

72-014

逐次型パーソナル推論マシンの
設計思想とそのアーキテクチャ

西川 宏 横田 実 山本 明

瀬 和夫 内田俊一

逐次型パーソナル推論マシン^{*)}の

設計思想とそのアーキテクチャ

西川宏, 横田実, 山本明, 瀧和男, 内田俊一 (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では人間のよう
に考えることのできる、推論機構を
有する計算機システムの実現を目指
して研究開発が進められている。

そのためには論理型プログラミング
が研究開発の環境を早期に確立する
ことが必要であり、本プロジェクト
の一環としてICOTが中心になり
パーソナル逐次型推論マシン(Personal Sequential Inference Machine: Ψ)
の研究開発を行っている。

Ψ の研究開発は本プロジェクトの
研究開発を進めるのに必要なソフト
ウェア開発用ツールとなること、ま
た論理型言語向きの新しいマシニア
ーキテクチャ研究開発のテストベン
チとなることであり、実行速度20~
30k LIOPS^{*}, 最大実装記憶容量16M語
の高性能パーソナル推論マシンとし
て開発される。更に新しい試みとし
て Ψ の上に実現されるオペレーティ
ングシステムを論理型言語によつて
記述することを目指して、 Ψ の機
械語である Ψ 言語^{*}をProlog-like
な言語仕様とした。したがって
 Ψ は述語論理型言語を機械語として
直接実行できるマシンであると共に、
オペレーティングシステムのための
ハードウェアサポート機能を備えて
いることが最大の特徴である。

本論文では Ψ の研究開発に要求さ
れている様々な目標を述べると共に、
 Ψ のシステム概要、アーキテクチャ
を決定する上での設計方針、並びに
アーキテクチャ上の特徴について述
べる。

2. 背景と開発目標

本プロジェクトの基盤となる論理型言語
の例としては、一階述語論理に基づくプロ
ログラミング言語PROLOGが最も良く知られて
おり、中でもDECIO PROLOGコンパイラ版[1]は
最も実行効率が良いとされている。

しかしながら既存商用マシン上で実現さ
れている処理系は従来型言語を実行する場
合に比べて大きなメモリ空間を必要とする
うえ、実行効率もあまり高くない。従つて
論理型言語の持つ高い記述性、非決定性とい
う特徴を活かすためには専用マシンア
ーキテクチャによるハードウェアサポート
が不可欠である。

更にコンピュータシステムを一つの概念
で統一できればシステムを単純化でき、理
解の容易さ、開発の容易さにつながるため
従来のオペレーティングシステム全体も論
理型言語で記述できるであろう。

ところが、論理型言語を用いる場合の最
大の利点の一つである非決定性は、オペレ
ーティングシステムのような現実的処理が
主体のシステム記述に用いるのは得策では
ないとされている。このために、低レベル
のシステム記述を論理型言語の枠組の中
で行なえるような機能拡張が必要である。

これらの背景から本プロジェクトでは逐
次型推論マシンの研究開発を進めているが、
その開発目標はオーに新しい論理型プロ
グラミング言語の研究、試作/評価のため、
及び論理型言語による種々の応用システム
(例えば自然言語処理、専門家システム等)
の研究、開発のためのソフトウェア開発用
ツールとなることである。オーには論理型
言語の実行に最適な新しいマシニアーキ
テクチャを開発するための基盤となること
がある。

開発用ツールとしての目標を達成するた
めには以下のような要求事項を満たすこと

^{*}) Logical Inference Per Second

が心配である。

- (1) 大規模な応用プログラムを実行可能とするだけの十分な実行速度、メモリ空間を備えていること。
- (2) ツールとしての高い操作性、作業性と変更するための良好なマシンインターフェースを備えていること。
- (3) ソフトウェア開発作業の分散化、あるいは既存計算機システムとの間での機能分担、ソフトウェアの流通を図るためのネットワークとの結合機能を備えていること。
- (4) の性能については最も著目している DEC-10 PROLOG を基準とするのが妥当であろう。これまでの PROLOG プログラム開発経験より実行速度、メモリ空間ともに DEC-10 PROLOG の 1 桁以上の性能が必要と思われる。特にメモリ空間の拡張は不可欠である。何故ならば実行速度が遅い場合にはその分だけ長時間計算機を動かすことによりとにかく解を得ることが可能であるのに対し、メモリ空間が狭い場合にはそれを使い切ってしまうとそれ以上実行不可能となることである。この意味で仮想記憶システムの導入が望ましい。

(4) の良好なマシンインターフェースを実現するためにはビットマップディスプレイ装置、ポインティングデバイス(マウス等)を用いた多生ラインドウシステムを備えることや、日本語入/出力装置をサポートすることが必要である。またシステムの利用形態としてはパーソナルマシンとして個人で専断して利用できる方が高い操作性が期待できる。そのためには(3)のネットワークインターフェースをサポートすることにより共有データベースへのアクセス等のシステム拡張が不可欠となる。

