

47—S 002

コンパイラ・ジェネレータの開発

昭和 48 年 3 月



財団法人 日本情報処理開発センター

この報告書は、日本自転車振興会から競輪収益の一部である
機械工業振興資金の補助を受けて昭和47年度に実施した「コ
ンパイラ・ジェネレータの作成」の成果をとりまとめたもので
あります。

序

当財団は情報処理に関する各種の研究開発をおこなっておりますが、本報告書はその一環としてとり上げた「コンパイラ・ジェネレータの開発」の成果をまとめたものであります。

現在、コンパイラ言語はコンピュータのプログラミング言語として最も広く使用されておりますが、新機種のための、あるいは新しい言語のためのコンパイラを、その都度開発する努力は多大なものであります。この労力を軽減し、かつ信頼性のさらに高いコンパイラを得るために、コンパイラ・ジェネレータが有効な手段であることはよく知られております。

しかし、ジェネレートされたコンパイラの効率の面、ジェネレータへの入力の方法などにまだ種々の問題があり、いま一步、実用化への努力が必要とされております。

本研究開発は、これらの問題点の原因を解明し、実用に供し得るコンパイラ・ジェネレータの試作を目ざすものであります。

この報告書が各方面に利用され、わが国の情報処理技術発展の一助として寄与できるよう願います次第であります。

昭和48年3月

財団法人 日本情報処理開発センター

会長 難 波 捷 吾

ま え が き

コンパイラ言語の構造と構文解析手法に関する研究は1960年頃から主として大学・研究機関において行なわれ、その成果として各種のコンパイラ・ジェネレータが開発されてきた。

コンパイラの作成における多大の労力を軽減し得るという理由でコンパイラ・ジェネレータには多くの期待がかけられているが、実用上ではいまだ多くの問題点が残されている。

例えば、セマンティックス記述においては、シンタックス記述のような形式的な記述方法がないこと、構文解析手法が一意的であるコンパイラ・ジェネレータによって作成されたコンパイラは、構文解析が各コンパイラ言語のレベルに最も適した手法で行なわれる hand-made のコンパイラよりコンパイル速度、メモリ占有量等において劣っていることなどである。

コンパイラ・ジェネレータ JPAC (JIPDEC Automatic Compiler generator) においては構文解析手法としては言語の複雑さによって SLR(k) 文法と L(m)R(k) 文法のいずれかが自動的に選択され、セマンティックス記述としてはシステム記述言語 L S D 及び標準セマンティックス・ルーチンを用いる方法を採用した。

主眼としてはジェネレートされたコンパイラの効率が実用に耐え得るレベルを目標としている。

本報告書は2部にわけ、第1部は主として米国において開発された各種のコンパイラ・ジェネレータを紹介し、第2部にはJPACの内部および外部仕様を記述してある。また付録として1968年までのコンパイラ・ジェネレータの著名な紹介論文である Translator Writing Systems の全訳を掲載した。

本研究開発は、2カ年計画でおこなっており、47年度には仕様の決定、システムの設計および作成の一部を終了し、48年度には標準セマンティックス・ルーチンの整備、JPACによるコンパイラと hand-made のコンパイラとの性能・労力などの比較をおこなう予定である。

総 目 次

第1部 コンパイラ・ジェネレータの調査

1. 概 説 1
2. 各種コンパイラ・ジェネレータ 4
3. 構文解析手法 87

第2部 コンパイラ・ジェネレータの作成

第Ⅰ編 コンパイラ・ジェネレータ JPAC 99

1. JPAC の概要 99
2. JPAC の動作概略 101
3. JPAC の言語仕様 105
4. STAGE 内部処理 124
5. 遷移テーブル類 157
6. 中間テーブル類 163

第Ⅱ編 LSD 外部仕様書 171

1. プログラム要素 171
2. データ要素 175
3. 式 187
4. 宣 言 191
5. 各ステートメントの説明 200
6. LSD と FASP のリンク 224

参考文献一覧 235

付 録 239



第 1 部

コンパイラ・ジェネレータの調査

コンパイラ・ジェネレータあるいはコンパイラ・コンパイラの研究は、1960年頃から主として大学、研究機関に於ておこなわれてきたが、1968年 Stanford 大学の J. Feldman, D. Gries によって、米国におけるそれまでの研究成果がまとめられた。

"Translator Writing Systems"

