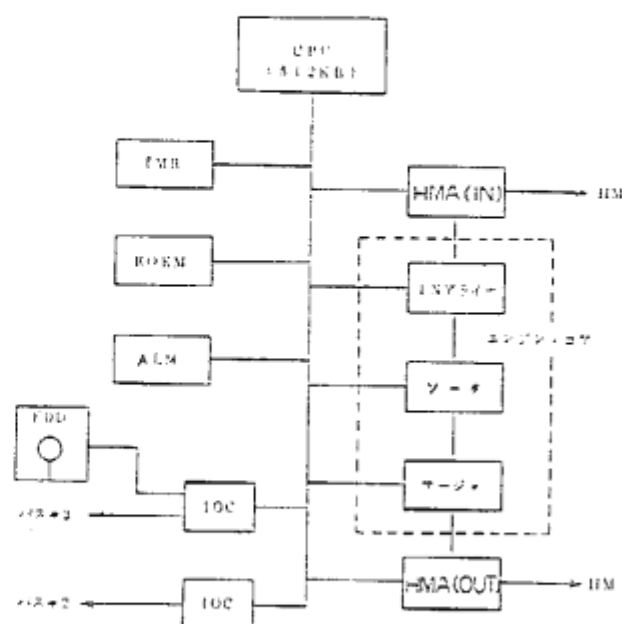
HMAアダプタ(以下、HMAと略記する)は、RDBEとHMとのインタフェースを制御する機能を有すると同時に、HMA(IN)は、エンジン・コアへのデータの入力ポート、HMA(OUT)は、エンジン・コアからのデータの出力ポートとして動作する。

HMAが、エンジン・コアの入出力ポートとして動作するときのデータ転送は、HMA/エンジン・コア間及びCPUのメモリ/エンジン・コア間で可能である。

本マージには算術演算系の機能を設けなかったもので演算演算が必要な時には、CPUでの算術演算機能を用いてその処理を行うように設計した。

4. おわりに

本稿では、Deltaの中核ユニットであるRDBEのハードウェア構成について述べた。RDBEを構成する各モジュールの詳細およびRDBEの制御ソフトウェアについては、別紙で報告する。誤謬・本誌に、ご指導いただきました「COT」KBMグループの方々に感謝いたします。



TMR : 計時装置
ROEM : 外部固定記憶装置
ALM : 警報装置
IOC : 入出力制御装置

図2 RDBEのハードウェア構成

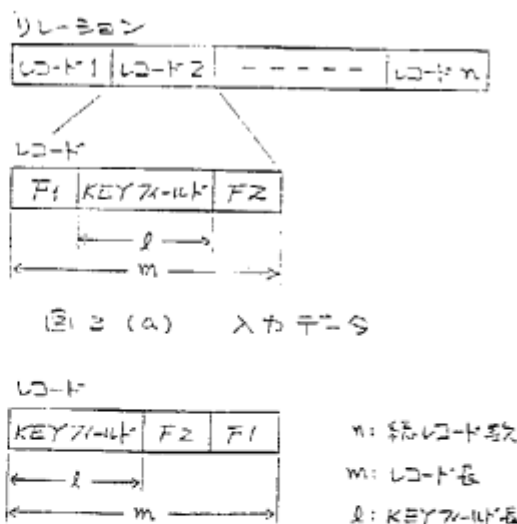


図3 (b) INアダプタの出力

〈参考文献〉

- [1] 角田、崇山、横田 他「RDB Delta(1)-(3)」情報処理学会全国大会予稿、PP4f-6-8, 1983
- [2] 神谷、徳井、安部、星野、伊藤 他「関係データベースエンジンの開発(その2-6)」本予稿、PP4f-6-10, 1984

関係データベースエンジンの開発 (その2)

4F-6

— ソー タ —

神谷 茂雄

岩田 和秀

酒井 浩

松田 進

(独逸 総合研究所)

(独逸 青島工場)

1. はじめに

通産省第5世代プロジェクトの一環として、関係データベースマシン“Delta”¹⁾のユニットのひとつである関係データベースエンジン²⁾の開発を行った。

本稿では、 n レコードのソートに対し、

ハード量: $O(\log_2 n)$

処理時間: $O(n)$

入力完了から出力開始までの遅れ: $O(\log_2 n)$ の性能を持ち、さらに、null値を含み、Keyが任意の位置にある最大4096バイト長のレコードに対し、昇順、降順の安定なソート処理ができるソータと、ソート処理がベースとなる関係代数演算と集合演算をソータと同速度でパイプライン的に処理できるマージャのハードウェア構成の概要を報告する。

2. ソータ

本ソータは、関係データベースの処理に応用するという観点から、以下に述べる点に留意して設計されている。

- (a) 関係データベースの処理では、レコードのKeyフィールドを対象としたソート処理と、集合演算のようにレコード全体を対象としたソート処理がある。本ソータは、この両方の処理に対応できるようにするため、ソートの対象となるレコード長を制御プログラムで指定できるようにし、2~4096バイト長のレコードの昇順、または、降順ソートできるようにした。さらに、Keyフィールド長も、2~4096バイトの長さのレコードに対してソート処理できるようにした。
- (b) 関係データベース処理では、Keyフィールドの値が同値の時と、null値の扱い方を考慮する必要がある。そこで本ソータでは、Keyフィールドの値が同値の場合は、レコードの入力順に出力する安定なソートを行い、null値を含む場合は、正常値のレコードの後にnull値のレコードを入力順に出力するようにした。
- (c) 関係データベースの処理では、種々のデータ形式が扱われるが、ソータのハードウェアは、2進数(符号なし)のみを扱うことにしている。そのため、整数と浮動小数点数は、ソータの入

力側(INアライナ)と出力側(マージャ)で、2進数へのデータ型式変換と、その逆変換を行うようにした。

- (d) パイプライン方式でソート処理を行うためには、レコードの先頭にKeyフィールドが位置している必要がある。そこで、Keyフィールドでソートする場合、ソータの入力側(INアライナ)と出力側(マージャ)で、Keyフィールドをレコードの先頭に移動させる処理と、元に戻す処理を行わせることにした。

このように、本ソータは、入力側にINアライナを、出力側にマージャを接続して、これらを連動動作させてソート処理を行わせている。そこで、2-wayマージ・ソート方式³⁾⁴⁾を採用し、sequential input/sequential output型のソータを構成した。

ソータ&マージャのブロック図を図1に示す。

本ソータは、12段のソート・セルからなり、これにより、最大4096レコードをソートできる。

また、ソート・チェッカは、ソート結果が正しいか否かを監視する。

3. ソート・セル

ソート・セルは、前段のソート・セルまででソートされている 2^{n-1} 個のレコードを2組入力し、 2^n 個のレコードをソートし、後段のソート・セルに出力する機能を有している。ここで、 n の値は、ソート・セルが置かれている入力側からの段数であ

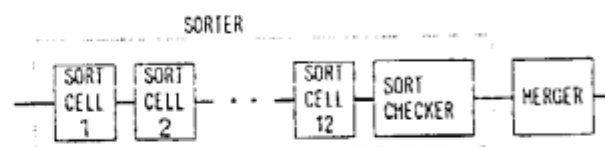


図1 ソータ&マージャのブロック図

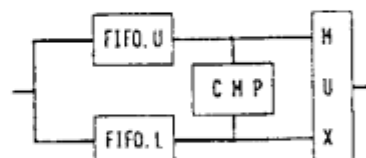


図2 ソート・セルのブロック図

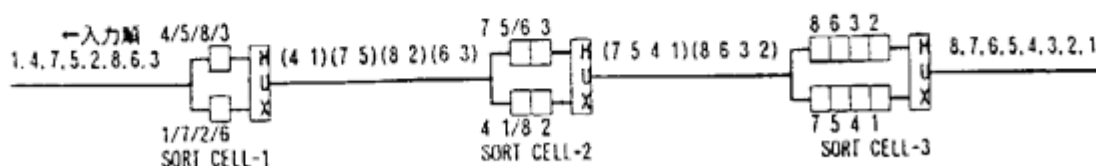


図3 ソート・セルの動作例

る。ソート・セルのブロック図を図2に示す。

ここで、FIFO, U/Lは、FIFO (First-In First-Out) 機能を有するメモリで、さらに、現在読み出し中のレコードの先頭から再度読み出し可能な機能も有している。n段目ならば、それぞれ、 2^{n-1} レコードを記憶できる容量を持つ。

CMPは、FIFO, U/Lから読み出されて来た値を比較する比較回路である。この比較結果により、いずれか側のレコードが選択回路MUXで選択され、後段のソート・セルに出力される。

次に、図3に8レコード (レコード長=2バイト) からなるリレーションの昇順ソートの動作例を示す。

4. マージャ

マージャは、ソート処理をほどこされている1つないし2つのリレーション間で関係代数演算や集合演算の基本となる命令セットを高速に実行する専用ハードウェアである。マージャのブロック図を図4に示す。

4.1 演算部

演算部は、CPUから指示された命令セットを実行するメイン・コンポーネントである。

Buffer, U/Lは、それぞれが1つのリレーションを記憶できる容量を持っている。2リレーション間の命令では、各レコードは同期して読み出されKeyフィールドの比較が行われる。

Buffer, U/Lは、レコード更新制御回路により制御されている。レコード更新制御回路は、FIFO機能と、現在読み出し中のレコードの先頭から、最初のレコードから等、前に遡って読み出す機能を実行する。

CMPs出力制御は比較回路と出力制御回路とからなり、前者は命令の種類と比較結果によりBuffer, U/Lのどちらか、あるいは、両方のレコードを出力すべきかを出力選択部に知らせる。

4.2 出力選択部

出力選択部は、演算部の補助動作として、演算部が出力対象指示を行ったレコードのみを選択出力する。出力する時、2進数に変換されているデータ型を元のデータ型式、つまり、整数や浮動小数点数に戻し、Keyフィールドの位置を元に戻す。

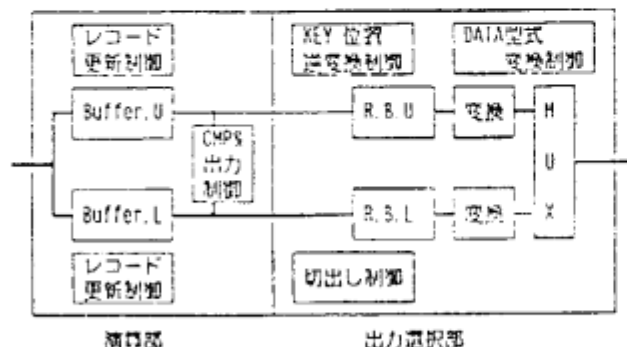


図4 マージャ・ブロック図

また、出力結果は、元のリレーションのレコードの全ての属性を必要とせず、その一部のみでよい場合が多い。そこで、必要な属性のみを切り出す処理「切出し制御」も行う。この指示がレコードの入力後にくるので、1レコード分のバッファR, B, U/Lでバッファリングする。

このように、マージャは演算部も、出力選択部も、ソート・セルと類似の回路で実現できるので、ソート・セルと同じ実行速度でパイプライン処理を行う。

4.3 命令セット

マージャで処理する命令セットを以下に列挙する。なお、²⁾ 関係演算、たとえば四則演算は、CPUが処理する。

(a) PASS系命令

対象リレーションをマージャに入力するか、出力するか、単に通過させるか、あるいは、重複を除いて出力する命令である。

(b) SORT系命令

ソータの容量を超えるソート、あるいは、重複を排除命令である。

(c) RESTRICT系命令

ひとつの対象リレーションの内から、条件に合うレコードのみを出力する命令である。

(d) COMPARE系命令

2つのリレーション間で順に比較し、条件に合うレコードの組のみを出力する命令である。

(e) JOIN系命令

2つのリレーション間でレコードの総当りを行い、条件に一致するレコードの組を全て出力する命令である。

5. おわりに

以上のハードウェア式は、TTLのHSI/SSIを用い、基板16枚で構成しており、既に完成し納入済みである。現在は、8000GateのGate Arrayを用いてソート・セルのLSI化を計っている。

謝辞 本開発に関し御指導いただきましたICOT KBMグループの方々に感謝いたします。

<参考文献>

- 1) 角田、栗山、横田他、「RDBM "Delta" (1)-(3)」情報処理学会第26回全国大会予稿、1983、4F-6~-8
- 2) 松田、酒井、安部、伊藤、星野他、「関係データベースエンジンの開発(その1、3-6)」情報処理学会第29回全国大会予稿、1984、4F-5、-7~-10
- 3) Todd "Algorithm and Hardware for a Merge Sort Using Multiple Processors" IBM J. R&D, vol 22, no5, September, 1978
- 4) 比田井他、「ソートをベースにしたデータベースマシン」情報処理学会第23回全国大会予稿、1981、4F-9

関係データベースエンジンの開発（その3） ～ハードウェアの制御方式～

▽ 西井 浩 岩田 和男 安部 公朗 神谷 良雄
(株式会社 東芝 総合研究所)

1 はじめに

通産省第5世代プロジェクトの一環として、関係データベースマシン"Delta"[1]のユニットの一つである関係データベースエンジン(RDBEと略記する)の開発を行った[2][3]。