オペレーティングシステムの論理型言語による記述をはじめとして、

本プロジェクトで開発予定の各種ソフトウェアは実際に逐次型推論マシンを用いて開発する必要があるので、それを初め早期に完成させることが要求されている。

論理型言語向きマシニアーキテクチャの開発開発のオメガとしては次のような項目が要求されるであろう。

- (1) 各種評価データの収集機能を備えること。
- (2) 言語仕様の変更に対処できるだけの、また新しい解釈/実行方式の実験、評価ができるだけの柔軟性を備えること。
- (3) 次の研究ステップにつながる新しい試みを積極的に導入すること。
- (4) の評価項目としては論理型プログラムのプログラム特性、実行特性の抽出から、ハードウェアレベルでの設計パラメータの有効性の評価まで含まれる。後者ではキャッシュメモリのヒット率をはじめとするメモリアクセス特性が重要な課題となる。
- (5) の柔軟性についてはマイクログラム方式や機械語の設計方法により達成することができる。論理型言語として最も著目している PROLOG に関しても今後の機能拡張が研究・提案されている。また並列実行へ向かっての並列 PROLOG (Concurrent Prolog) が最近注目されており言語仕様のレベルでの試行錯誤が今後とも必要である。従って機械語命令設計時には充分な柔軟性を考慮すること、あるいはマイクログラム制御の適用等が要求される。

(3) の実験的試みについては前述の開発ツールとしての性格と相反するものであり、この程度新しい試みを取り入れるのは難しいが(4)の評価機能を有効に利用して積極的な実験・研究をすべきであろう。

本プロジェクトでは、まもなくのソフトウェア開発に用いるためのパーソナルツールとしての要件を満たすパーソナル逐次型推論マシン(Personal Sequential Inference Machine: PSI)をまず開発し、その後機能拡張した高機能版を開発する計画である。

3. Ψ システムのアウトライン

Ψ はできる限り早期に、しかも実際にツールとして使えるマシンとして開発されることを要求されている。そのためには現状で最も実行効率が良いと考えられているスタックを用いた論理型言語の解釈/実行方式を採用すると共に極小型、シンプルに実現する方針とした。

Ψ の設計目標はVAXに代表される32ビットスーパーミニコンピュータ並みの性能で、Lispマシン等のビットマップディスプレイ装置を有するスーパーパーソナルコンピュータ並みの使い易さを兼ね備えたパーソナルマシンとなることであるが、性能については今後の開発ソフトウェアの大規模化を考慮してパーソナルマシンとしては若干高めに設定されている。

具体的な目標性能としては実行速度がDEC-10 PROLOG コンパイル版(DEC2060)と同程度、メモリ空間は1桁以上大きい4M語を標準実装として最大16M語と設定した。この理由は既に述べたように実行速度の高速化よりも使用可能なメモリ空間の拡大の方が論理型言語プログラム環境を充実させる上で急務であるとの判断による。

図1に Ψ のシステム構成を示す。 Ψ は基本的にローカルエリアネット

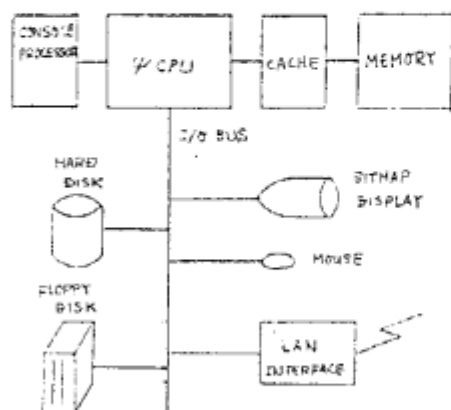


図1 Ψ システム構成図

ワークに結合されたワークステーションとして利用されることを前提にしている。従って Ψ 自身はあまり多くの入出力装置を備えてはいない。基本システム構成としては良質なマンマシンインターフェースを実現するためのビットマップディスプレイ装置、マウスと固定ヘッドディスク(300kバイト)装置を備えている。但し、 Ψ はサウンドアロンでの利用も可能であり必要に応じてプリンタ等を付加することも勿論可能である。

CPUはTTLにより実装することとし、十分な柔軟性を得るためにマイクロプログラム制御方式を採用し命令サイクルは200nsec以下を目標とした。アーキテクチャはPROLOG-like な言語仕様を持つ機械語、被言語KL8の解釈実行に最適化するように設計されており、最大の特徴は各語にタグを持たせ、その判定を高速に行なうハードウェア機構を持っていること(タグアーキテクチャ)、及びメモリアクセスの効率化を図っていることである。メモリアクセスはスタックを多用する論理型言語の解釈/実行方式では主要なオペレーションであり、その効率化が直接 Ψ の実行性能に影響する。従って Ψ ではキャッシュメモリの導入によりメモリアクセスの高速化を図った。仮想記憶については、現段階では論理型プログラムの動作特性がどのようなものかわかっておらず、また仮想記憶システムへのカーページコレクションの導入の技術的困難さも考え合わせ、 Ψ では仮想記憶の採用は見合わせることにした。そのかわり実用規模のプログラムを動作させその挙動を測定するために充分と思われる主記憶を用意することにした。その大きさはDEC 10 PROLOGの使用経験等を踏まえて、標準4M語、最大実装16M語であれば適切であると判断した。また Ψ ではシステムセキ上げ、保守のためコンソールプロセッサとして1ボードマイクローコンピュータを内蔵している。