Communication ACM, Vol. 11, No 2

がそれである。

この報告書では Translator Writing Systems (TWS) 以後に発表されたシステムを中心に、それぞれのシステムのねらいおよび手法の概要を紹介するが、代表的なものはあえて TWS と重複させた。

なお TWS の全訳を付録として巻末に収録した。

目 次

1. 機 説	1
2. 各種コンパイラ・ジェネレータ	4
2.1 XPL	4
2.2 MPL/I	16
2.3 TREE META	19
2.4 META PI	25
2.5 VITAL	37
2.6 CO	58
2.7 CABAL	65
2.8 GULP	73
2.9 INSCAN	80
2.10 CANONIC TRANSLATOR	84
3. 構文解析手法	87
3.1 Weak precedence 文法および右順位文法	87
3.2 Total precedence 文法	90
3.3 SLR(k) 文法	94

1. 概 説

ランゲージ・プロセッサの開発には非常に多くの労力を要することが知られている。たとえば、超大型コンピュータの FORTRAN コンパイラ作成のために要する労力の試算では 8 ~ 15 人年を必要とするという報告がなされている。

ここ数年来この労力を軽減するために種々の方法が試みられている。現在コンピュータ・メーカーの提供しているコンパイラの多くはまだアセンブラ言語で記述されているが、FORTRAN または PL/I のような汎用言語の使用、APL, LISP などの特殊言語の使用、あるいは BPL のようにシステム記述言語を作成し、これによって記述することなどにより労力の軽減をはかることができる。

これ以外にメタ言語によりコンパイラを作成する方法もあり、汎用のコンパイラ言語等によるシステム作成よりも労力の削減が期待できるといわれている。

コンパイラ言語は Context Sensitive 文法の類に属するので Canon システムのように、この類の文法を記述し、解析する手段を提供するものがある。

しかしながら、その処理速度が実用化するにはほど遠いシステムになってしまったため、ほとんどのものは、コンパイラの記述を Syntax 記述と Semantics 記述に分離して考え Syntax は Context Free 文法の範囲で記述し、Context Sensitive 文法になる部分是对応する Semantics という形で記述する方法をとっている。

Syntax 記述に関しては、バックス記法 (BNF) というシンプルな方法があり、BNF あるいは extended BNF 等を使用することでこと足る。

Syntax の解析には、TREE META のように recursive 機能を使用することにより下向きに解析するものと、順位文法等の Context Free 文法の一種を使用して上向きに解析する手法とがある。

TREE META のような方法は構文解析においてあともどりをおこなわなければならない。あともどりとは生成規則の適用の各段階を、1 段階ずつ逆戻り

して、それまでに試されなかった規則を適用することによって、他の可能性を調べることである。

この方法は、解析に必要なテーブル類は小さくなるが、解析速度が遅いという欠点を持っている。

順位文法等の方法では、あともどりは起こらないため解析速度を速くすることができるが、Syntax 記述をその文法にあうように書き変えなければならないこと、およびテーブルが大きくなる欠点を持っている。

構文解析に関しては、どちらの方法をとるにしろ決定的によいものはないので今後の研究が望まれるが、FORTRAN クラスのコンパイラならば解析できるところまでできている。

Semantics の記述に関しては Syntax 記述における BNF のようなよい記述方法が開発されてはいない。このためこの部分に関する技術の確立が待たれる。

Semantics の記述に必要な機能は、コンパイラ記述言語に必要な機能とほぼ同様で次のようなものがあげられる。

1. テーブル類のとりあつかいができる。
2. スタックのとりあつかいができる。
3. ビット、文字のとりあつかいができる。
4. 整数の簡単な演算ができる。
5. アドレスを作り出すことができる。
6. リスト、tree 等のあつかいが何らかの形でできる。
7. 入出力機能

現在までに開発されたコンパイラ・ジェネレータは、上記機能を有する言語を用意して、これによって自由に Semantics を記述するもの (XPL, MPL/I) と、これ等の機能を持つ標準ルーチンをいくつか用意して、これ呼び出