本稿ではRDBE全体としての機能とそれを実現するためのハードウェア(HMアダプタとエンジン・コア)の制御方式について述べる。

2 RDBEの機能

RDBEの機能は、関係代数演算(JOINなど)の他、データのクラスタ化、更新などDeltaに必要な処理のうち、エンジン・コアによる高速処理が可能なもの、および利用可能ではないがDeltaでの複数プロセッサによる機能分散の設計方針からしてRDBEで処理すべきものを構築している。

RDBEの機能は、次のような5つ組で表現される。

(演算種別, 演算対象,

レコード切出し, TIO 付加, CPU による演算)

演算種別は、文字どおり演算の種類を表わすものである。表1にその一覧を示す。表1からわかるように演算種別は、関係代数演算そのもの、あるいはそれと同程度に論理的な演算からなる。そして、その大部分は、RDBEのエンジン・コアで高速処理される。ただし、エンジン・コアの機能上の制約により、Joinをはじめとする多くの演算では一度にひとつの比較条件しか処理できない。従って条件が複数の場合は、ひとつの条件による演算に引続いて、別の条件による演算を行うようにするか、あるいはCPUによる演算の部分に、ふたつめ以降の条件を記述する必要がある。

演算対象は、REで処理するレコードを格納しているHMのバッファID(高々ふたつ)、演算対象フィールドの情報(データ型など)、および結果を格納すべきHMのバッファID(常にひとつ)からなる。

レコード切出しは、射影演算に相当するものである。例えば一般にJOINという演算を行なう場合、もとのリレーションのすべての属性が必ずしも必要というわけではなく、そのうちの一部が必要であることが多い。そこでRDBEでは図1に示すように、もとのリレーションの一部だけを切出

演算種別	説 明	実現手段
Pass	CPU による演算等を使用	HW
Restrict	属性と定数との比較による選択	HW
Compare	属性どうしの比較による選択	HW
Join	θ -Join を行なう	HW
Sort	ひとつの属性について並べかえ	HW
Unique	ひとつの属性について重複除去	HW
H-Sort	複数の属性について並べかえ	HW
H-Unique	複数の属性について重複除去	HW, SW
Set	集合演算	HW, SW
CHECK	集合の包含関係の判定	HW, SW
Aggregate	集約演算	SW
Zone-Sort	クラスタ化に使用する。	HW
Delete	データの更新に使用する。	HW

表1 演算種別の一覧表

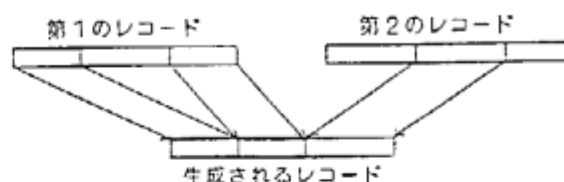


図1 レコード切出し機能

すことが、エンジン・コアでJoinと同時にこなせるように設計されており、ここでは切出す属性を指定する。なお、片方のレコードを切出さないように指定することにより、SEMI-JOIN が実現できる。

TIO 付加は、演算の結果生成される個々のレコードに、一連の通し番号を付加する機能を制御するものである。TIO 付加の有無とTIO の初期値が指定できる。

CPU による演算は、エンジン・コアの機能を補う目的で用意されている。詳細は[2]で報告する。

3 ハードウェアの制御方式

3.1 HHアダプタとエンジン・コアの制御方式

ROBEの機能は、エンジン・コアがほとんど処理しているので、エンジン・コアの制御方式だけを示せばよいように見える。しかし、実際にはエンジン・コアは入口に達したレコードを処理し、演算結果を出口に準備するだけであるので、それをHHとの間で転送するHHアダプタも含めて制御方式についても述べる必要がある。それらは次のとおりである。

HHアダプタ(IN)は、HHからデータを受取り、それをエンジン・コアへ送る。JOINなどの演算では、演算対象となるふたつのリレーション中のレコードを、ブロック・マルチプレクスして転送する。つまり、まず第1のリレーション中のレコードを一定個数まとめて受取り、それをエンジン・コアに渡す。そしてエンジン・コアが全部のレコードを受取り終わると、次に第2のリレーション中のレコードを受取り、それをエンジン・コアに送るという方式である。その切替えは、ROBE制御プログラムがHHアダプタ(IN)を制御することにより行なう。

エンジン・コアは、HHアダプタ(IN)から一定個数まとめて渡されたレコードを、処理対象となるフィールドの大小の順に並べかえて内部のバッファに格納したり、あるいは、先に内部のバッファに格納しておいたレコードと、今渡されて並べかえたレコードを順次大小比較することにより、HHアダプタ(IN)からまとめて渡されたレコードについて、JOINなどの演算を行なう。従ってエンジン・コアは、常にHHアダプタ(IN)と同期して制御すればよい。

HHアダプタ(OUT)は、エンジン・コアからデータを受取り、一定個数ずつまとめてHHに送る。エンジン・コアから出力されるレコード数は、演算の種類やデータの値の分布によって変わるので、HHアダプタ(OUT)の動作はエンジン・コアとは非同期となる。なお、HHアダプタ(OUT)は通常エンジン・コアからの出力をひとつのストリームとしてHHに送ればよいが、非常に大量のデータの並べかえの場合に限り、2ウェイ・マージを行なう都合により、2つの中間結果をブロック・マルチプレクスして転送する必要がある。その切替えは、ROBE制御プログラムがHHアダプタ(OUT)を制御することにより行なう。

3.2 大量データの処理方式

ここでは、エンジン・コアで一度に処理できない大量のデータの処理方式について述べる。処理方式は、次に示すように演算の種類に応じて4通りある。

(1) Pass型

レコード間の比較を必要としない演算では、エンジン・コアで処理できる量ごとにブロック化して区切ってHHから入力することにより処理する。すなわち、

for $i = 1, 2, \dots, n$

A_i の処理

ただし、 A_i は i 番目のブロックを表わしている。

(2) Join型

ふたつのリレーションに関する演算で、それぞれのリレーションの各要素のすべて対について調べればよい演算は、それぞれのリレーションをエンジン・コアで処理できる量ごとにブロック化し、すべてのブロックの対について処理する。すなわち、

for $i = 1, 2, \dots, n$

for $j = 1, 2, \dots, n$

A_i と B_j の処理

(3) Difference型

差集合 ($A - B$) を求めるような演算では、 B が大量の場合、 B をエンジン・コアで処理できる量ごとにブロック (B_i) に分け、 $((A - B_1) - B_2) - \dots - B_n$ を求めることにより処理する。

(4) Sort型

大量データの並べかえは、まずエンジン・コアの容量 (最大 128KB) ごとに並べかえを行なう。次に、2ウェイマージを繰返すことにより並べかえを行なう。

4 結論

ROBEは、関係代数演算をはじめとする関係データベースの処理を、主にハードウェアで高速処理する。そしてその指定方法は、ハードウェアの制御を意識する必要がないという意味で非常に論理的である。

ROBEのハードウェア制御方式では、エンジン・コアだけでなく、HHアダプタとの関連が重要である。特にエンジン・コアの容量を超える大量データについては、HHアダプタ(IN)がエンジン・コアで処理できる量ごとにわけて受取ることにより処理する。

謝辞

本開発に関し、ご指導頂いた ICOT KBHグループの方々に感謝いたします。

参考文献

- [1] 角田、柴山、横田他、「ROBE Delta (I) ~ (III)」, 情報処理学会第26回全国大会予稿, 4F-6~8, 1983
- [2] 安部他、「関係データベースエンジンの開発 (その4)」, 本予稿, 4F-8, 1984
- [3] 岩田他、「関係データベースエンジンの開発 (その1, 2, 5, 6)」, 本予稿, 4F-5, 6, 9, 10, 1984

関係データベースエンジンの開発（その4）
CPUによるデータ処理方式

秀郎公司 藤井 浩 岩田和秀 神谷深雄 (以上、東京綜合研究所)
 杉田 道 田澤文英 (以上、東京實施工場)

1 はじめに

係のEを、にぎを
関トBに、タていMで理
ッD、ドーし行H理処
てニRンデ力て処も
しユ(マ・入理いはで
とのンコスら処用でU
環”ジる一かでをけP
ーanくベMA)だC
のてエ。てタHエTアの
トー、たれーてウUコE
クスうらでいドO・B
エD一行送る用ー(ンD
ジ”ベをらなをハタジR
ロンタ晃かと)うブン、
プシー開P象Nいダエは)
代マデのC対Iとア、て照
世・無)の(アMたい查
5ス関。は理タコHまつ1
第一るるE処アコを、に(図
省べあるB、ダン果るの、
産タで記Dてアジ結すも、
通一つ略RいMN理力いう
デーと、づHエ処出な行



図 1 RDBE の処理の概要

本報告では，RDBEのCPUによるデータ処理とその実現方法について述べる。

2 ハードウェアとソフトウェアの機能分担

RはをコでC(1)(2)(3)(4)Cの現象にがド規

常か果るけの通M結くわE、HのてるBとをそれいDるタ。らてR取一い送し、受テ行ら現はド・をか実てドス理Pでアニー如CIAつマバで、エにコタアシウ能ら一コカド機。含う。濃たアをドに大さドにかデ・しーるるにぞ4則れ一箇一うを小一AAPるん。ハべいルか、四さル小ハよ畢がハコCなじるて速てブ値2の合ク最るる度を・とンすべにスター（点結）や舞する煩作をは象工送す次理のU算数で舞值演理工用演ジE対。転を、短つ消小OR演大の効巧便のン。つ。あはてで敷れてD効力MNなうひい舞浮AND集にれUI一舞れのなに要個さR。入HMはP(1)つ度(4)とこPハ演こ在と(1)つの定のはらをコでC(1)(2)(3)(4)Cの現象にがド規

- [illegible]

理るMるバ様の一
処あをすの仕極下
のがム理め能一
ム要テ始た機Uハ
テ必イ次ののM
イウア順えaやす
ア行、で替へ算理
てでは位をへ除始
のつて単ムe乗を
うかアトテDのら
伴向コイイ、しれ
をへ・バアたうそる
げ側ン2、まど、あ
上Bジへで。数りて
析Sン測のる然お替
は、MEBるなのて困
はらのしてとれは
舞かInして要イさ登
演測現らつて必バ定開
のB、かなが8根の
(2) Sし剣にア、もア
はしSよつて処ウ

(3) についても、比較演算をするためには、比較の種別とそれの数値を記憶するメモリが必要である。

たのがるるた
うけ値がちる。
行分の各がもめる
るブそて要は求ある
す一し必能をで
るてそる機難要
けグい。めの均必
分、つる求め平が
アてにあをたやリ
ーいドで値る値ム
ルツル要均す大メ
グに一必平現最
ずムイがや実のベ
まテフ定値を毎す
イの判大(2)ノ納
はアての最(2)一
てたベか、ルを
いっすうて(1)グ
つあるどい、結
にりなかつめと聞
(4)隣といにたこ中
に象しブののに
め対等一こんめ

3 CPUによる処理

CPUによる処理には、次のような種類がある。
(1) UNIQUE演算 (Null値同志は互いに異なる値と

- (1) UNIQUE 唯一 (Null 値とは互いに異なる値として扱う)。
これは、2 の (1) に対応する。
(2) アイテムは、内のフィールドに値する。変更。
これは、2 の (2) に対応する。
(3) アイテムは、指定された条件で絞り込み。
これは、2 の (2)、(3) に対応する。
(4) 集約演算。
これは、2 の (2)、(3)、(4) に対応する。

4 實現方法

4. 1 コンパイラ方式の採用

Eを二記
 Bドジ主
 Dコーン
 Rコーエ
 し・Bそ
 しはD
 折くれR
 解えこ、
 をジ。は
 ドブた
 マンオし
 コ能成イ
 コ能作ア
 した可を
 う行ラれ
 と美イさ
 け接バ力
 受向ン出
 らでこら
 CPサア
 の生・憶
 CC成コ
 の生・憶
 発-Compiler

4.2 テキスト処理方式

ここでは、条件に合うアイテムを出力する較り込みの処理を例にとり、CPUによるデータ処理の大まかな流れを示す。

- (1) CPから受取ったコマンドを解釈して、CPUでは、処理が必ず必要かどうかで、それを判断する。必要ならば、コンパイル・ジョブを生成する。
- (2) HMから、HMアダプタ(IN)を用い、処理

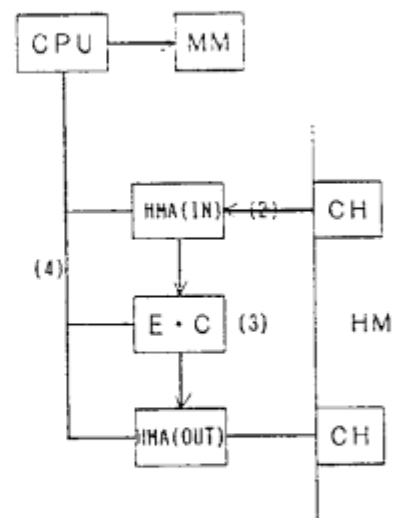
- 対象データベース(3)のエンジン・コアを記憶装置に格納し、主記憶装置に転送する。データは、主記憶装置からエンジン・コアに転送され、エンジン・コアで処理される。処理結果は、エンジン・コアから主記憶装置に転送され、主記憶装置からエンジン・コアに転送される。このように、データは、主記憶装置とエンジン・コアの間でやり取りされる。
- (3) データベース(3)のエンジン・コアを記憶装置に格納し、主記憶装置に転送する。
- (4) データベース(3)のエンジン・コアを記憶装置に格納し、主記憶装置に転送する。
- (5) データベース(3)のエンジン・コアを記憶装置に格納し、主記憶装置に転送する。
- (6) データベース(3)のエンジン・コアを記憶装置に格納し、主記憶装置に転送する。