入出力インターフェースについては極力標準化を図ることによってユーザの要求に合わせた種々のデバイスを自由に付加できることを意図し、IEEE標準仕様であるIEEE-796バス(Multi-Bus)を採用した。

入出力制御は中のCPUが直接制御する方式とした。この理由はシステム構成の単純化を図ることにあるが論理型言語によるマルチウィンドウシステムなどの研究実験を行うために核言語KLXによって低レベルの入出力制御を行うことも考慮しているためである。

Yは推論エンジンとしてのCPUのまわりに研究開発の努力を注いでおり、入出力バスは1本設けているが特に高速アクセスが要求されるデバイスとの接続や特殊デバイスの接続のために32ビットパラレルI/Oポートを設けている。

4. 設計上の選定点

核言語KLXは中の上での最も低いレベルの言語であり、従来のアセンブリ言語に相等する。コンパイラをはじめとして中のオペレーティングシステムの核部分は核言語KLXで記述される予定である。

しかしながらKLXは述語論理に基づく高い言語仕様を持っているため中の機械語命令もどのように設定するかはマシンアーキテクチャを設計する上での重要な問題である。

以下に核言語KLXを高速に実行するために備えるべきマシンアーキテクチャについて、その設計上考慮したいくつかの選定点と中での設計方針について述べる。

4.1. 核言語KLXの設定

Yを実際に活用するためにはハードウェアと共にオペレーティングシステムの開発が必要である。オペレーティングシステムは記述には決定性処理と履歴依存処理の機能が必要であるが、これらの機能は本来述語論理の体系には異質のものである。従って、システム記述用の目的のためには別の言語体系を導入する方が普通である。しかしこのように述語

論理ベースの言語でシステム記述用の別言語の2つの枠組を設けるとシステム全体の均質性が崩れ、ユーザーにとってシステムを統一的に理解することが困難になる。ことにアーキテクチャレベルでは、2種の異なる言語体系をサポートしなければならずそのハードウェア機能も複雑になることが予想される。

そこで中の機械語である核言語KLXはこれを本来異なる2種の体系を統合させ、簡素かつ統一的なハード、ソフトのインターフェースを作るよう設計された。即ち、PROLOGに代表される述語論理型言語の特徴の一つである宣言的にプログラムが書けるという良さと低レベルのシステム制御の記述のための組込述語を導入して機能拡張を図った。[2]

従って核言語KLXは2つの側面を持っている。一つはユーザー定義の述語であり、これはバックトラックを基本とするPROLOG-likeな非決定的な制御の対象となるものである。もう一つは組込述語であり、これについては必ずべき処理は決定しておりバックトラックの対象とはならないものである。しかしいずれの場合もデータの受け渡しはユニフィケーションを基本とする。

4.2 核言語KLXの実行について

核言語KLXは実行時に決定される動的要素を多く含んでいるにもかかわらずコンパイル方式を採用したのは次の理由に依る。

1) 組込述語は決定的(deterministic)に実行される。従ってその内部表現はできるだけコンパクトにすれば、処理効率の向上とメモリ使用の効率の点から望ましいこと。

2) TR O (tail Recursion Optimization) [3] はコンパイラのレベルで対処しなければマシンレベルで動的に処理することが難かしいこと。

これらの理由により核言語KLXはコンパイルされ、マシン内部形式(オブジェクトコード)に変換される。

このオブジェクトコードを実行するマシンの制御方式としては、ハードで直接実行

を行なうハードワイヤド方式とマイクロコードにより解釈/実行を行なうマイクロプログラム方式が考えられる。どちらの方式を採用するかはオブジェクトコードのレベルと密接に関連している。ここではマイクロプログラムによりオブジェクトコードを解釈/実行する方式を取ったがそれは次の理由からである。

- 1) コンフィケーション、クローズの実行順序制御等はメモリアクセスを中心であり、これらは分解され、低レベルのオブジェクトコードで処理されるよりも、一つのオブジェクトコードで処理される方が命令フェッチに要するメモリアクセスが少なくなる。このようにハードウェアの単純化、簡素化のために命令フェッチユニットを持たないマシンでは、オブジェクトコードのレベルは高い方が望ましいこと。
 - 2) このソフトウェア支援ツール、ハードウェア実験機としての性格を考えると、容易に機能の追加、修正のできるマイクロプログラム方式が望ましいこと。
 - 3) 前述の処理の処理は決定しておりこの中にはハードウェアで直接実行した方が明らかに望ましいものもあるが、プログラムの実行順序制御、実行環境を変更する等の命令は、コンパクトなオブジェクト表現を取る限りマイクロプログラムによる実行が望ましいこと。
- まとめると、このオブジェクトコードの表現の高さのためとこのマシンとしての性格から、この実現方式としてマイクロプログラム制御を採用した。