す方式 (GULP, META) のいずれかによっている。

後者は、前者よりもユーザにとってはコンパイラを作る場合に楽ではあるが
どのような標準ルーチンを用意すればよいかが問題となる。

2. 各種コンパイラ・ジェネレータ

2.1 X P L コンパイラ・ジェネレータ

X P L コンパイラ・ジェネレータは、スタンフォード大学においてW.M. McKeeman 他によって1968年にIBM/360に対してインプリメントされた。

これは、B N Fで書かれたSyntax規則を読み込んで、後のプロセスで必要となるテーブルを作り出すANALYZER, Semantics 記述言語であるX P L言語のコンパイラXCOM, およびジェネレータの対象となるコンパイラのソース・プログラムを読み込んでANALYZER で作り上げたシンタックス・テーブルにしたがい必要なセマンティックス・ルーチンにコントロールをわたす。X P L言語で作られたSKELTONより構成されており、その関係は(図2-1)X P Lコンパイラ・ジェネレータの構成に示すとおりである。

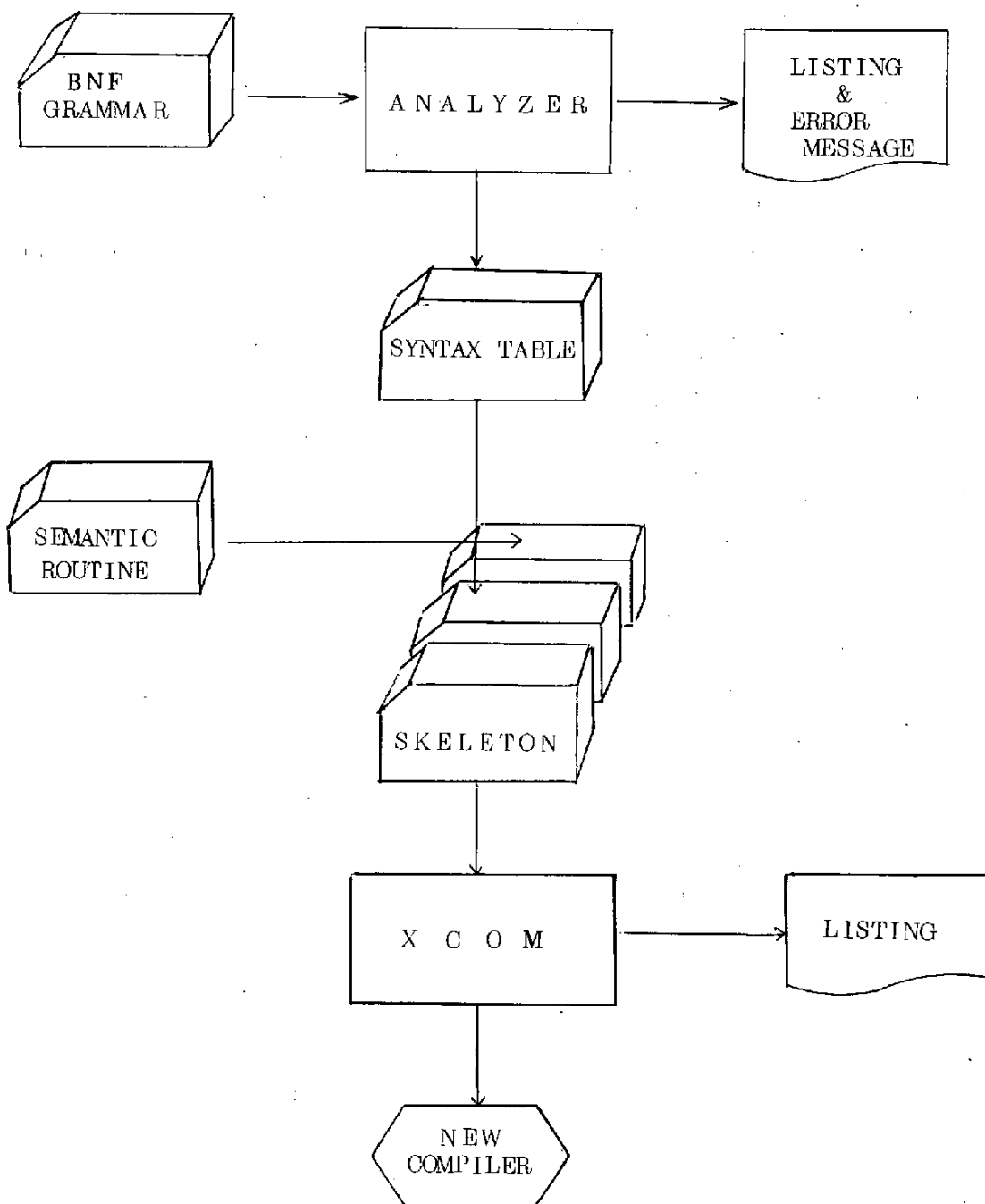


図 2-1 XPL コンパイラ・ジェネレータの構成

ANALYZER

ANALYZERは構成図より明らかなように、単なるシンタックス・テーブル・ジェネレータであって、これが受け付ける文法は $(2, 1; 1, 1)$ 位のMixed Strategy Precedence文法(以後MSP $(2, 1; 1, 1)$ と略称する)である。MSP $(2, 1; 1, 1)$ というのは、スタック中の記号と新しく入力したterminal記号を見てstackをするか、あるいは生成規則を適用してreduceするかという決定を $(2, 1)$ 位のprecedence文法によっておこないreduceの決定がなされた場合において、生成規則の選択を $(1, 1)$ 位のbounded contextによって確認する手法を意味している。

precedence文法の解析手法は、通常順位マトリックスをもちいる(順位関数を作る場合もあるが、常にこの関数ができ上がるとはかぎらない)。

順位マトリックスの大きさは、 $(1, 1)$ 位の場合

$$(\text{記号の数}) \times (\text{terminal記号の数})$$

となるが、 $(2, 1)$ 位の場合

$$(\text{記号の数})^2 \times (\text{terminal記号の数})$$

となり、マトリックスは非常に大きくなってしまう。

幸いにして、 $(2, 1)$ 位の文法において、stackあるいはreduceの決定のほとんどは $(1, 1)$ 位でおこなうことができる。このために、順位マトリックスとしては $(1, 1)$ 位に対応する大きさだけ用意する。 $(1, 1)$ 位の場合、マトリックス要素としてはstack, reduce, エラーの3種のみでよいが、 $(2, 1)$ 位の場合には、これに不明という情報を追加しておく。