5 おわりに

以上、RDBEのCPUによる処理とその実現方法について述べた。データベースのエンジン・コアを記憶装置に格納し、主記憶装置に転送する。データは、主記憶装置からエンジン・コアに転送され、エンジン・コアで処理される。処理結果は、エンジン・コアから主記憶装置に転送され、主記憶装置からエンジン・コアに転送される。このように、データは、主記憶装置とエンジン・コアの間でやり取りされる。

謝辞

本開発に関し、ご指導戴いたJCOT KBMグループの方々に深く感謝致します。



HMA(IN) : IN側のHMアダプタ

E.C. : エンジン・コア

MM : 主記憶

HMA(OUT) : OUT側のHMアダプタ

CH : ブロック・マルチプレクサ・チャンネル

図2 CPU処理以前のデータの流れ(その1)

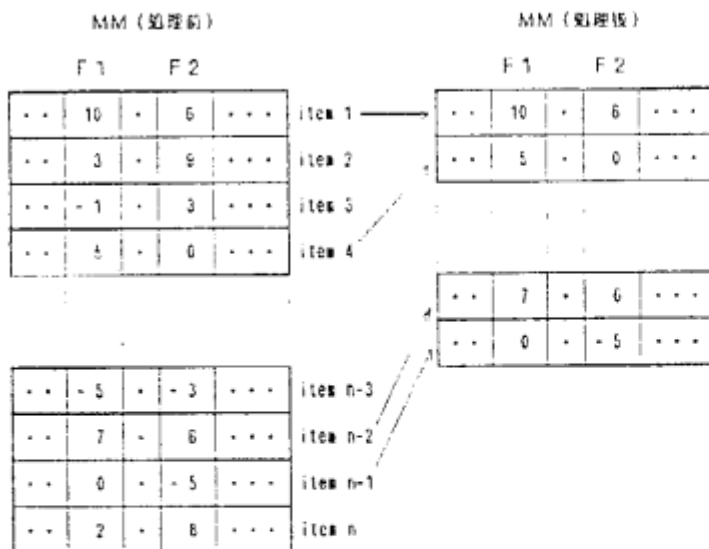


図3 CPUによる処理の例

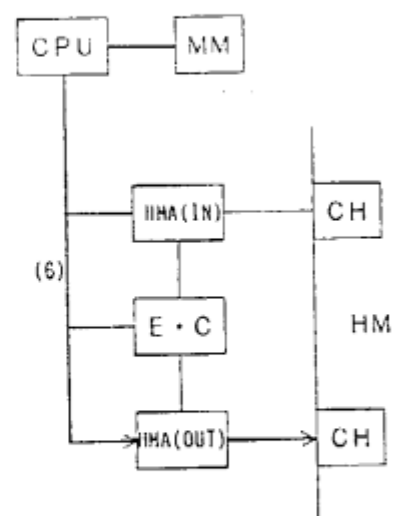


図4 CPU処理後のデータの流れ(その2)

関係データベースエンジンの開発 (その5) ～関係データベースエンジンに対する処理指定～

伊藤 文英 星野 昭夫
(株式会社東芝 資機工場)

沼井 浩
(株式会社東芝 総合研究所)

1. はじめに

通産省第5世代プロジェクトの一環として、関係データベースマシン“Delta”[1]のユニットの一つである関係データベースエンジン(RDBEと略記する)の開発を行った。[2]

本稿では、Deltaに対するホストからの問合せに基づき、コントロールプロセッサ(CPと略記する)からRDBEに対して、どのような処理が指定されるかについて報告する。Deltaでは関係を属性方向のデータの集まりとして格納するため、RDBEに対する処理指定にもその特徴が反映されている。

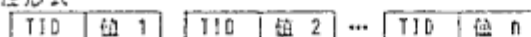
2. 処理されるデータの形式

RDBEが処理するデータのアイテムの構造を図1に示す。データの形式には次の2種類がある。

- ① 属性形式
- ② タブル形式

図1 データの形式

属性形式



タブル形式



各アイテムには、属性形式のデータの個でのタブル方向のつながりを示すために、タブル識別番号(TIDと略記する)がつけられている。TIDはDelta内でのみ使用され、ホストからは意識されない。

属性形式とタブル形式の間でのデータ形式の変換は、CPの指示により、データの置かれている階層構

造メモリサブシステム(HMと略記する)でおこなわれる。ただし、属性形式からタブル形式に変換する場合は、属性形式の各データはTIDで同様に並べられていなければならない。

3. RDBEにおけるデータ処理

RDBEにおけるデータ処理は次の2種類に大別される。

- ① ホストからのデータベース問合せに基づく演算
- ② 上記演算の入力となるデータ、あるいはホストに出力するためのデータを作成する演算

3.1 データベース問合せに基づく演算

Deltaがホストに提供するコマンドと、それを実現するためのRDBEに対する命令の対応を表1に示す。各演算は次のようにおこなわれる。

(1) 制約系演算

- ① 定数との比較条件による制約では、条件となる定数をCPから与え、属性形式のデータの中から条件をみたすアイテムを選択する。
- ② 属性同士の比較条件による制約では、TIDで同様に並べられた2つの属性形式のデータのアイテムを順次比較し、条件をみたすアイテムの組を選択する。
- ③ 算術演算の結果による制約では、条件中に含まれる属性を全て含むタブル形式のデータの中から、条件をみたすアイテムをCPUでの処理により選択する。
- ④ 準結合では、対象および条件となる2つの属性形式のデータを用意し、対象となる属性形式

表1 データベース問合せに基づく演算

コマンド	機能概要	RDBEに対する命令	入力データの形式
SELECTION	定数との比較条件による制約	定数とのJoin	属性形式
"	属性同士の比較条件による制約	Compare	属性形式
"	算術演算の結果による制約	Pass+CPU	タブル形式
JOIN	θ 結合、自然結合	Join	属性形式
SEMI-JOIN	準結合	Restrict	属性形式
集合演算系	和集合、共通部分集合、差集合	Set	タブル形式
統計処理系	総和、最大値、最小値、平均	Aggregate + CPU	属性形式
算術演算系	属性と定数との演算	Pass+CPU	属性形式
"	属性同士の演算	Compare + CPU	属性形式
検査系	集合の包含関係	Check	タブル形式
SORT	複数属性による並べ替え	H-Sort	タブル形式
UNIQUE	重複タブルの削除	H-Unique	タブル形式

のデータの中から条件をみたすアイテムを選択する。

(2) 結合系演算

θ結合、自然結合では、対象となる2つの属性形式のデータから、条件をみたすアイテムの組に新しいTID（結果リレーションのTIDとなる）をつけたものが出力される。

(3) 集合演算、検査、SORT、UNIQUE

これらの演算はタプルを1つの要素とみなすので、タプル形式のデータに対し、それぞれの演算をおこなう。

(4) 統計処理

属性形式のデータに対し、CPUで統計処理をおこなう。

(5) 算術演算

- ① 属性と定数の演算の場合は、定数をCPから与え、属性形式のデータに対してCPUで演算をおこなう。
- ② 属性同士の演算の場合は、TIDで同様に並べられた2つの属性形式のデータのアイテムを順次入力し、CPUで演算をおこなう。

3. 2 データを作成する演算

データを作成する演算には次のものがある。

- ① TIDによる制約
- ② TIDによる結合
- ③ TIDによるソート

(1) TIDによる制約

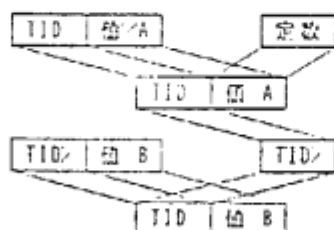
データベース問合せに基づく制約系演算は、条件となる属性のデータのみを使っておこなわれる。ある属性のデータにかかった制約を他の属性のデータに反映する場合は、制約がかかった属性のデータのTIDを条件として、他の属性のデータをTIDで制約する。

この様子を図2に示す。あるリレーションの属性Aにかかった制約を他の属性Bに反映させる場合は、制約後の属性AのTIDを使って属性Bを制約する。

図2 TIDによる制約

////: 比較フィールド

制約による制約



TIDによる制約

(2) TIDによる結合

データベース問合せに基づく結合系演算は、条件

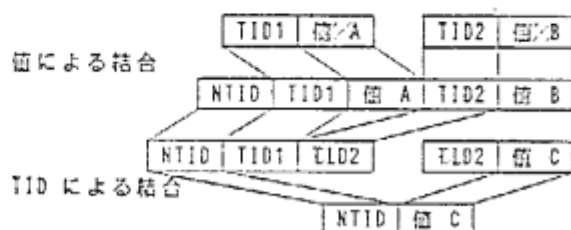
となる属性のデータのみを使っておこなわれる。結果として、条件をみたすアイテムのTIDの組に新しいTID（結果リレーションのTIDとなる）をつけたものが出力される。結合条件に含まれない属性の結果リレーションのデータを作る場合は、このTIDの組とその属性のデータの間にTIDによる結合をおこなう。

この様子を図3に示す。2つのリレーションは属性Aと属性Bに関する条件で結合されている。属性Bと同じ対象リレーションに含まれる属性Cの結果リレーションのデータを作る場合は、TIDの組と属性Cのデータの間にTIDによる結合をおこなう。

図3 TIDによる結合

////: 比較フィールド

値による結合



TIDによる結合

(3) TIDによるソート

TIDで同様に並べられた2つ以上の属性形式のデータを作成する場合は、それぞれのデータに対して、上記のTIDによる制約、結合をおこなったあと、TIDでソートする。タプル形式のデータを作成する場合は、全属性についてこの処理をおこなったあと、CPからHMに対してデータ形式の変換が指示される。

4. まとめ

CPからROBEに指定される処理には、次の特徴がある。

- ① Deltaでは関係を属性方向のデータの集まりとして格納するため、演算は必要な属性のデータのみを使っておこなわれる。
- ② ホストからのデータベース問合せに基づく制約による演算と、データを作成するためのTIDによる演算がある。

また、Deltaは知識ベースマシンのデータベースとして使われるので、知識ベース特有の問合せに対して効率のよい処理をおこなう方法を検討中である。

謝辞

本開発に関し、御指導頂いたICOT KBHグループの方々に感謝いたします。

参考文献

- [1] 角田、柴山、横田他、「ROBE Delta (I) ~ (II)」, 情報処理学会第26回全国大会予稿, 4F-6 ~ 8, 1983
- [2] 松田、神谷、酒井、安部、星野他「関係データベースエンジンの開発 (その1~4, 6)」, 本予稿, 4F-5 ~ 8, 10, 1984

関係データベースエンジンの開発 (その6)

MF-10

～エンジンの利用制御演算系の位置づけ～

鹿野 啓夫 伊藤 文英
(株)東芝 青柳工場

塩井 浩
(株)東芝 総合研究所

1. はじめに

通産省第5世代プロジェクトの一環として、関係データベースマシン“Delta”[1]のユニットの一つである関係データベースエンジン[2] (RDBEと略記する)の開発を行った。

本稿では、関係データベースエンジンについての報告の(その6)として、RDBEについてのまとめの意味で、関係データベースマシンDeltaシステム全体の中でのデータおよびその演算子群の概念的構成について述べ、RDBEの機能的な位置づけの説明をしたい。

Deltaは、ホストとのインタフェースプロセッサ(IP)、DB制御プロセッサ(CP)、エンジン(RDBE)、階層構造メモリ(HM)、Delta全体の監視プロセッサ(MP)を構成要素とする機能分散と、エンジンを複数個装備することによる負荷分散を組合せたシステムである。

2. Deltaが保持するデータの単位と演算系

Deltaでは関係データに関する関係演算にソート・マージアルゴリズムを用いるという目的から次のような特徴を持つデータ単位・演算系を持つ。(図1、表1参照)

(1) カラムデータ

Deltaでは、各リレーションを行方向(関係方向)に見てタブルを並べたデータ形式にするのではなく、カラム方向すなわち属性方向に分割して各属性毎のまとまりを単位として扱っている。

このような分割により、各属性ごとに独立したバリエーション化が可能となり、またそれぞれ独立してソート処理が行え、属性と属性の間でのマージ処理などにも効果的となる。このことは、属性を単位として単純化されたハードウェア向きアルゴリズムの演算子系が用意できることを意味する。

ただし、その見返りとしてタブルとしてのまとまりが論理的に保存されるように、それぞれの属性のそれぞれのバリューごとにタブル番号(タブルID)が必要となることも事実であり、“タブル形式への再構成”のような演算が増えることになる。しかしながら、これはハードウェアによる演算装置作成による高速化が可能になることから、パフォーマンスの点では十分補償されてしまう部分である。