次に、マイクロインタプリタによる核言語KLの実行方式について述べる。実行処理方式を逐次型に限っ

て考えると、次の2つに絞ることが可能である。即ち、ヒープをベースにした実行メカニズムとスタックをベースにしたそれである。ヒープを基本とする処理方式では述語の呼び出しに伴い生成される変数領域は新しく作られるだけで解放つまり消滅はしない。一方スタックを基本とする処理では変数領域は、通常呼び出し元への復帰時に消滅する。

一般に通常の言語の処理方式としてスタックを基本とするものが多いのは、呼び出された手続きがバックトラックの対象とはならないので、実行終了時点の呼び出し元への復帰時にはその変数領域を解放してもよいからである。一方核言語KLの仕様ではバックトラック制御を行なわれ、クローズによってはバックトラックの対象となり得るものが存在する。従ってその場合にはそのクローズの変数領域は解放できない。

さらにコンフィケーションでは変数で構造体をバインドすることがある。この場合構造体への参照ポインタが変数には格納されるが、その構造体を含む変数領域を解放してしまうと変数の参照先がなくなってしまい、いわゆる dangling reference になってしまう。

このように核言語KLの仕様では変数領域が割り当てられるスタックは単純には解放することができない。それにもかかわらずこのスタックをベースとする実行メカニズムを採用した理由を次に示す。

- 1) ヒープベースによる実行メカニズムではメモリ空間の消費が激しく、ガーベジコレクションが実行される頻度が増すこと。
- 2) ヒープベースを採用することによれば広いメモリ空間を根拠することが望ましく、従って仮想記憶方式が有効であるがここでは採用を見合わせたこと。
- 3) DEC-10 PROLOG コンパイラ版は基本的にスタックをベースにした実行メカニズムを採用しているが高い処理効率を得られていること。

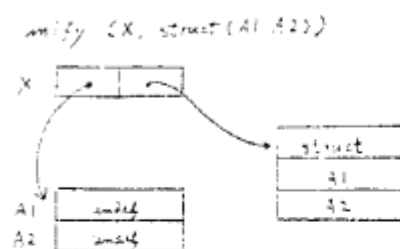
4.3 sharing / copy 方式について

ユニファイケーションによって変数に値が代入されるという操作は基本的にスタック上に生成された変数セル(一語)への値の代入として実現される。

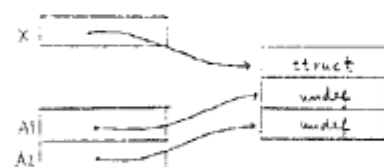
しかし代入すべき値が構造体データの場合は、1語対複数語の対応関係を作り出す必要がある。最も単純な方法は代入すべき構造体データの複製を作つてそれへのポインタをセットする copy 方式であるが、複製を作らずに相手のデータ構造を共有することにより実行効率を上げることが意図した structure sharing 方式[4]もある。いずれの方式も各長所、短所を持っているがここでは後者の structure sharing 方式を採用したがその理由を以下に述べる。

図2は変数 X と構造体 $\text{struct}(A1, A2)$ がユニファイされたときの sharing 方式と一般的な copy 方式を用いた場合に生成される環境をそれぞれ示したものである。

図2-(a)に示される sharing 方式では構造体のみな型 (skeleton) と実際の値とは分離して管理され、両者へ



(a) structure sharing



(b) structure copy

図2 変数 X と構造体 $\text{struct}(A1, A2)$ のユニファイケーション

のポインタを指した molecule によって構造体データが表現される。従つて sharing 方式では molecule の生成だけの手間で済む。一方図2-(b)に示される copy 方式では、構造体のフレームが新しく生成され、 $A1$, $A2$ に対応する変数セルからそれぞれそのフレーム内変数セルへのポインタが格納される必要がある。即ち copy 方式では構造体の要素数に比例した複製するたりの手間がかかることになる。

この状態でさらに変数 X と構造体 $\text{struct}(A1, A2)$ がユニファイされる場合を考えてみる。どちらの方法によつても構造体の名前 (functor) の一致を調べるのは同一であるが各要素へのアクセスの仕方に差が生じる。即ち struct の $A1$ 要素である変数 $A1$ とアトム $A1$ のユニファイケーションが実行される場合の $A1$ のセルアドレスの生成に要する手間が異なるのである。sharing 方式では構造体のみな型を先ずアクセスすることによりその $A1$ 要素 $A1$ の値領域である $A1$ 変数セル位置を知り、その変数領域を示すベースに加算することで求めるセルアドレス $A1$ が得られる。つまり1回のメモリアクセスと1回の加算が必要である。これに対して copy 方式では X の持つ構造体フレームへのポインタに $+1$ するだけでよく1回の加算で求めるセルアドレスが得られる。

このように構造体同志のユニファイケーションが行われる場合には copy 方式の方がオーバーヘッドは小さく、sharing 方式では構造体要素へのアクセス毎に余計なメモリ参照が必要となる。4のようにアドレス情報 が 32 ビットまで持つマシンでは1ワード内に molecule を格納することは難しい。 molecule の表現に2語要する場合には、図2-(a)を例にとると変数セル X には molecule へのポインタが格納されることになり、 molecule 読み出しのためにさらにメモリアクセスを要することになる。