不明の場合には、あらたに別のテーブルを調べて、stack, reduceの決定をすることになる。すなわち、スタックの頭の2個の記号と、新しく入力した記号の計3個の記号をテーブルにしておき、stackあるいはreduceを決定する。この時、3個の記号の連なりがエラーであるかどうかのチェックをしないとすれば(エラーチェックを省略しても、生成規則を適用しようとして生成規則のテーブルをサーチした時、登録されていないのでエラーの発見はできる。)

stackあるいはreduceのいずれか一方だけをテーブルに登録しておけばよい。このため、あまり多くない情報の追加で(2, 1)位のPrecedence文法の処理が可能となる。

ANALYZERの入力は、BNFで記述した文法であり、上記の思想により次のものがカード出力される。

1. V: terminal, nonterminal 記号テーブル。
2. C1: (2, 1)位Precedence文法に対応するstack, reduce, 不明, エラーを示すマトリックス
3. C1 TRIPLES: C1において不明の場合に調べるテーブルであってstackすべき場合のみがのっている。
4. PRTB: 生成規則の選択のためのテーブル
5. PRDTB: セマンティックス・ルーチンの番号テーブル
6. HDTB: 生成規則が適用された時、どのようなnonterminal記号に還元されるかを示すテーブル
7. Context-Case: reduceの決定がなされた時、生成規則の選択にあたってContextを調べる必要の有無を示すテーブル
8. Left-Context, Context-Triple: 生成規則の選択にあたってそれぞれ(1, 0) Context, (1, 1) Contextを調べる必要のある時のテーブル。
9. その他の補助情報

ANALYZERにおける特徴としてあげられるものに、入力BNF記述の自動変更がある。bounded contextアルゴリズムにおいて“insulting comma” — 区切り記号としての役割を果たすカンマ等 — の問題がある。この問題に関しては、カンマに対応するnonterminal記号を定義し、それに

よって影響を受ける生成規則を書き換えることによりしばしば解決できる。このためXPLにおいてこのような場合には自動的に新しいnonterminal記号および生成規則の追加，変更をしてANALYZERが受け付けられるように自動変更している。

SKELTON

SKELTONはANALYZERで作り上げたマトリックス・テーブル等にしたがって構文解析をし，生成規則の適用(reduce)が決定すると，生成規則番号をパラメータに持ってユーザが用意したセマンティックス・ルーチンにコントロールをわたす役割をする。すなわち，目的としたコンパイラの共通メインルーチンとみなすことができる。