このカラムデータは、階層構造メモリ装置に記憶される。他の制御系からはカラムデータの物理的構造とある程度の低レベルの論理構造を意識せずに、その生成、消滅、参照等の制御ができるようにしたい。そのために、階層構造メモリ自身がそのインデリジェンスとして、それらの制御用演算子系を装備することになった。これによりHM以外の制御系からは各属性に対応するカラムデータが抽象データ型と

して扱える。

(2) ディレクトリデータ

階層構造メモリは、リレーションを分解した一段プリミティブなカラムデータを抽象データ型化して扱っている。したがって、リレーションとしての意味単位に再構成できるためにはさらにカラムデータとリレーションとの対応関係や、リレーション自身に関する情報が必要であり、これらをまとめてディレクトリデータと呼ぶことにする。

ディレクトリデータに対する直接のアクセスは、その物理構造に依存した検索、生成、更新、消去等の操作であり、これらの操作用演算子系はDB制御プロセッサであるCP上に用意されている。

(3) リレーションデータ

(1)、(2)よりカラムデータとディレクトリデータに対するアクセスは独立した演算系によることになっているが、関係データベースマシンとして扱うのはリレーションデータであり、カラムとディレクトリの双方のデータを同時に関連させて扱うことにより始めてリレーションとしての意味的にまとまったデータ操作が行える。

したがってリレーションというデータ単位でアクセスを行うような演算系がDELTA内部においても必要である。この演算系は(1)、(2)で述べた2つの演算子系を組合わせさらに意味処理を行うことでリレーションという1段上のレベルのデータ型を抽象化している。

以上のデータはDelta内に保持され、ホストからの要求により消去されない限り存続するという意味で、パーマネント・データである。

3. 関係データベースエンジン

3.1 演算の過程で生成、操作するデータ

(1) サブタブル

ホストからDeltaへの関係演算指令の一つ一つに対してDelta内では種々の演算を組合せた演算子シーケンスを実行する。その実行過程では必要最小限個数の属性、即ち、カラムのみを組み合わせて進んでいくので、演算の最終目的の完成タブルから見ればサブタブルと呼ぶべきデータ形式が現れる。

サブタブルデータこそが関係データベースエンジンの演算対象となるものであり、サブタブル型を抽象データ型化する演算子のセットを提供している。(3.2節参照)

(2) サブディレクトリ

このサブタブルデータは演算の過程で作られる一時的なものであるがホストからの指示によっては新しいカラム型データとしてパーマネント化されたり、ホストからの次の関係演算指令で演算対象と指

図1 Delta 内データ・演算系の構成

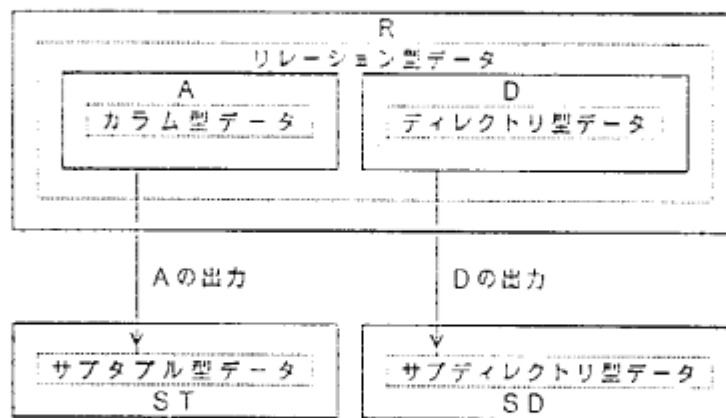


表1 各抽象データタイプの演算系の機能と演算担当プロセッサ

抽象データタイプ	機能	演算担当プロセッサ
A	カラム型	関係データベースエンジン (RE)
D	ディレクトリ型	DB制御プロセッサ (CP)
R	リレーション型	"
ST	サブタブ型	関係データベースエンジン (RE)
SD	サブディレクトリ型	DB制御プロセッサ (CP)

定されたりする。

このためにこのサブタブ型データに対してもそれ自身のスキーマ構造や、元になったリレーション、演算の履歴等の情報を保持しておかなければならない。これは、2章で述べたディレクトリ型データを元にして作成されるものであり、サブディレクトリと呼ぶことにする。

サブディレクトリにはディレクトリデータと同様のスキーマ情報も含まれるが、前述のような固有の情報が、ディレクトリデータとは異なる操作用演算子で用意されている。

以上のデータ型は、2章で述べたパーマネントデータを元にして生成されたデータであり、トランザクションという仕事単位の間だけ存在するテンポラリデータである。

3.2 エンジン演算系

エンジンは、このサブタブ型データを抽象化する演算制御系プロセッサとして位置づけられている。(図1、表1参照)

ホストからDeltaへ送られた各関係演算コマンドはDelta内部演算子の列(すなわち図1、表1の演算子の列)に変換される。DB制御プロセッサがその演算子列の実行制御を行い、各演算子をその属するプロセッサに指令して実行させる。エンジンに対しては、関係演算の下位レベルであるストリーム演算系の演算子が指定され、階層構造メモリ上のサブ

タブ型データに対する操作を行う。すなわちRDBEとHHとを合せると、演算子系とデータ実体の双方を包含しているという意味で、サブタブ抽象データ型マシンと呼べる。

4. おわりに

Delta内でのデータと演算セットの組み合わせをデータタイプの抽象化の観点から整理してみた。現実のインプリメントでは関係演算系以外の種々の制御や例外処理的な機能をサポートするためにデータタイプの抽象化の考え方から逸脱した部分もあるが、全体としてはほぼ本書の内容に沿ったものとなっている。

謝辞

本開発に関し、御指導頂いたICOT K&Mグループの方々に感謝致します。

参考文献

- [1] 角田、柴山、横田他 「RDBH Delta (1) - (3)」, 情報処理学会第26回全国大会予稿, pp 4F-6-8, 1983
- [2] 松田、神谷、酒井、安部、星野他 「関係データベースエンジンの開発(その1-5)」, 情報処理学会第29回全国大会予稿, pp4F-5-9, 1984

PROLOGソースレベル・オブティマイザ

ユ・フ・シ

— 決定性の検出とその利用 —

沢村・一・竹島 早・加藤昭彦

(富士通国産研)

1. はじめに

Prologプログラムをソースレベルで最適化(改良)するPrologソースレベル・オブティマイザを試作した[1, 2]。Prologはバックトラック機構に基づく非決定性プログラミング言語であるためプログラムのデータ及び制御の流れは極めて複雑である。

それゆえ、Prologプログラムの複雑なふるまいを反映したプログラムの最適化規則には様々な前提条件が必要になることが多い。このような前提条件の中で、オブティマイザを構成する上で最も重要であったのは述語の決定性である。

しかしながら、述語の決定性を一般的に検出することは不可能と思われるので、本稿では決定性の部分クラスとしてのc-決定性なる概念を定義し、その次のような最適化手法への利用について述べる：

- (i) 部分評価(インライン展開、局所的最適化手法)
- (ii) カットの自動挿入

2. 決定性とc-決定性

[定義 1] 決定的述語

述語が決定的であるとは、その定義節の両方一つの節で成功するのみで、バックトラックしたとき再び成功することはないことをいう。

[定義 2] c-決定性

述語がc-決定的であるとは次の条件を満たすときをいう：

- (i) 組み込みの決定的述語はc-決定的述語である。
- (ii) 述語名 H に関する定義体を次のものとする：

$$H1 :- \Gamma 1.$$

$$\vdots$$

$$Hi :- \Gamma i, [!,] \Delta i., \text{ここで、!は最右とする}$$

$$\vdots$$

$$Hn :- \Gamma n.$$

このとき、次の条件が成立するならば H はc-決定的述語である；各 i に対して、

- ① H_i の本体にカットが存在しないとき：

$\Gamma i, \Delta i$ はc-決定的で、さらに H_i は定義節の最後の節である。

- ② H_i の本体にカットが存在するとき：

Δi はc-決定的である。

[系] c-決定的述語は決定的述語である。

[注意] c-決定性はインライン展開[1, 2]によって保存される。

[例 1] c-決定的述語の例

$$p(a) :- \text{write}(a1), nl, !, \text{write}(a2).$$

$$p(b) :- q, \text{write}(b).$$

3. c-決定性の最適化への利用

我々のPrologオブティマイザではインライン展開を次のような自然な方策の下に行っている：

「述語の呼びを、呼出機構を渡す等式を付随した代替節の導言によって置き換える」。

しかしながら、これが安全に行えるのは呼ばれる側の述語の定義体の中にカット記号が現れないときである。というのは、バックトラックが起こったとき、展開前と展開後のプログラムではカットによってその振る舞いに違いが生じるからである。たとえば、

展開前： $p :- q, a, r.$

$$p.$$

$$a :- b, !, c.$$

$$a :- d.$$

展開後： $p :- q, (a = a, b, !, c; a = a, d), r.$

$$p.$$

$$a :- b, !, c.$$

$$a :- d.$$

において、展開前のプログラムで c が失敗したとすると制御は q に移るが、他方展開後のプログラムでは制御は p の呼びを失敗させることになる。

(i) インライン展開における利用

このような問題を避け、上記のようなインライン展開を可能にする一つの方法はc-決定性の概念を用いることである。すなわち、次の最適化図式が成立する。

$$a1 :- \Gamma 1.$$

$$\vdots$$

$$ai :- \Gamma i, p(\Delta i), \Delta i.$$

$$\vdots$$

$$an :- \Gamma n.$$

$p1(\Delta 1) :- \Theta 1., \Theta i (1 \leq i \leq m)$ のいずれかにカットが含まれているものとする

$$pn(\Delta n) :- \Theta m.$$

$$\vdots$$

$$a1 :- \Gamma 1.$$

$$\vdots$$

```

a1 :- Γ1, ( p(A) = p1(A1'), Θ1' ; ... ;
           p(A) = pm(Am'), Θm' ), Δ1,
...
a2 :- Γ2,
p1(A1) :- Θ1,
...
pm(Am) :- Θm.

```

ここで、 $A1'$, $\Theta1'$ はそれぞれ $A1$, $\Theta1$ をリネームしたものを表す。さらに次の条件を満たす:

① $\Gamma1$ の中にカットが存在しないとき:

$\Gamma1$ の中のすべての述語は c -決定的であり、 $a1$ は定義節の最後の節である。

② $\Gamma1$ の中にカットが存在するとき:

$\Gamma1$ の中の最も右にあるカットの右側において $\Gamma1$ に含まれるすべての述語は c -決定的である。

【例 2】 カットを含む定義体によるインライン展開

```

q(a) :- !, write(a), p(X),
q(b) :- write(b), p(Y), write(Y),

```

は例 1 の p によって次のように展開が可能となる。

```

q(a) :- !, write(a), ( p(X) = p(a),
                       write(a1), !, write(a2);
                       p(X) = p(b), q, write(b) ),
q(b) :- write(b), ( p(Y) = p(a),
                       write(a1), !, write(a2);
                       p(Y) = p(b), q, write(b) ),
                       write(Y).

```

(2) カットの自動挿入における利用

バックトラックしたとき再び成功することがないことが分かっているとき、適切な位置にカット記号を挿入しておけば不必要な redo を避けることができる。上の図式に認識して次のようなカットの自動挿入が可能となる。

(i) 上の図式において条件が満たされているとき:

$\Gamma1 \Rightarrow \Gamma1, !$

とする。すなわち、次の最適化図式が成立する

```

a1 :- Γ1,
...
a1 :- Γ1, !, ( p(A) = p1(A1'), Θ1' ;
...
           p(A) = pm(Am'), Θm' ), Δ1,
...
an :- Γn,
p1(A1) :- Θ1,
...
pm(Am) :- Θm.

```

```

a1 :- Γ1,
...
a1 :- Γ1, !, ( p(A) = p1(A1'), Θ1' ;
...
           p(A) = pm(Am'), Θm' ), Δ1,

```

```

an :- Γn,
p1(A1) :- Θ1,
...
pm(Am) :- Θm.

```

(ii) さらに p が c -決定的であるとき:

$\Delta1 \Rightarrow !, \Delta1$

とする。すなわち、次の最適化図式が成立する

```

a1 :- Γ1,
...
a1 :- Γ1, p(A), Δ1,
...
an :- Γn,
p1(A1) :- Θ1,
...
pm(Am) :- Θm.

```

```

a1 :- Γ1,
...
a1 :- Γ1, !, ( p(A) = p1(A1'), Θ1' ;
...
           p(A) = pm(Am'), Θm' ),
           !, Δ1,
...
an :- Γn,
p1(A1) :- Θ1,
...
pm(Am) :- Θm.

```

以上の最適化図式の妥当性のチェックは現在の Prolog のオペレーショナルな意味 (インタプリタ) の下に確認している。さらに形式的な意味論の下での同値性の証明は今後の研究をまたねばならない。

4. あとがき c -決定性は日常の Prolog プログラミングにおいてもよく現れ、そのチェックも容易に行える決定的述語の一つのクラスを定義している。一般に、述語の決定性は単にそれが Prolog のような非決定性プログラム言語の最適化 (改善) に有効になるばかりでなく、効率の良いプログラムを書くときの一つの指針ともなるものである。

本研究の一部は第 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。

日頃御指導、御教誨いただく北川勉男会長に感謝いたします。

参考文献

- [1] 沢村・竹島・加藤: PROLOG ソースレベル・オブティマイザ—最適化手法のカatalog—, Prolog Conf. '84.
- [2] 竹島・沢村・加藤: Prolog ソースレベル・オブティマイザ—インライン展開を中心として—, 本大会予稿.