ところが、今度は変数 $A1$ とアトム b が単にユニファイされる場合を考えてみる。sharing 方式では変数セル $A1$ へのアクセスは直接であるのに対し、copy 方式では変数セル $A1$ のポインタを参照するために向

格アクセスとなり sharing 方式に比べてオーバーヘッドは大きくなる。

以上のようにどちらの方式が優れているかはユーザーと構造体をどの様に扱うかという特性に強く依存しており一概には判定できない。4ではこれまでの PROLOG 使用経験より copy 方式での複製の手間の方が、sharing 方式での構造体要素へのアクセスに要する手間より大きいと考えて sharing 方式の採用に踏み切った。この問題は引き続きプログラム特性の評価もふまえて研究する必要がある。

4.4 ユニフィケーション高速化のサポート

枝言語 KL8 の解釈/実行の中心はユニフィケーションであり、4の実行効率を上げるためには、これを高速実行するためのハードウェアサポート機構が必要である。ユニフィケーションの処理の中心はメモリ参照とそのデータの属性判定である。従ってその高速化は結局

- 1) メモリ(特にスタック領域)アクセスの高速化、
- 2) データタイプの判定の高速化に帰結する。

スタック領域へのアクセスの高速化は単純な LIFO 形式のハードウェアスタックを用いるのが一般的である。しかしバックトラック処理を行なう枝言語では単独に領域をポップアップ(即ち解放)できない。この場合のスタック動作を忠実にサポートするハードウェアスタックの機構はかなり複雑なものになることが予想され、その有効性等については充分な検討が必要であると思われる。

従って4ではスタックアクセスの高速化のために1クロース分の変数フレーム(即ちスタックトップフレーム)をCPU内のフレームバッファ(高速な専用レジスタ)に割

り当てて、その中でユニフィケーションが実行されるようにした。これはTRCを実現するに当たっての引数レジスタ[3]としての使用を兼ねることもできかなりの高速化が期待できる。

一般のメモリアクセスの高速化にはキャッシュメモリを設けることにより対処することとした。キャッシュメモリの長所はスタックアクセスのみならず命令読み出しも含めたすべてのメモリアクセスを均等に高速化できる点にある。キャッシュメモリの方式としてはライトスルー方式とライトバック方式が代表的であるが、若干性能がよいとされているライトバック方式を採用した。

2) のデータタイプの高速化のためにはすべてのデータにタグを付与するタグアーキテクチャを採用した。即ち、すべてのデータにそのデータタイプを表現するタグビットを付加することにもタグビットを高速に判定するための専用ハードウェアを設けた。

4.4 オペレーティングシステムのサポート機能

4は効率のよいプログラム作成環境もユーザーに提供することが求められておりこのための対話型オペレーティングシステムの開発が予定されている。このオペレーティングシステムはエンドユーザインタフェースシステム(コマンド・インタプリタ、ウィンドウシステム等)、プログラミングシステム(エディタ、デバグガ等)あるいはファイルシステム等の各種のソフトウェアシステムから成る。システムの均質性を重視していることからオペレーティングシステムについても述語論理の枠組で記述されるため、4のサポートは必要である。

4のオペレーティングシステムサポート機能については、前述述語を効率よく実行するという以外にハードウェアレベル、ファームウェアレベルでのサポートを設けている。即ち、記憶管理システムにおけるメモリ領域の割り当て、ガーベッジコレクション処理等は直接マイクロプログラムで記述される。また、プロセス管理システムでの

プロセススイッチ機構もマイクロプログラムで記述される。ψではさらにプロセス情報も高連専用メモリ内に格納している。これは核言語KLMの解釈/実行のために生成される実行環境が通常の言語のものよりも大きく、プロセス環境を主記憶に格納する場合のメモリアクセスオーバーヘッドが小さくなるためである。

このようなファームウェア、ハードウェアレベルのサポートその他にオペレーティングシステムの核部のためのサポートとしてカーバッドコレクションの対象とならない領域を設けた。この領域に格納されたプログラムはカーバッドコレクションの実行中でも処理が可能となっており、特にここに格納されたプロセスはSupra G.C. プロセスと呼ばれている。この特別な領域を設けた理由は、カーバッドコレクションの実行途中においてもオペレーティングシステムの核部は実行できる形にしておいた方が対話型システム環境では望ましいためである。

5. ψのアーキテクチャ

4章で述べた様々な設計上の考慮に基づいてψのマシンアーキテクチャが定められている。

ψは基本的にはワードマシンであり、データ、機械語命令は一語を単位としてアクセスされる。またデータタイプ判定の高連化のためにすべてのデータにはタグが付与されておりタグマシンでもある。

以下にψのアーキテクチャ上の特徴のなから、アドレス空間とデータ及び機械語命令の内部表現について述べる。

5.1 アドレス空間

ψは32ビットのアドレス空間を持っており、メモリアクセスは図3に

示すように8ビットのエリア番号と14ビットのページ番号、10ビットのページ内オフセットという論理アドレス形式でアクセスが行われる。

ψの実際のメモリシステムは仮想記憶ではなく実記憶システムであるが、核言語KLMのプログラムからは論理的に16M語まで拡張できる256枚の"エリア"と呼ばれる論理アドレス空間により構成されているように見える。各エリアは全く独立した空間でありψのオペレーティングシステムにより管理される。このようなエリアの概念を導入した目的は次の理由に依る。