(図2-2)によりMSP(2, ; 1, 1)における構文解析アルゴリズムの概略を示し(図2-3)でXPLが実際におこなっている方法を示す。なお，(図2-3)では説明を簡単にするために，実際に使われている記号をや、変えている。

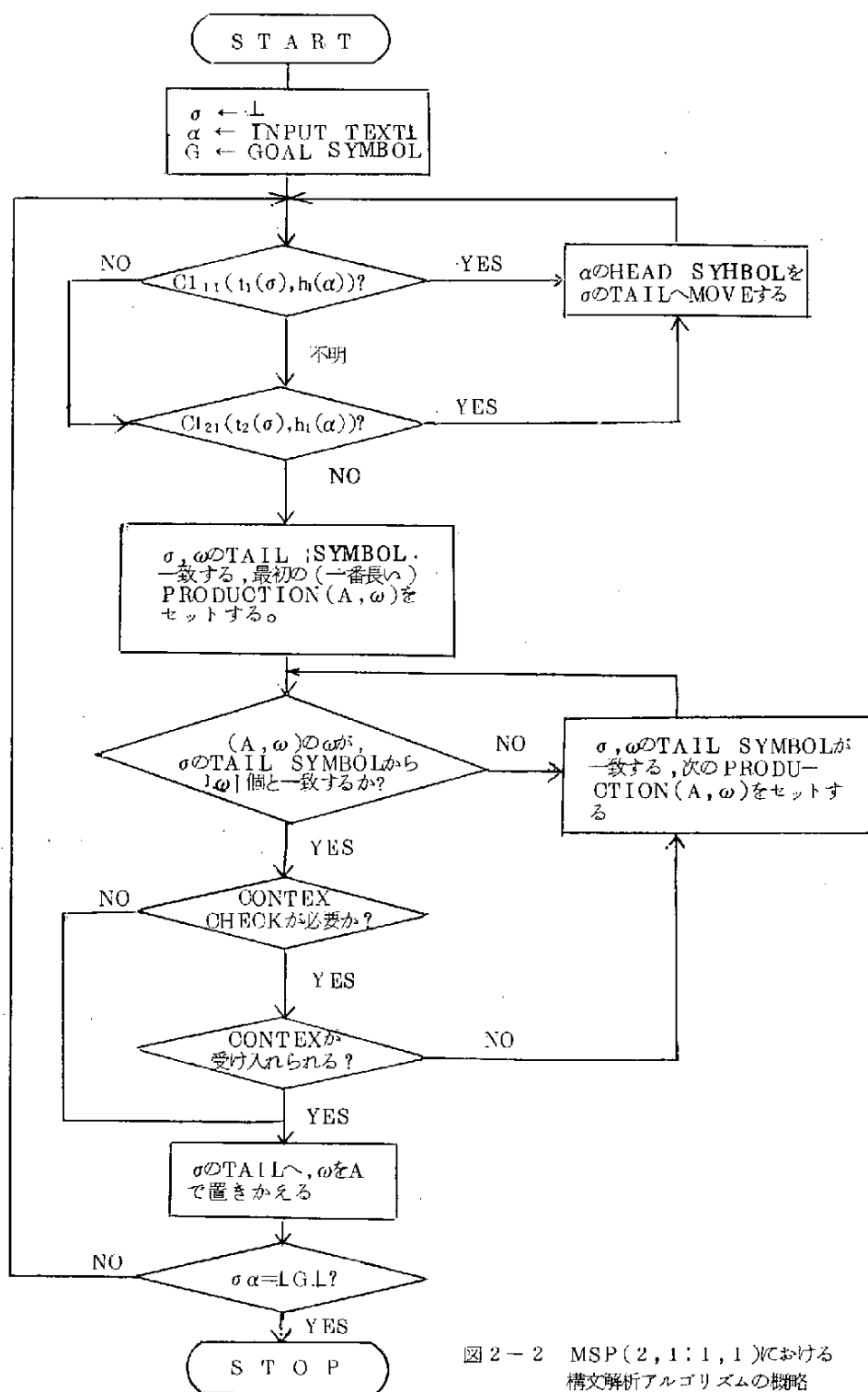


図 2-2 MSP(2,1:1,1)における