42-6

PROLOGソースレベル・オブティマイザ — インライン展開を中心として —

竹島 卓・沢村 一・加藤 昭彦

(富士通国研)

1. はじめに

Prologプログラムをソースレベルで最適化(改良)するPrologソースレベル・オブティマイザを試作した(1, 2)。Prologのような論理型言語の最適化に対する研究は未だ少ないので、我々がPrologプログラムをソースレベルで最適化するに当たって考察した事をまず論じ、ついで具体的方針としてのインライン展開を中心とした最適化について述べる。

2. Prologの非論理性に起因する最適化の困難性

Prologプログラムを最適化しようとするとき、Prologの次のような実行メカニズム、言語要素は論理的な最適化をかなり制限する原因となっている。

① Prolog特有の実行メカニズム： 上一下、左一右、バックトラック、とくにcutによるバックトラック制御；これらは論理式レベルでの自由な最適化を阻害する。

② 組み込みのevaluable述語： side effectを持つ述語(入出力述語)、メタ述語、カット、否定、repeat(トートロジー)、etc.；これらはオブジェクトとメタ言語をともに考慮に入れた最適化を必要とする。

3. 最適性の判断標準

最適化の有用性の尺度、すなわち何を基準にして最適化されたかということを確認しておくねばならないことは言うまでもない。次の三点から考察する。

(1)外的規準(処理方式、環境) (2)内的規準(表現)

(3) 計算方式

最適化方法といわゆる複雑性の研究は目的とするところはよく似ている。しかしながら、複雑性の研究では(1)、(3)は問題とならない。プログラムの最適化問題ではオーダよりもむしろ、係数の改良に主張が置かれるので、(2)のみでなく(1)、(3)も最適化を考えるときの考察対象とされるべきであろう。

4. プログラム変換論とプログラムの最適化

狭い意味においては、プログラムの変換論とプログラムの最適化の目的は同じと考えられるが、広義には、プログラムの変換論とは一つのプログラミングparadigmであって、stepwise refinement (Dijkstra, Wirth) の概念にこれまでの伝統的な最適化技術を結びつけたものを意味することが多い。そして、このparadigmの下では、我々は非効率的であるかもしれないが、明確で正しいプログラムを書

き、それから、それを同値性を保存する変換を経て、明確さには欠けるかもしれないがより効率的なよいプログラムへと変換する。

以下で述べられるProlog言語の最適化はPrologのプログラミングの規範を示すことでもなく、またeurekaを多用するdrasticな変換でもない。それはどちらかと言うと伝統的な従来言語の最適化の精神に通じるものと言える。従って、プログラムの変換論とは明確に立場が異なっていることを強調しておくねばならない。

5. 最適化の方針(インライン展開を中心として)

我々のPrologの最適化問題に対する取組方は実用及び経験主義的である(少なくとも演えきめではない)。

従来言語の最適化技法の中でインライン展開(サブルーチン展開)は最も重要な最適化法として知られている。その大きな目的は、①起動機構(サブルーチンの呼出機構)の解消、および②他の最適化技法の適用可能範囲の拡大にあるとされる。

一方、Prolog言語に対してのインライン展開はつぎのような特徴をもっている。

(i) Prologは述語(手続き)の列からのみ書かれる言語であって、すべての実行文が展開の対象となってしまう。

(ii) Prologではそもそも呼び出し機構は述語定義の探索とユニフィケーションの二つからなり、その呼び出し機構を完全に解消することは一般に不可能である。すなわち、述語の呼び出し機構を呼び出しの述語とその述語定義の単一化可能性にすりかえる必要がある。

これらはProlog特有の言語形式、計算方式に依っている。(i)はPrologの最適化を考えるさいにインライン展開がその出発点となるべきことを暗示するし、(ii)は起動機構を新たな表現によりプログラムテキストの中に表現することを要求する。

実際、最も素朴な場合のゴールGの展開形式はつぎのようになる。

$$\begin{array}{l} H1 : -\Gamma 1. \\ H2 : -\Gamma 2. \\ \vdots \\ G \quad Hn : -\Gamma n. \end{array}$$

$$G = H1', \Gamma 1' ; \dots ; G = Hn', \Gamma n'.$$

ここで、 H_i は H_i のremaining termである。時出候補の一部は $G = H_i$ として残存する。

このようなインライン展開によって変えさせられたプログラムに対していく種類の最適化手法が適用されることになるわけであるが(展開して式を長く保っておく他最適化を受けさせるのにいつも都合がよいというわけではないが)、適用の順序が問題になり、個々の最適化手法が有効であってもその順序を誤るとせっかくのインライン展開の効果は半減してしまうことになりかねない。

オブティマイザに関する研究の現状からいって、2節で論じた最適化に関する判断基準すべてに適合する最適化システムを構成することは不可能に近い。ここでは理想的な最適化システムの実現に向けて、first stepとしての方針を取り、つぎのような点に留意してシステムを試作した。

(a) 展開レベル

現在我々のいく種類の最適化手法を総合的に見ると、インライン展開は帰納的なプログラムの場合1レベルの展開で十分であると考えられる。それ以上展開しても他の最適化手法によってdrasticにプログラムが最適化されるという期待は薄いようである。また展開のし過ぎはメモリーの過大な消費による負荷が高くなる。

(b) 時間 vs. 空間効率

現段階の研究の現状からいって、時間効率にのみ焦点を合わせた最適化システムの構成とし、空間効率の評価はシステムの使用経験を踏まえて考察していくことにする。

(c) interactive vs. automatic (システム構成に当たって)

Prologオブティマイザへの入力としてfullなPrologのソースプログラムを対象とするために、適時最適化のための指示を使用者が与えるという方式にする。

5. Prologソースレベル・オブティマイザの概要

Prologソースレベル・オブティマイザはつぎの3つの主要機能からなる。

① 最適化に必要な情報をプログラムテキストから抽出する機能

② 各種最適化手法に対応する機能

③ 入出力機能

オブティマイザの全体の処理はユーザからの指示に従って、直線・終端再帰、一般の帰納的プログラムの各インライン展開(3.2節参照)の前段に各種の個別の最適化手法(3.3節参照)を適用する形で進行する。

直線プログラムのインライン展開以外のインライン展開では、そこでさらに最適化を進めるか、終了させるかの判断を人間にもとめるようにしている。

ここで各種の最適化手法とはつぎのものである。

i. ユニフィケーションの部分実行 等式の形をしたゴールの計算を部分的に進行させる。

ii. ゴールの簡単化 運営中あるいは運営中の置換ゴール

ルの除去、冗長なtrue, failの除去、不実行部分の除去、共通ゴールのくくりだし。

iii. 置換除去 冗長置換除去、等式代入。

iv. ゴールの統合化 複数の等式ゴールを統合する。

v. 節の分解 選言をもつ節を分解する。

これらの最適化手法は無条件で適用できるとは限らず、所要の条件を満たしている必要がある。基本的な条件はプログラムの形式(直線、終端再帰、一般)に依存するものでありこの形式の認識が必要である。とくに重要な条件はプログラムの決定性である。決定性情報によらない最適化は極く小さなプログラムのクラスにかぎられる。決定性検出とその利用については別に本大会で詳述している。

7. あとがき

主として経験的かつ直観的にPrologプログラムのソースレベルでの最適化手法について検討した結果、多くの有効な最適化手法が見いだされた。それらは従来言語の最適化手法と比べて論理的には複雑である。我々はPrologプログラムの変換論とは異なる立場から検討してきたが、各種最適化手法が同値性を保証していることの形式的な証明については今後の課題としてまだ触れていない。

我々のPrologプログラムのインライン展開の方法とBurSTALLらのプログラムの変換論におけるいくつかの概念との間にはいくつかの表面的な類似点が見いだされることについて触れておかねばならない。例えば、プログラム変換論におけるunfoldingはその行為だけではインライン展開と同じであるが、unfoldingは後に abstraction によってfolding されることを意図して行われるのでインライン展開の主要目的①、②とは明らかに異なる。またfoldingは我々の最適化手法では等式ゴール列の統合化(3.5章)といくらか似ているが、我々の方法では現在のところ abstraction という発見的手法を用いることはしていない。

他方、インライン展開もunfoldingも形式面から見ると共に計算を部分的に進行させているわけであるから、partial evaluationの異なる形態であるとも見なせる。

本研究の一部は第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。

参考文献

- (1) 沢村・竹島・加藤: PROLOGソースレベル・オブティマイザ最適化手法のカatalogー, Prolog Conf. '84.
- (2) 沢村・竹島・加藤: Prologソースレベル・オブティマイザ決定性の検出とその利用, 本大会予稿.

PROLOGインタプリタの構成

記 名 長 野 隆 司 藤 本 剛 木 村 康 則 岸 本 光 弘
(富士通株式会社)

1. はじめに

LISPマシンALPHA (3) の上に、PROLOGマシンアーキテクチャの研究と高性能PROLOG処理環境の実現を目的として、PROLOGインタプリタシステムを開発したので報告する。

2. 開発のねらい

このPROLOGインタプリタシステムの主な開発のねらいは次の通りである。

(1) PROLOGの高速処理 (高速インタプリタ)

PROLOGの実行に適した内部表現やデータ型を導入し、インタプリタの実行部やガーベジ・コレクタをファームウェア化することにより、高性能インタプリタを実現すること。特に、4項で述べる各々の高速化技法を用いて、実用プログラム開発に必要な高速性を実現すること。

(2) PROLOGマシンアーキテクチャの評価

PROLOGマシンアーキテクチャ (実行制御、メモリ制御等) に関する基礎データを収集、評価すること。

(3) LISPシステムとの最適インタフェース

PROLOGの世界からLISP機能が述語と同様に使えること。

(4) DEC-10PROLOG互換

標準的PROLOGであるDEC-10PROLOGとデバッグ機能も含めて構文、機能の互換性があり、実数もサポートすること。

3. 処理系のシステム構成

図. 1にPROLOG処理系のシステム構成を示す。リーダー、プリンタ及び組込述語の大部分はLISPで記述されている。PROLOGの実行部 (リゾルバーと呼ぶ) やガーベジ・コレクタはファームウェア化されている。PROLOGドライバルーブ及びユーザ定義述語は図. 2のようなベクトル型に内部表現されている。

(1) PROLOGの起動

LISPシステムから、PROLOGのドライバルーブを初期ゴールとしてリゾルバーに渡すことにより、PROLOGシステムがスタートする。

(2) 入出力系

リーダーはバックトラックを含む変形SLR (k) パーザでDEC-10PROLOGの全構文をサポートしており、端末やファイルから外部表現を読み込んで、高速処理に適した内部表現へ変換する。

プリンタは内部表現を外部表現へ変換して端末やファイルへ出力する。

(3) リゾルバー

リゾルバーはユニフィケーション、組込述語コール、リターン、バックトラック等の機能を持ち、内部表現に変換されたドライバルーブや入力ゴールを組込述語やユーザ定義述語を用いてリダクション処理するマイクロルーチン (2K語) である。

(4) LISPシステムとのインタフェース

リゾルバーからLISP関数がコールでき、関数UNIFYを用意している。それにより、LISP関数との論理変数を介した通信が高速に行える。

(5) ガーベジ・コレクタ

ガーベジコレクタはファームウェア化 (約2K語) されており、LISP、PROLOG両システムに共通である。LISPオブジェクトの他にグローバルスタック、トレイルスタックの回収、コンパクション機能を持つ。

(6) スタック

DEC-10PROLOGと同様にローカル、グローバル、トレイルの3種のスタックを用い、構造データの表現はStructure-Sharing法 (1) を用いている。ローカルスタックは

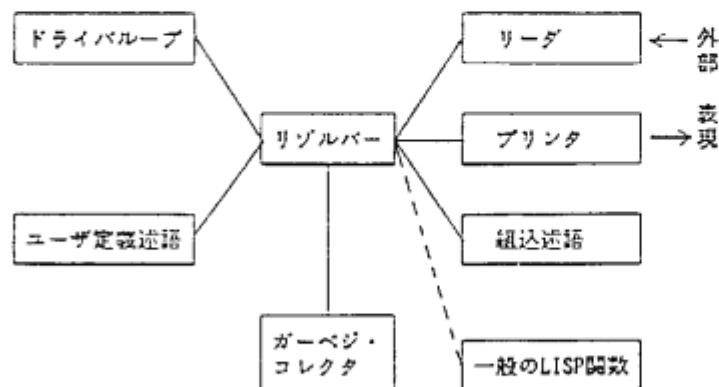


図. 1 PROLOG処理系のシステム構成

使用頻度が高く、スワッチ技法も使えるのでCPU内蔵のハードウェアスタック（実容量8K語、仮想容量64K語）に割り付けている。グローバル、トレイル両スタックは大量のため主記憶に割り付けている。