- 1) 核言語KLMの解釈/実行のために必要とされる各種スタック領域を論理的に別空間、即ち別々のエリアに割り当てることによって、同一空間に複数のスタックを生成する場合に生じるスタック領域同士の衝突の問題を論理レベルからなくしたかったこと。
- 2) ψは論理型プログラミング環境にプロセスを導入することを目標としておりプロセス毎に独立の空間を与えたいため。これらが"エリア"の概念を導入した理由である。

プログラムの実行環境として必要とされる各種スタックを別々のエリアに割り当てるとして1つのプロセスの実行には4枚のエリアが必要とされる。プログラム格納領域(ヒープ)は共有を許しており、これに4エリアを割り当てるとすると、256枚のエリアからは、最大63個までのマルチプロセスをサポートできる。

各エリアの空間は更に1K語の大きさの"ページ"を基本単位として管理されている。即ちエリアがスタックやヒープとして使用される場合にその領域の増減の基本単位は1K語である。

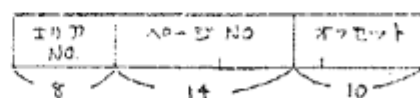


図3 論理アドレス

図4にアドレス変換のメカニズムを示す。論理アドレスから物理アドレスへの変換はテーブルを2回参照することで行なわれる。即ちまずエリア番号によりエリアテーブルがひかれこれによりページテーブルのベースが生成される。ページテーブルベースとページ番号を加算した結果でページテーブルをアクセスすると物理ページアドレスが得られる。これをページ内オフセットと結合することにより計24ビットの物理アドレスが生成される。

大型汎用機ではアドレス変換の高速化のためにTLB(Translation Lookaside Buffer)を持ち、求める論理ページエントリがその中にあれば主記憶中の変換テーブルを参照することなく高速に物理ページエントリを求める方式がよく用いられる。

ここではこの方式は採用せずにエリアテーブル、ページテーブルを専用の高速メモリ内に持つこととし、この専用ハードウェアによりアドレス変換は常に1マシサイクルで行なわれる。この方式を採用したのは以下の理由に依る。

- 1) ミスヒットが発生したときTLB方式では主記憶上の変換テーブルを参照せねばならず、それによるオーバーヘッドが存在すること。
- 2) 4のキャッシュメモリは論理

アドレスで管理されており、キャッシュミスヒット時には常に論理アドレスから物理アドレスへの変換が必要である。4ではキャッシュミスヒット時に備えてキャッシュメモリへのアクセスと並行してアドレス変換を行うよう設計しており、ミスヒットが判明した時点で対応する物理アドレスが生成されている。このようなアドレス変換方式でTLB方式を採用すると更にTLBでのミスヒットも管理せざるを得ず、メモリシステムが複雑になる。

3) ガーバッジコレクシヨニ実行中では、すべてのメモリ空間を探索する必要があり、その参照アドレスポインターにはあまり局所性のないことが予想される。従ってTLB方式ではあまりヒット率を高くすることはできず、TLBミスヒットによるオーバーヘッドを伴うことが予想される。

4) エリアテーブルのエントリは256語であるが、ページテーブルのエントリはエリア当たり16kページを含むので論理的には $256 \times 16k = 4M$ 語必要である。しかし4では仮想記憶方式を採用していないためシステム全体で使用できるページテーブルエントリの総和は実記憶のページ数以上大さくならない。即ち4の最大実装主記憶の容量は16M語であるからページテーブルエントリの総和は16k個を越えないため変換テーブルをすべてアドレス変換ユニット内の高速メモリチップ上に格納することが可能である。

このような理由によりTLB方式は採用されなかったが、仮想記憶方式を用いる場合には有効な手段である。

エリアは論理的に各最大16M語まで拡張可能であるが、あるプログラムが使用するエリアがそれぞれ最大どれだけのページを必要とするかは事前に予測できない。更に1つのプロセスは少なくとも4つのエリアを使用するかプロセス毎にその使用状況は異なる。従って1枚のページマップメモリ上にすべてのエリアに対する使用中のページに関するページテーブルを格納する方式では、エリアのページ数の増加により他のエリアのページテーブルとの衝突が生じるこ

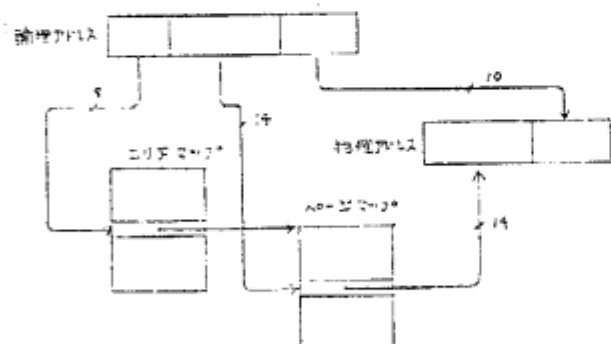


図4 アドレス変換メカニズム

ともあり得る。これをなるべく避けるためには実装ハードウェアの取扱程度ハードウェアが登録できるハードウェアメモリを用意してその中にはエリアに対するハードウェアテーブルをなるべく分散させて配置するのが望ましい。ここでは標準実装全記憶4M語に対しハードウェアメモリの大きさは16K語とし標準実ハードウェアの4倍を確保することとしている。もしもハードウェアメモリの中でハードウェアテーブル同士の衝突が生じた場合には割り出し(Trap)が生じハードウェアテーブルの再配置を行なう。このアルゴリズムも種々考えられるがどの方式がよいかは今後の課題である。