構文解析アルゴリズムの概略

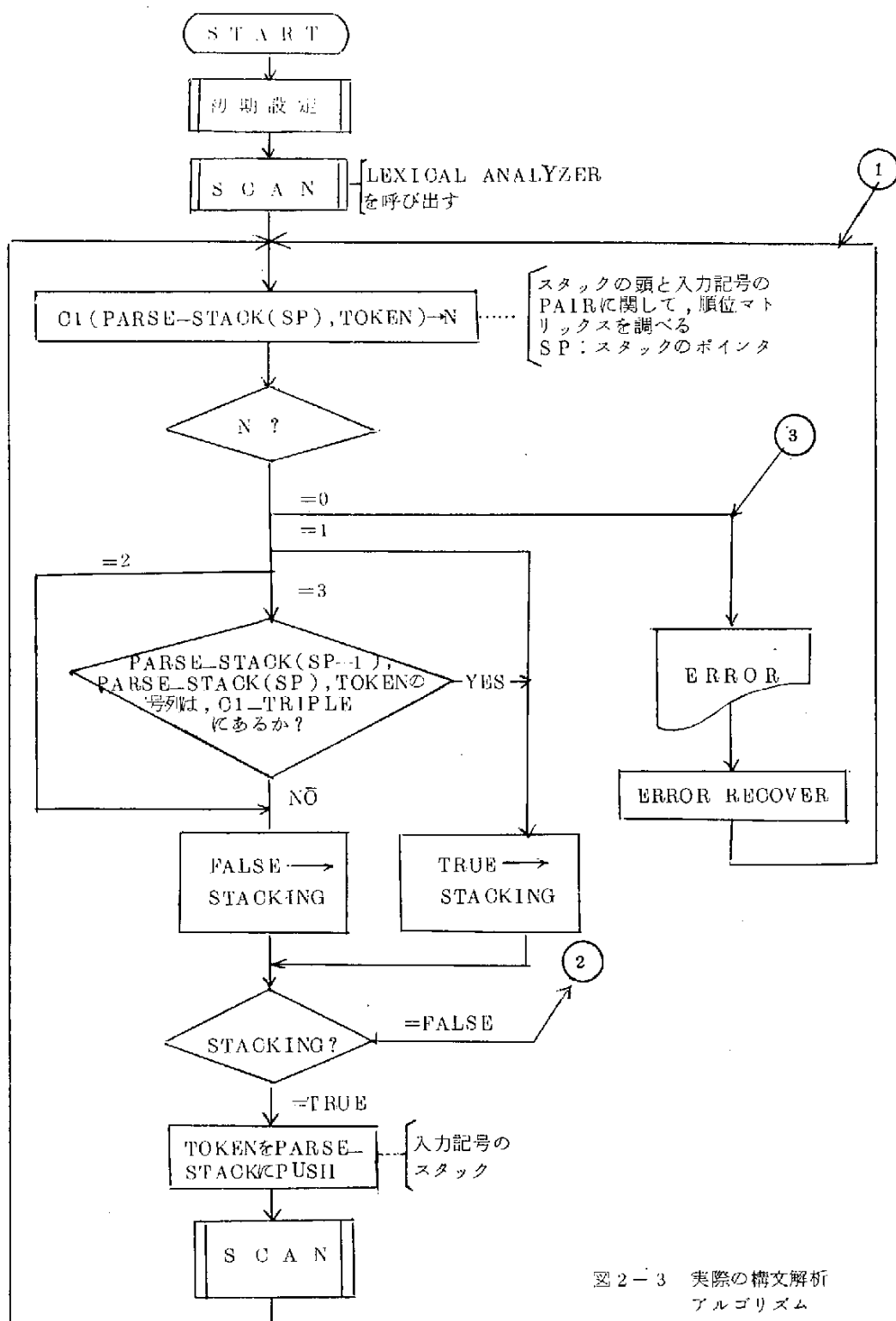
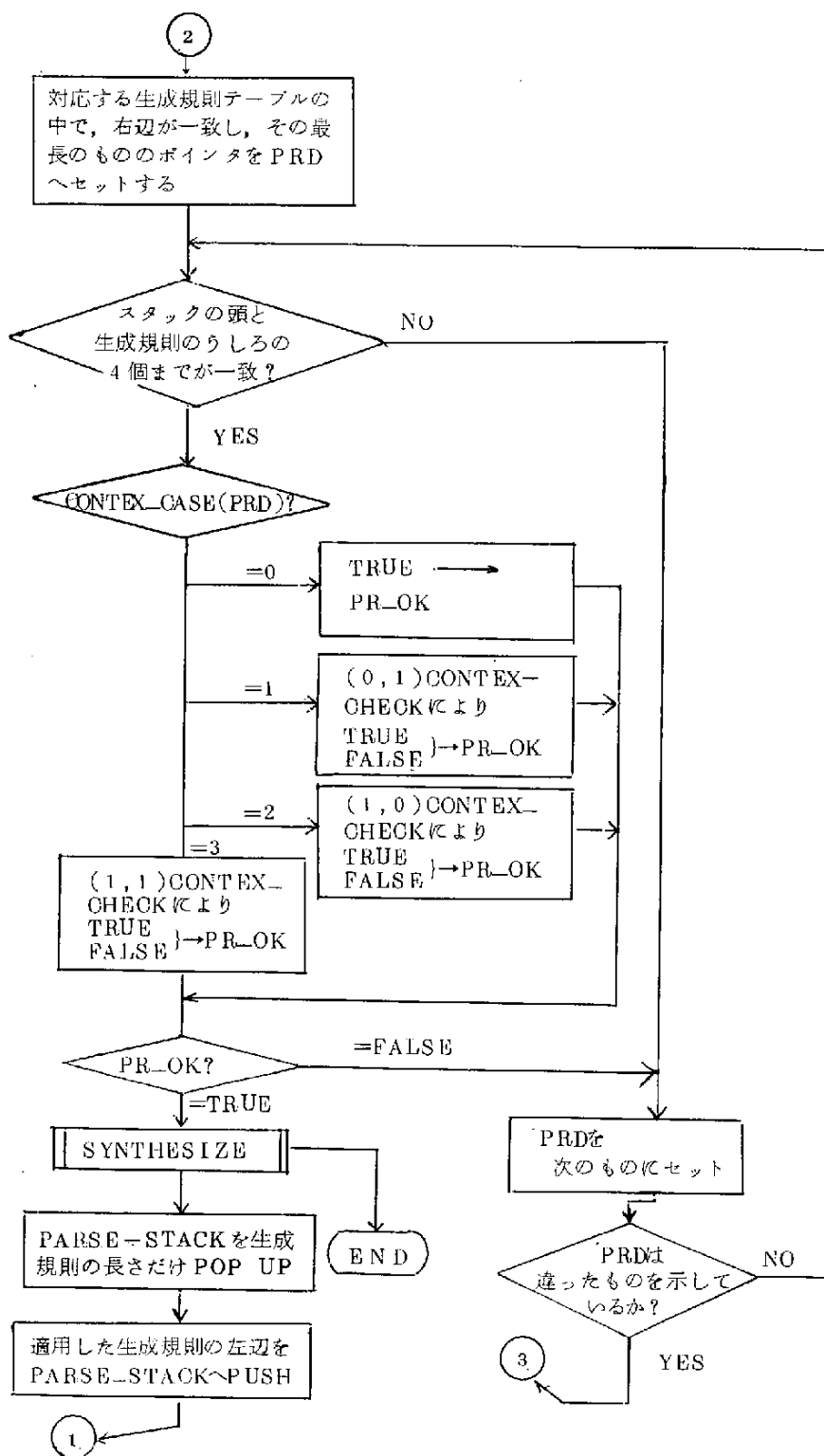