(7) 述語及びクロースの内部表現

述語は述語名シンボルの型化リストに引数個数インデックスして格納している。ユーザ定義述語の各クロースは図2のようにベクトルとして内部表現している。一方、大部分の組み込み述語はLISPのSUBR関数となっている。

4. 高速化技法

このPROLOGインタプリタは、次のような方法を用いて高速化をはかっている。

- (1) ローカルスタックのハードウェアスタック割付
- (2) CONSをLISPと同様のリスト表現

他のスケルトンはベクトル型の内部表現を用いているが、CONSは特に使用量が多いのでLISPと同様のリスト表現を用いてアクセスの高速化とメモリ使用量の効率化をはかっている。

- (3) 定数CONS、定数スケルトンの導入

変数を含まない構造体専用のデータ型を導入し、不要なモレキュール作成を省略している。

- (4) トレイルスタックへの不要なプッシュの省略

ローカル、グローバル両スタックのバックトラック点をレジスタ上に保持することにより、バックトラック点以降の変数への束縛時のトレイルプッシュを省略している。

- (5) Deterministic 述語の高機能リターン処理

Deterministic 述語が終端ゴールの場合、次に実行すべきゴールの親フレームとバックトラックフレームの内、新しい方のフレームの下に次ゴール用のフレームを作成す

る。終端ゴールでない場合には、親フレームを流用する。これにより、終端ゴールでのローカルスタック長の短縮と非終端ゴールでの高速化をはかっている。

5. 性能

本PROLOGインタプリタによるD.H.D.Warrenの3種のベンチマークプログラムの実行時間を下表に示す。

NATIVE REVERSE-30	35.9 msec
QUICK SORT-50	45.6 msec

この性能はDEC-10PROLOGインタプリタ（DEC C2060）に比べ、5倍以上高速である。

尚、4項で述べたCONSのリスト表現及び定数CONS、定数スケルトン型の導入は上記プログラムの場合、6～7%高速化に寄与している。4項(4)、(5)の効果は更に大きな応用プログラムでの評価が必要である。

6. おわりに

LISPマシンALPHA上に高速PROLOGインタプリタシステムを開発した。今後、各種応用プログラムでの推論マシンアーキテクチャに関する基礎データの収集、評価を行う予定である。

尚、本研究は第5世代コンピュータプロジェクトの一環としてICOTの委託で行ったものである。

謝辞

日頃御指導頂く御橋部長、林室長並びに本研究に御協力頂いた関係諸兄に深謝致します。

参考文献

- (1) D.H.D.Warren: Implementing Prolog.
D.A.I. Research Report no.39-40 (1977).
- (2) G.M.Roberts: An Implementation of PROLOG.
A thesis of Waterloo University (1977).
- (3) 服部、林、秋元、丹羽、品川、徳木、木村、佐藤:
ALPHA-高性能LISPマシン、情報処理学会
記号処理研究会資料23-1 (1983).

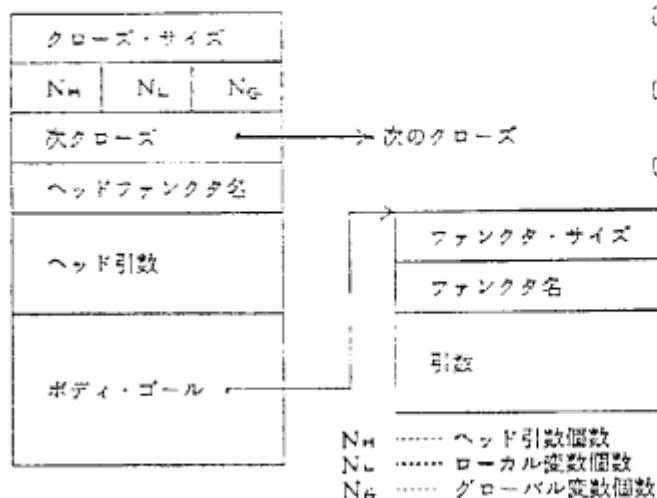


図. 2 クローズ及びファンクタの内部表現

73-1

PROLOGの並列処理システム

板倉 具弘

佐藤 健

増沢 秀徳

(富士通株式会社)

1. はじめに

我々はプロログを並列に処理する方式を研究している。前回の全国大会では、プロログを節単位に処理するモデルを提案した⁽¹⁾。その後、このモデルに基づいた並列処理システムを試作した。

今回は、試作した並列処理システムについて、特徴・構成・結果を報告する。

2. モデルの特徴

モデルの特徴については文献(1)に詳述してあるので、ここでは列挙するに留める。

- ・OR並列処理方式を採用する。AND関係は逐次に処理する。
- ・リテラルとユニフィケーションが成功する節単位で、UNIFY プロセスを作る。
- ・UNIFY プロセスに解を保持する機構を付ける。

3. システムの構成

今回試作したシステムは、プロセッサ要素としてパーソナルコンピュータを複数台使い、互いに他の総てのプロセッサとデータと通信できる形態とした。

具体的には、SUNワークステーション5台をイーサネット⁽²⁾で結合した。(図1)

各SUNには、同一の並列処理系を搭載した。処理系は互いにデータをやりとりしながら、全体としてプロログを処理する様にした。

4. 並列処理系の構成

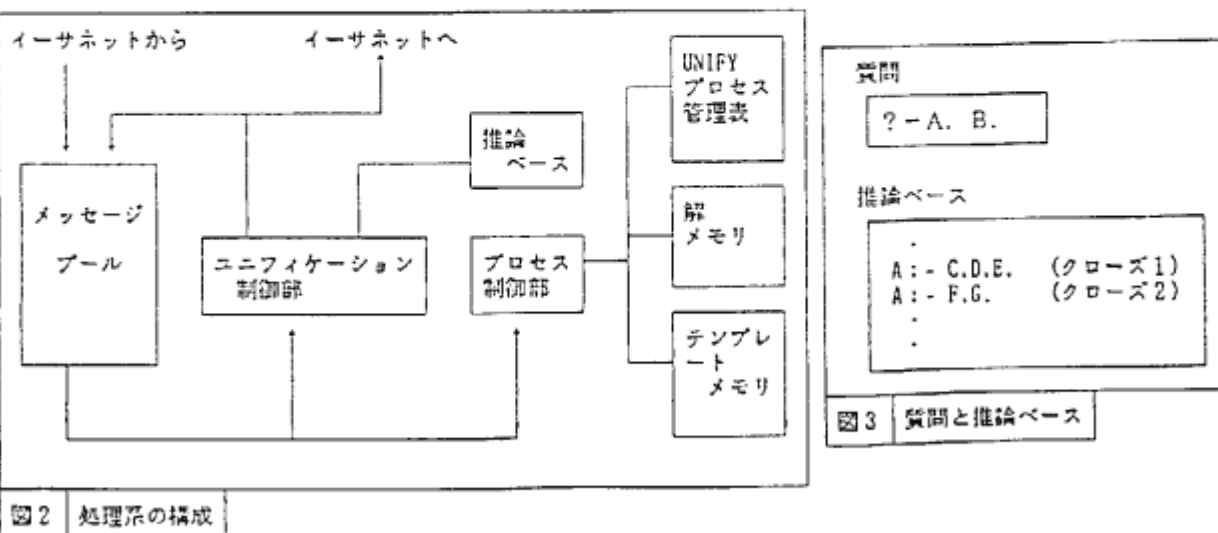
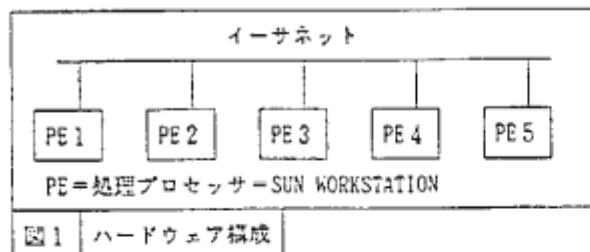
ここでは、各プロセッサ上に搭載した並列処理系について述べる。図2に示す処理系の構成を参照しつつ、動作の概略を説明する。

(1)処理系の基本動作

①質問メッセージに対する処理

メッセージプールに質問メッセージ?-A.B. が在るとする。

ユニフィケーション制御部は、?-A.B. を取り出し、その最も左のリテラル、すなわち Aについて推論ベースを使用してユニフィケーションを行う。



・一つ以上のクローズとユニファイに成功した場合:

質問メッセージ (?-A,B) を制御する為の情報
を UNIFY プロセスとして UNIFY プロセス管理表に登録
し、またそのテンプレート (A,B) をテンプレートメモ
リに格納する。

例えば、クローズ1でユニファイが成功したとす
ると、そのボディを質問メッセージ?-C,D,E.として、
メッセージプールへ、或いはイーサネットを介して他
のプロセッサへ転送する。

・一つのクローズも成功しなかった場合:

親の UNIFY プロセスに fail メッセージを転送する。

②解メッセージに対する処理

プロセス制御部はその解メッセージから対応する UN
IFY プロセスを見だし、その UNIFY プロセスの制御
情報から得られるテンプレート (A,B) を取り出す。

そして、その解をテンプレートに反映させて作成し
た B に関して、ユニフィケーション制御部は①と同様
の処理を行う。

③メッセージの選択

メッセージプールからのメッセージの取り出し順は、
種類によらず、受け付け順に取り出すことにした。

(2)処理の分配

他のプロセッサに依頼する仕事の単位を大きくする
事を旨とし、UNIFY プロセスを次の戦略で分配した。

- ・質問が組み込み述語かファクトなら、自プロセッサ
のメッセージプールに入れる。
- ・質問がボディのあるルールなら、自分のプールを見
て、ルールがあったら、別のプロセッサに質問を送
る。自分のプールにルールが無かったら、そのプー
ルに質問を入れる。

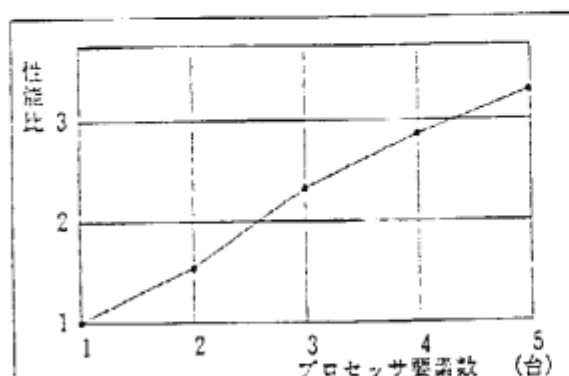


図4 1台に対する性能比

6. 実行結果

試作したシステムで、プロセッサ数数の効果を見て
みた。例として8クエリ間のプログラムを実行した時の
結果を図4に示す。図中の値は、1台のプロセッサの
性能に対する複数プロセッサを使用した場合の性能比
である。

グラフより以下の事が言える。

- ・3台のプロセッサで実行したら、1台のときの2倍
以上 (3×0.77)、5台では3倍以上 (5×0.65)
の性能向上が見られる。

今回の試作に使った処理系の性能と通信路の性能は
どちらも充分なものではなかった。処理系のユニフィ
ケーション当りの時間と、1回の通信に要する時間の
比は3対1程度と近接していた。しかし、5台の場合
の実行結果では、ユニフィケーションが約8000回、通
信が約100回であり、全処理時間に占める総通信時間
は無視できる値であった。

しかし、グラフを見ると、台数効果が十分には出て
いない。その原因は、処理の分配方式が適切でなく、
プロセッサの稼働率に偏りがある為と考えている。

6. おわりに

試作したPROLOGの並列処理システムについて述べた。
節単位処理モデルに基づく並列処理システムを複数
のパーソナルコンピュータを使用して試作した。その
試作実験により、並列処理効果を確認した。

今回の実験においては、処理系が初期版であり、通
信路もイーサネットという様に、各要素の性能はあま
り高くない条件での結果であった。処理系に改良が加
えられ、そして更に高速の通信路が使用される場合の
効果については今後の調査が必要である。

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、I C O T の委託で行ったものである。

ここで、日頃から様々な指導を賜っている、釜島部長
と堀馬室長に感謝の言葉を捧げたい。

参考文献

- (1)佐藤ほか「節単位処理モデルの提案」情報処理第28
全国大会5B-8.P.1141

並列環境における Concurrent Prolog 実現法

尾内理紀夫 麻生盛敏
(財) 新世代コンピュータ技術開発機構

1. はじめに

我々は、リダクション方式並列推論マシン[ICOT 84]の研究開発を行う過程において、ソフトウェアシミュレータを開発し、その上で、Prologプログラムを走行させ、各種データを収集している。このシミュレータ上では、Prologプログラムだけでなく、Concurrent Prolog (CP) [Shap 83]プログラムをも実行できる。これは、いわば並列環境でのCPの一種の処理系を実現している訳で、これについて述べる。

2. 並列環境

並列環境と、本稿において前提とする事項について述べる。

①プロセス：CPのClauseは、次のような形をしている。

head :- guard | body.

guard, body はリテラル列である。プロセスは、(結合情報+guard リテラル列+bodyリテラル列) からなる。また、Parallel And Operator で結合されたりテラル列がreduction可能となった時は、各々のリテラルを独立したプロセスとして生成させる。(reduction可能かどうかは、Parallel And Operator, Sequential And Operator, Commit Operator により指示される。)そして、プロセスを単位として処理ユニット(並列環境から処理ユニットは複数ある)に割りつける。

本並列環境では、reduction 可能なリテラルのみをeager copyしてreductionする。例えば、bodyリテラル列が、

b1&b2&b3 ('&'はSequential And Operator)

のとき、b1のみをcopyし、reductionする。その結果、子プロセスが生成されれば、子プロセスから親プロセス b1&b2&b3 へ、backポインタを張る。これは、子プロセスから親プロセスへ結合情報を返す時に使用する。以後、reduction 可能なguard リテラル、bodyリテラルをreducible subgoal と呼ぶことにする。

プロセスの持つ各種tag 内にC-tag(Commit tag)がある。プロセス生成時、C-tag はOFF であり、最初にguard が成功した子プロセスがこれをONにしてcommitする。guard が成功した子プロセスはcommitする前に親プロセスのC-tag をチェックし既にONならば、消滅する。(C-tagがONになった直後に他の兄弟プロセスを殺しにはいかない。)

②チャンネル：

規定1：clause内のParallel And Operator で区切られたリテラル間の共有変数はチャンネルである。

(例) go:-p1(X), p2(X), p3(X).