4.2 データ表現

4はノーマルマシンであるが、今後の論理プログラムの大規模化にも対応できるように、上述のアドレス空間の拡張とともにデータ表現能力も強化されている。

4の一語は図5に示すように32ビットのデータ部と8ビットのタグ部から成る40ビットで構成される。タグ部の上位2ビットはカーバシンコレクション用のマークビットであり、次の6ビットはデータを識別するためのデータタグである。データ部は32ビット幅のデータもしくは論理アドレスを表現している。4の取り扱うデータのすべてにはタグが付与されており、全記憶上のすべてのデータは一意的に識別される。

ユーザーが取り扱うデータタイプとして4は次のものを用意している。

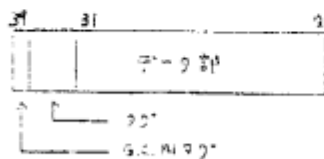


図5 ワードフォーマット

1) シンボル

シンボルは校書語KLMのアドレス情報を表現することを目的としているが、直接文字列表現との対応関係はない。この対応付けはオペレーティングシステムにより管理されている。

2) 整数、実数

通常のマシンが扱う整数、実数と同一のものである。

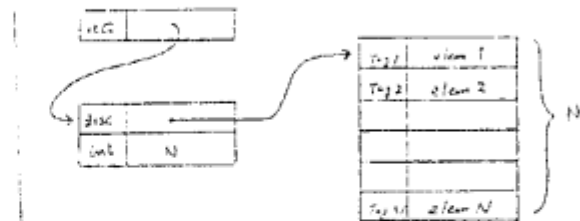
3) ベクター

ベクターは決められた要素数からなる一次元配列であり、2進木リストをはじめとする各種構造体データを表現するのに用いられる。ベクターは標準形として図6に示すようにベクター記述子を介してアクセスされる。ベクター記述子のキーワードはベクターの要素数を示すものである。

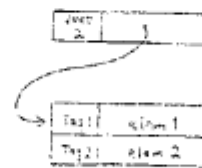
しかしこのような表現方法ではベクターへアクセスする際に必ずベクター記述子の読み出しが必要でありオーバーヘッドがある。例えば2進木リストのような要素数の少ない小規模のベクターは頻繁に利用されることが予想され、このようなベクターのアクセスを高速化する目的で4ではデータタグ部に要素数情報を埋め込んだ「ダイレクト型ベクター」を設けた。

4) ストリング

ビット列や文字列の表現のために4ではストリングが用意されている。ストリング



(a) 標準ベクター表現



(b) ダイレクト型ベクター表現

図6 ベクターの表現

の表現はベクターと同じように記述子を介してストリング本体へアクセスする方法を用いる。

ベクターとストリングの違いは次のとおりである。

- (a) ベクターの中には変数を含むことが許されているがストリング中には含むことは許されない。
- (b) ベクターの各要素は1語の任意データでよく別のベクターを要素として持つことも可能であるが、ストリングの各要素は1ビット、1バイト、2バイトのいずれかでなくてはならない。

次に、プログラムの内部形式中に現われるデータタイプとしては次のものがある。

- 1) ローカル変数
- 2) グローバル変数

これらは共に述語もしくは構文体のローカル変数、グローバル変数であることを示すとともに変数セル位置を指示するのに用いられる。

- 3) 述語呼び出し

クロース中に表われるゴール呼び出しの述語名はコンパイラにより呼び出すべきクロースに対応する内部形式へのポインタに置き換えられ、タグ部にはコードを示す情報をつけられる。

- 4) 組込述語

中はワードアクセスを基本としているためクロースの内部表現も1要素/語の形式で表現される。しかしながら組込述語の処理に於いては引数データの読み出し処理が通常のユニフィケーションに比べて簡単である。このことを利用して組込述語に対しては図7に示すように基本的には1語のコンパクトな内部表現形式を採用した。タグ部では組込述語であることが示され、デ

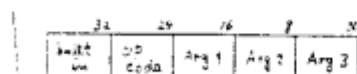


図7 組込述語のフォーマット

ータ部にはオペコードころつまでのコンパクト型の引数が格納される。

ここではさらにプログラムの実行時にマイクロインタプリタによりダイナミックに生成されるデータタイプも存在する。代表的なものをいくつか上げると、

- 1) しつアレス

変数同定がユニファイされたときに生じるポインタを示すものである。

- 2) もしキュール

変数と構文体がユニファイされたときに生じるものであり図8に示す表現がとられる。

- 3) コントロールブロック

クロースの実行に伴って発生するクロース環境情報を保存するためのデータブロックであり、その中にはユニフィケーションを行なうのに必要な実行環境情報やバックトラック、呼び出し元クロースへのリターンを制御するための情報が格納される。