(','はParallel And Operator)

規定2：clause内のリテラル間の共有変数のうちの少なくとも一つにread only 変数がある時、これら共有変数はチャンネルである。

(例) go:-p1(X), c1(X?), c2(X?).

X はチャンネルであり、X?を特にreadチャンネルと呼ぶ。

このように、チャンネルとして使用される変数(チャンネルと呼ぶ)と、それ以外の変数(変数と呼ぶ)とを区別する。また、チャンネルの性質は実行時に動的にinheritする。たとえば、clause側引数内変数は、reducible subgoal 側のチャンネルとunifyされるとチャンネルとなる。(3章の例を参照)

③Message Board：チャンネル用のメモリとしてMessage Board (HB)を設ける[Onai 83]。producerプロセスとconsumerプロセスとは、HBを介してメッセージをやりとりする。HB内の一つのチャンネル用セルの論理構成は次のとおり

チャンネルの値	suspend プロセス リスト
---------	------------------

suspend プロセスリストは、このチャンネルが原因でsuspendしたプロセスが登録され、値が書き込まれた時、その値がsuspend プロセスに送られる。

④OR並列実行：reducible subgoal とclauseとのunify はOR並列に実行される。

⑤AND並列実行：Parallel And Operatorで区切られたリテラルはそれぞれ独立したプロセスとしてAnd 並列に実行する。

3. Concurrent Prolog 実現法

チャンネルは、論理的に二語の“チャンネル情報”でデータ表現される。第一語は、HB内該当チャンネルあるいは親プロセスのチャンネル情報へのポインタであり、第二語は、local データ領域である。guard が成功し、commitした時に、このlocal データをHB内該当チャンネルセルあるいは親プロセスのチャンネル情報内local データ領域に書き込む。以下、例を用いて説明する。

(例 3.1) p(*) reducible subgoal 側
チャンネル情報
*Xmb *Xmbは、HB内セルへのポインタ
local データ領域 (以後略して p(**Xmb))

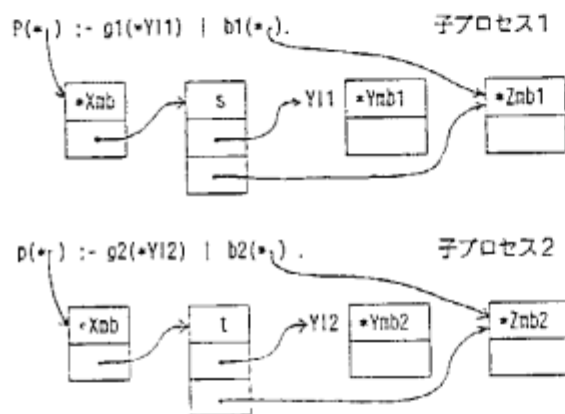
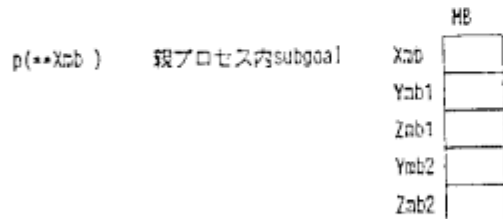
p(s(Y, Z)) :- p1(Y) | b1(Z).

OR関係にあるclause側

p(t(Y, Z)) :- p2(Y) | b2(Z).

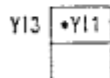
(まだY, Z はチャンネルではない)

unify 時にチャンネルの性質はinherit するので、各Y, Z は、チャンネルとなり、OR clause ごとにHB内にセルが確保される。一方、g1, g2 のY に関しては、並行関係にあるプロセスが存在しないので、このg1, g2 のY は、子プロセスレベルではチャンネルとしては使用されない。よって、このg1, g2 のY のために、HB内に新たにセルは確保せず、headチャンネルとチャンネル情報を共有する。(本並列環境では、reduction 時には引数はeager copyされる) reduction 結果は、次のとおり。

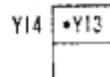


これら、同一の親を持つプロセスでも、同一処理ユニットに存在するとはかぎらないから、g1, g2がそれぞれ成功すると、これらのプロセスは、親プロセスのC-tag チェックを行う。そしてC-tag を先にCNにしたプロセスのみがcommitし、local データをHB内のXnb に書き込む。以下のようにguard g1がネストする場合は、f1, f2 には、やはり並行関係にあるプロセスはないので、チャンネル情報をheadチャンネルと共有する。各local データの内容は、プロセスが成功した時、結合情報として親プロセスのlocal データ領域に返されていき、g1(•Y11)が成功し、commitする時に、local データがHB内のXnb, Ynbに書き込まれる。ネストの先で、プロセスがsuspend した時は、チャンネル情報を逆にたどり、HB内チャンネルセルにアクセスし、値が既に書き込まれていれば、取り込み、まだならばsuspend プロセスリストに登録する。

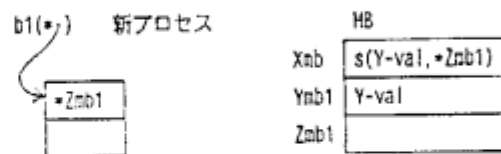
g1(•Y13) :- f1(•Y13) | b3.



f1(•Y14) :- f2(•Y14) | b4.



本例で、例えば、子プロセス1のguard が先に成功し、その時Y11 のlocal データ値がY-val とすると、commitした後はHB への書き込みが起り、次のようになる。

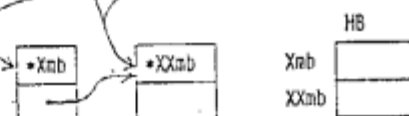


(例 3.2) guard 内のあるチャンネルに関し並行プロセスが存在し、各並行プロセス内チャンネルのうちに、少なくとも一つreadチャンネルでないチャンネルがある時は、そのチャンネルのために新たにHB内にセルを確保する。

p(•Xnb) reducible subgoal 側
p(X):- g1(X),g2(X?) | b(X). clause側

reduction の結果、

p(•) :- g1(•), g2(•?) | b(•).



となる。g1とg2は独立したプロセスとなり、HB内のセルXXnbを通じてメッセージをやりとりする。そして、g1, g2 が共に成功し、commitした時にXXnbの値をXnb に書き込む。

なお複数producerプロセスが存在する時、各プロセスは、あるチャンネルへの書き込みをする。この時、該チャンネルのHB内チャンネルセルは、一つであるから、同一値の書き込みは許されるが、異なる値の書き込みは、failになる。

A. おわりに

チャンネル用メモリ(Message Board) およびreducible subgoal のeager copyによる、並列環境でのCPの実現法について述べた。これはリダクション方式並列推論マシンのソフトウェアシミュレータに組込まれており、シミュレーション結果については、機会を改めて報告する。最後に日頃御指導いただく村上田男第一研究室長、御討論いただいたICOT研究員の皆様に感謝する。

【参考文献】

- [Shap 83] E.Y.Shapiro : A Subset of Concurrent Prolog and Its Interpreter, ICOT Technical Report TR-003, 1983
- [ICOT 84] 電子計算機基礎技術開発成果報告書 推論サブシステム編 (財)新世代コンピュータ技術開発機構, 1984
- [Onai 83] 尾内, 麻生: 並列推論マシンにおけるGuard と入力annotationの制御機構, 第58回情報処理全国大会

CAD向けPrologインタープリタ CADLOG

光本 圭子 森 啓 藤田 友之 後藤 敏
(日本電気(株) C&Cシステム研究所)

1. はじめに

最近、人工知能、あるいは、知識工学の研究が活発になり、そのプログラム言語として、Prologが注目され、実用化に向けて、改良、拡張の検討が行なわれている[1,2]。ここで、大規模問題を取り扱うシステムを、Prologですべて表現することを考えると、現在のハードウェア環境や、Prologの処理系では、実行速度、メモリ使用量の点から見て問題がある。著者らは、FORTRAN 言語との結合機能をPrologに付加することによって、実行時間の短縮、蓄積されたソフトウェアの有効な活用を可能にした。

本稿では、Prologの実用化の1つとして、FORTRAN 言語とのリンク機能を持つPrologインタープリタCADLOG について報告する。

2. CADLOGの特長

CADLOGの構文規則は、DEC-10版 Prolog[3]に類似している。文字は、大文字、数字、記号の3種であり、変数は '\$' を付加して記述する。

2.1 FORTRAN 言語とのリンク機能

手続き的処理が多い問題に対する記述しやすさや、プログラムの実行効率の面から考えると、Prologに手続き型言語(FORTRAN)との結合機能を付加することは、有効な手段である。FORTRAN で書かれたプログラムとのリンクは、外部ファンクション述語と呼ばれる述語によって対応づけられる。Prologの中で、この述語は頭文字に '\$' を付加して他の述語と区別している。FORTRAN 言語では、Prologのような引数の評価ではなく、必ず入出力を規定しなければならない。したがって、CADLOGでは、外部ファンクション述語の引数の入出力の整合がとれなければ、FAILとして、バックトラックが生じる。

例えば、FUNCという節を下記のように表現する。

FUNC(\$X,\$Y):-\$PROC1(\$X,\$Y). サブルーチンPROC1(X,Y)に対応する外部ファンクション述語
FUNC(\$X,\$Y):-\$PROC2(\$X,\$Y). サブルーチンPROC2(X,Y)に対応する外部ファンクション述語
FUNC(\$X,\$Y):-\$PROC3(\$X,\$Y). サブルーチンPROC3(X,Y)に対応する外部ファンクション述語
(注) $\underline{X}$ は入力変数、 $\overline{Y}$ は出力変数を表す。

このFUNC述語を用いて、SUB_GOAL1 という節を次のように定義する。

SUB_GOAL1(\$X,\$Y,\$Z):-FUNC(\$X,\$Y),FUNC(\$Z,\$Y).

この節において、第1引数が定数、第2、3引数が変数とユニファイされた例を示す(図1)。

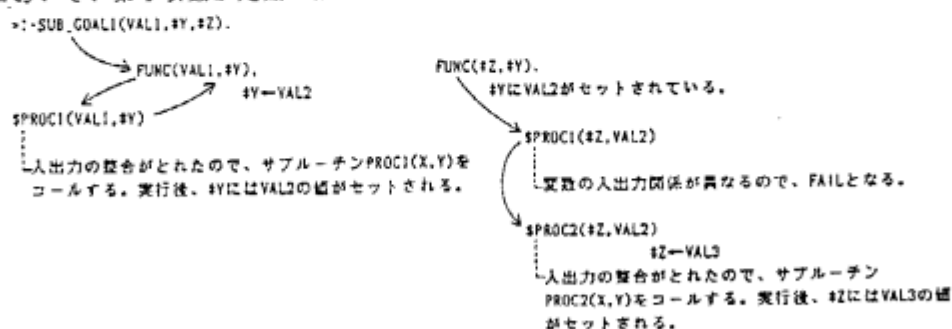


図1 外部ファンクション述語に対応するFORTRAN サブルーチンがコールされる様子

2・2 FORTRAN 言語によるシステムのデータ構造内のデータの整合性

FORTRAN 言語とのリンク機能を持っているため、この言語で構築されているシステム内のデータ構造を使用することができる。これによって、現在Prologでは実現不可能な大規模データ処理が効率よく処理でき、実行時間も短縮される。FORTRAN で構築されているシステム内のデータ構造を用いる場合、このデータ構造を作っている変数と、Prologプログラムの中に表われる変数とで、システム全体の状態を表現することになる。そこで、Prologのバックトラックが生じた時に、Prologプログラムの変数に対してのみユニファイされる前の状態に戻す処理を行なうと、Prologの変数が表現する状態と、データ構造内のデータが表現する状態とに食い違いを生ずる。したがって、プログラマは、どの外部ファンクション述語がデータ構造内のデータを更新するか常に意識し、バックトラックが生じると、どのような処理を行なわなければならないか考えて、プログラムを組む必要がある。これは、プログラマに負担をかけるばかりでなく、プログラムを読みにくくする。CADLOGでは、データを更新する外部ファンクション述語に対応するFORTRAN サブルーチンに、常に更新前のデータを保持する処理を組み込んでおき、バックトラックが生じてこの述語まで戻った時に、保持しているデータを用いて、データの改復処理を行なう機能を持っている。例を図2に示す。このように、FORTRAN で構築されたシステム内のデータ構造を用いる場合でも、CADLOGはデータの改復処理用サブルーチンを組み込むことにより、常にデータ構造内のデータの整合性を保つことができる。

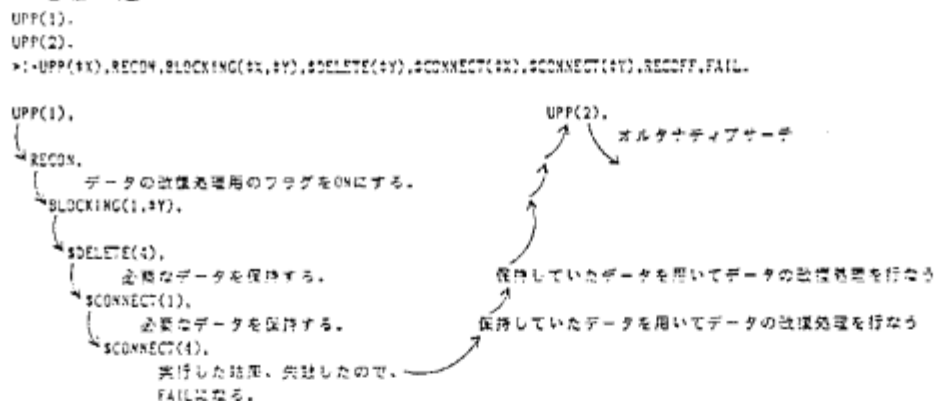


図2 データの改復処理機能の例

2・3 外部ファンクション述語の登録、変更、削除、表示

外部ファンクション述語を新たに登録したり、対応するFORTRAN サブルーチンに変更が生じたり、不必要な外部ファンクション述語を削除するために、ユーティリティが用意されている。このユーティリティを用いて、ユーザは、外部ファンクション述語の登録、変更、削除が容易にできる。

3. おわりに

FORTRAN 言語とリンク機能を持つPrologインタープリタ CADLOG について報告した。この機能により、大規模問題を取り扱うシステムでも、蓄積されているソフトウェアを有効に活用しながら、人工知能の手法を取り入れたシステムの構築を可能にする。現在、CADLOGは、NEC MS上でインプリメントされ、LSI配線設計エキスパートシステム[4]を構築する言語として、使用している。

参考文献

- [1] 中島, "Prolog/KR の概要", 情処, 記号処理, 18-5, 1982.
- [2] 梅村, 中嶋, 小長谷, 幅田, 島津, 横田, "文字処理を導入したProlog:Shapeupについて", Proc. of the Logic Programming Conference, 1983.
- [3] W.F.Clocksinn and C.S.Mellish, "Programming in Prolog", Springer-Verlag, 1981.
- [4] 森, 光本, 藤田, 後藤, "配線エキスパートシステム", 情処第28回全国大会, pp 1073-1074, 1984.

Study of a Newspaper's Treatment of Proper Names

Satoru Ishii

Institute for New Generation Computer Technology

1. Introduction

People's names are important keywords in sentences, and care must be taken that they be correct. However, errors in names are difficult to detect. For this task, computerized name-checking system may be of great assistance in avoiding errors, and may reduce the proofreader's workload.[1]

With a view toward designing a computerized directory-consultation system, this study used one day's output of a typical Japanese newspaper to determine the following characteristics of the use of personal proper names:

- (1) Style in which names appear
- (2) Level of sentence understanding required
- (3) Types of directories required.[2]

2. Sample Data

The morning and evening editions of the July 19, 1983 Asahi Shimbun constituted the study's sample data. All sentences, including headlines and tables, were analyzed. Advertisements, TV and radio listings and serials were ignored. Assumed names such as pen names and screen names were treated as conventional names.

The data contained a high volume of sports reporting, especially, stories about amateur and professional baseball; this bias has not been overlooked.

3. Style of Name Usage

We established four categories of style,[12][13] and tabulated names falling into each category. The categories and the number of appearances are as follows.

Names are used:

with an organization and a post	79 times
with a post	227 times
with a title or an honorific	272 times
with none of the above	1293 times

The first category mainly included names of statesmen. The second category mainly included names of statesmen and sports figures. The third category included names of various persons falling into no particular grouping. The last mainly included names of sports figures and entertainers.

There were 12 names that could not be assigned to any of these four categories; these were ignored.

4. Contextual Understanding Required

In a computerized consultation system, in the case of government officials for instance, three fields must be generated through contextual understanding: the individual's name, the organization he belongs to and the post he holds. We divided the contextual understanding required into six levels and counted names appearing in each.

The levels, from easy to difficult, have been tentatively established as follows. Figures in parentheses indicate the percentage of appearances in each level.

Understanding of BUNSETSU required(10%): In the BUNSETSU in which the name appears, the organization and post also appear. Therefore, understanding of the BUNSETSU is required. BUNSETSU is a concept unique to the Japanese language, and has no English equivalent.-Ed.

Understanding of text format required(43%): The text is based on a well-defined format, the understanding of which is required.[9] A typical example is a table of baseball scores.

Understanding of relationship to another appearance required(8%): The name appears in the same text in a more detailed style, and understanding the relationship between these appearances is required.

Understanding of context required(14%): The organization and the post appear in the text, therefore, the context must be understood.

Understanding of text domain required(24%): The text domain, for example, sports or science, is necessary.

Impossible(1%): The name forms part of another word. In such cases, it is impossible to distinguish the name.

Understanding of the BUNSETSU is sufficient for only 10%of the appearances. For 90%of the appearances, more in-depth understanding is required.

5. Effectiveness of Directories

We divided names into nine classes, selected appropriate directories for each class and determined the effectiveness of each. The results are shown in Table 1.

Names of government officials and foreign statesmen were effectively located. Names of professors and directors could be located; however, since such names were few, relative to the size of the corresponding directories, confirmation of names in these classes would be inefficient in a computerized system. Names of creative artists were difficult to locate. Names of sports figures were effectively located, but this depends on the sports involved. To locate the names of entertainers, an extensive directory was required.

Of the entire sample, about 75%of the appearances could be located using a number of directories whose entries listed approximately 200,000 persons.

6. Conclusion

Using one day's output of a large Japanese newspaper, we studied the usage patterns of personal proper names. This work represents the first stage in the design of a computerized directory-consultation system.

(1) We divided the names into four categories and found that names unaccompanied by titles appeared most frequently.

(2) For generating database queries, the understanding of BUNSETSU was sufficient for only 10%of the appearances.

(3) Using adequate directories with approximately 200,000 entries, about 75%of the appearances could be located.

Table 1. Effectiveness of Directories

A: Class	B: Appear- ances (%)	C: Directories Used	D: Entries (in 1000s)	E: Located	F: E/B (%)	G: E/D (%)	H: F*G
Government officials	248(13)	[3]	5	219	88	4	352
Foreign statesmen	61(3)	[3]	1	36	59	4	236
Professors	65(4)	[4]	90*	65	100	0.07	7
Directors	37(2)	[5]	30**	23	62	0.08	5
Creative artists	123(7)	[3]	6	36	28	0.6	17
Sports figures	1035(55)	[7][8] [10][11] [14]***	50	976	94	2	188
Entertainers	71(4)	[6]	9	57	80	0.6	48
Newspapermen	27(1)	---	--	0	0	---	---
Others	199(11)	---	--	0	0	---	---
Total	1871(100)		191	1412	75	0.7	53

* This figure represents only professors and assistant professors. Names of other university instructors were assumed not to have been listed.

** This figure represents only directors. Names of other company officials were assumed not to have been listed.

*** Only the number of KOUKOU YAKYUU (high school baseball) teams was found. Names of players were assumed to have been listed.

References

- [1] Ishii, S.: "Study of Proofreading Techniques Used at a Japanese Newspaper", Proc. of the 28th National Conf. (II) 2H-7, pp.1205 ~1206, IPSJ (1984)
- [2] 石井 睦: "新聞における人名の使い方の調査", TM-0062, 新世代コンピュータ技術開発機構(1984)
- [3] 堀山定雄: "朝日年鑑1983年版", 朝日新聞社(1983)
- [4] 大学職員録刊行会編: "1983全国大学職員録", 廣潤社(1982)
- [5] ダイヤモンド社編: "ダイヤモンド会社職員録(全上場会社版)1983年", ダイヤモンド社(1982)
- [6] 著作権資料協会編: "出演者名簿1983年版", 著作権資料協会(1982)
- [7] "棋士名鑑", 昭和57年版将棋年鑑, pp.335~381, 日本将棋連盟(1982)
- [8] "棋士名鑑", 1982年版囲碁年鑑, pp.389~441, 日本棋院(1982)
- [9] "記事のフォーム", 記者ハンドブック第4版第3刷, pp.327~338, 共同通信社(1982)
- [10] "12球団新陣容", 朝日新聞1983年3月28日号, 朝刊, pp.28~29, 朝日新聞社(1983)
- [11] "昭和五十八年度九月場所東西幕下十両力士星取表"及び"昭和五十八年度九月場所東西幕下以下力士星取表", 相撲1983年10月号, pp.182~185, ベースボール・マガジン社(1983)
- [12] "人名、年齢の書き方", 記者ハンドブック第4版第3刷, pp.359~363, 共同通信社(1982)
- [13] "人名等の書き方", 新版記事スタイルブッカー新用字用語集-新版2刷, pp.443~452, 時事通信社(1983)
- [14] 朝日新聞1983年7月17日号, 朝刊, p.16, 朝日新聞社(1983)

知的プログラミング環境の実現へ： 第5世代プロジェクトにおけるアプローチ

横井 俊夫

(財)新世代コンピュータ技術開発機構・研究所

1. (知的) プログラム言語

プログラム言語のプログラミング手法やプログラミング環境への関わり方は、その時代、時代で様々な形をとってきた。コンパイラの制作が自動プログラミングのすべての話題であった時代から、手法や環境は、言語と独立して議論できる、そうすべきことが強調されさえたソフトウェア工学の時代へと移る。最近はもう少し言語を重視しようという風潮がみられるようになったが、そして、第5世代コンピュータは、プログラム言語を再び主役に引き立てる。

なぜ、主役の座を去るに至ったかの過去をたどる。

(1) プログラム言語の研究に大きな展開が無くなった。様々な言語の模倣はCOBOL, FORTRAN, PL/I等の商用汎用言語に集約された。新しい要求事項も、それらの言語の部分的な手直しにとどまり、根底からおびやかすような事態は起らなかった。また、その傾向は、互換性維持の立場からの強い要求によって、一層増長されることになる。

(2) ソフトウェアの基本的枠組を規定するのは、プログラム言語だけでなく、オペレーティング・システム、プログラミング・システム、データ・ベース・システム等を大きく前提として考慮せねばならなくなった。

第5世代における復活は、この2つの事態に大きな変化が生じた、生じつつあることに基づく。

(1) 知的機能を実現するための土台としての記号処理機能、モジュール化の新しい仕組であるオブジェクト指向機能、これらは、従来の言語機能を大きく越え、新しい枠組を求める。

(2) 新しい枠組に立つ言語は、従来の言語の機能をその一部に含み、新しい機能は高度なオペレーティング・システム、プログラミング・システムの構成を容易にし、データ・ベース機能が本来の仕組として組込まれており、ソフトウェアを総合的に規定する土台となる。

この新しい言語を方向づける基本的プログラミング・スタイルとして論理型言語を採用する。

2. (知的) ワーク・ステーション

新しい言語に則った高機能なワーク・ステーションは、プロジェクト全体の共通の研究・開発ツールであるが、それは、また、知的プログラミング機能を事実に関して明らかにする場、またシステムを具体的に検証する場としても必須のものである。逐次型推論マシンの名のもとに、高機能パーソナル・コンピュータ：PSI(マシン・サイクル200ns, 主記憶16M語)を開発し、現在、ソフトウェア・システム：SINPOSを作成中である。オブジェクト指向機能を論理型言語上に融合したESPを使用し、新しいモジュール分割によるシステム開発を実験している。モジュール・ライブラリの知識ベース化、知的Help機能等、具体的事例に即しての知的プログラミング・システムの実現を計画して居り、また、ハード、ソフト、言語の改良、拡張を勇力的に続け、プロジェクト全期をかけ大きな土台に成長させる予定である。

3. 知的プログラミング・システム

まだ、小規模な実験システムの試作や理論的検討の域を出ていない仕様記述・検証技術への本格的な取組み、実用化と称しうるレベルのシステムの開発を目指して展開される。これを基調とし、ソフトウェアの日本語化、さらに自然言語化とソフトウェア開発・制作支援の統合化が絡む。仕様記述、検証技術に向っては、様々な側面からの様々な試みが必要であるが、大きく仕様記述の側からと検証・自動証明システムの側からアプローチする。対応するメーカーの研究グループの努力に加え、仕様記述に関しては、東工大(根本教授)TELL/NSLプロジェクトとの共同研究、自動証明に関しては、基礎理論ワーキンググループ(広瀬教授主査)における研究活動が、大きな支えとなっている。