4.3 クロース表現

抜き語KLBにおけるクロースは PROLOG の仕様と同じであり、ハット部とボディ部の組合であるボディ部から構成される。1つのクロースはコンパイラにより図9に示すような語の並びである内部形式に変換される。クロースハットに対応

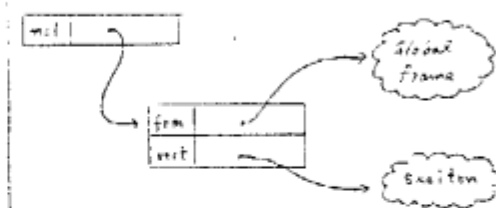


図8 もしキュールの表現

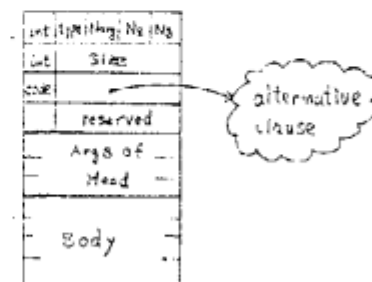


図9 クロース内部表現

する部分はヘッダ部とヘッド引数から構成される。ヘッダ部は4ワードから成り、ヘッドの引数個数を示す Narg, クローズに含まれるローカル変数, グローバル変数の数を表わす Nl, Ng とクローズのタイプ(コンソートフロー等)を示す Type フィールドから成る1語の他に、このクローズ内部形式のテーブルサイズを示す1語、他の置換肢へのポインタを格納した1語等から成る。クローズヘッドに含まれる引数は内部形式の引数部に展開される。

ゴールの結合のされ方には、次の3つのタイプを基本形式として認けている。

1) AND 結合

ゴールとゴールが探索木では and 木として表現されるものであり、各ゴールは呼び出し述語へのポインタとその引数が展開された形で表現される。この結合の例を図10に示す。

2) OR 結合

ゴールとゴールが探索木では OR 木として表現されるものである。実行順序はこの結合されたクローズの先頭以外はバックトラッキングの対象となるような制御を取る。

```
append (1, X, X).
```

```
append ([A|X], Y, [A|Z]) :-  
    append (X, Y, Z).
```

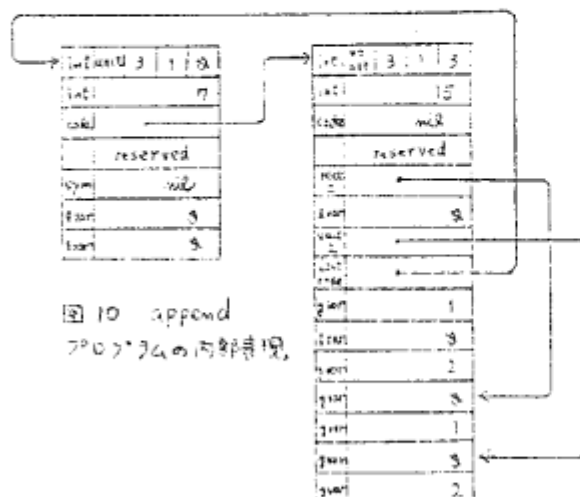


図10 append
プログラムの内部構造

3) CASE 結合

この形で結合されたゴール同志は Index 付けられたものと見なすことができる。この形で結合されたゴールの一つが実行時に選択されるが、バックトラッキングの対象とはならない結合のされ方である。

5. おわりに

本論文ではECOTを中心に、開発を進めているパーソナル逐次型推論マシン、その開発目標、設計思想及びアーキテクチャの概略としての機械語の内部形式、データ表現、メモリ空間について述べた。

現在アーキテクチャ設計は完了し、ハードウェア試作に向かって詳細設計を進めているところである。性能としては DEC-10 PROLOG コンパイラ版(DEC 2260)相当の性能を実現できる見通しである。また機械語の内部形式については TRC 等の最適化手法の導入において、現在の語を基準とした形式ではなく組合せ語と同様のコンパクトな表現形式も検討中である。今後、このマイクロプログラムレベルでの種々の最適化技法についても検討してゆく予定である。

最後に、本研究の機会を与えて下さった第一博研定所長、村上国男博士研究室長、討議を通じて有益な助言を与えて下さった近山隆夫はじめ ECOT メンバ諸氏に深謝する。

参考文献

- Warren, D.H.D: Implementing Prolog
- compiling predicate logic program
D.A.I. Research Report no 39-40 (May 1977)
- 近山他: 検査語オメガ版の仕様
Dバックプログラミングコミニティ'83 (5月1983)
- Warren, D.H.D: An improved Prolog
Implementation which optimises
Tail Recursion
Proc. of the Logic Programming
Workshop, Hungary (July, 1980)
- Boyer, R.S and J.S. Moore: The sharing of structure in theorem proving programs
Machine Intelligence Vol.1, 7
Edinburgh Up (1972)