図2-3 実際の構文解析
アルゴリズム



(図2-3)の中にあるSYNTHESIZEルーチンがセマンティックスの記述部分である。また、必要に応じてLexical AnalyzerであるSCANルーチンを組み込む必要がある。

XPL言語

XPL言語はPL/Iのサブセットに若干の変更を加えたものである。PL/Iに対する制限なり変更は、PL/Iの仕様の中でコンパイラ作成に必要な機能を残し、PL/Iの仕様にはないが、あった方が便利な機能を追加するという方向でなされた。

PL/Iから取り除いた仕様

- ① すべての変数は、それが使われる以前で宣言しなければならない。
- ② 配列の下限は、0である。
- ③ 属性等を指定する時などに、省略形はない。(CHARACTERをCHARと指定してはいけない)。
- ④ データのタイプはFIXED, CHARACTER, BITの3種類である。
- ⑤ 属性はカッコでくくって指定することができない。
- ⑥ ビットストリングのとりあつかいが違っている。
- ⑦ 文字ストリングはすべてVARY属性を持ち0文字目から始まる。
- ⑧ DOループは正でなければならない。
- ⑨ PROCEDUREはnon-recursiveであってvalue parameterのみを持つ。
- ⑩ DO, IF等はreserved wordで名標として使えない。

PL/Iを更に拡張した仕様

- ① LITERALLY 属性を持つデータがある。……指定された変数をコンパイル時に定数におきかえる。

- ② 2, 4, 8, 16 進数の表現ができる。
- ③ SHL(シフト・レフト), SHR(シフト・ライト)という標準関数がある。
- ④ FILE, INPUT, OUTPUTと書く(代入文等に)ことにより入出力ができる。
- ⑤ CASEステートメントにより一連のステートメントの一部を実行できる。
- ⑥ INLINEにより機械語の命令を書いて実行できる。

このXPL言語によって必要とするコンパイラのSemanticsを記述するわけである。すなわちコンパイラのおこなうべき処理手順をプロシーディアルに記述し必要なオブジェクトコードを出力することがXPL言語の役割である。

この言語の特徴はDOステートメントの拡張によりgo to ステートメントをほとんど使わないでコーディングできるために、他人がプログラムの解読する場合に非常にわかりやすい。

このシステムは、作成にあたって次のようなブート・ストラッピングによる方法がとられた。

まず次の記号の定義をする。

- (1) $(a \rightarrow b, c)$ は a から b への変換プログラムが c でコーディングされている。
- (2) $(a \rightarrow b, c) \times (c \rightarrow e, f) \rightarrow (a \rightarrow b, e)$ は $(a \rightarrow b, c)$ のプログラムをトランスレータ $(c \rightarrow e, f)$ で実行させることにより $(a \rightarrow b, e)$ ができあがることを示す。

XPLのインプリメントにあたってまず

$(XPL \rightarrow S/360ML, ALGOL)$

$(XPL \rightarrow S/360ML, XPL)$

の2つが作られた。そして

- ① $(XPL \rightarrow S/360ML, ALGOL) \times (ALGOL \rightarrow B5500ML, B5500EX) \rightarrow (XPL \rightarrow S/360ML, B5500ML)$
- ② $(XPL \rightarrow S/360ML, B550ML) \times (B5500ML \rightarrow B5500EX, B5500EX) \rightarrow (XPL \rightarrow S/360ML, B5500EX)$
- ③ $(XPL \rightarrow S/360ML, XPL) \times (XPL \rightarrow S/360ML, B5500EX) \rightarrow (XPL \rightarrow S/360ML, S/360ML)$
- ④ $(XPL \rightarrow S/360ML, S/360ML) \times (S/360ML \rightarrow S/360EX, S/360EX) \rightarrow (XPL \rightarrow S/360ML, S/360EX)$

注 S/360 : IBM/360
B5500 : Burroughs 5500
ML : 機械語
EX : 実行形式

のステップにより、IBM/360 にインプリメントし、③、④のくり返しをおこなって機能拡張をおこなった。

したがってXPL自身がXPLで作られているわけでありその中のXCOMコンパイラに関していえば、SKELETON が作り出すシンタックス・テーブルの大きさは約 $3.8K + \alpha$ (α はXPL言語の変数領域アロケーション方法により増加する部分であるが0.5Kバイト以下である) バイトでおさまり実用上問題ない大きさである。コンパイルスピードに関しては、資料がないので不明であるが、構文解析手法から考えてそんなに遅いとは考えられない。

XPLを使ってコンパイラを作成するのにどのくらいの工数が必要であるかということに関しては、スタンフォード大学のシステムプログラミングコースの学生が、このシステムを使うことにより3~6週間でコンパイラを作っているとの報告がされている。

XPLの欠点と思われる点は、XPL言語にはPL/IでいうEXTERNAL

属性がないので、新しいコンパイラを作成する場合に、すべてのソースプログラムを毎回コンパイルしなければならないこと、および<identifier><number>以外も terminal 記号としてあつかえるようにした方がよい点であろう。

XPL に関しては、内部仕様書、リストが公開されており調査するのに非常に便利である。

参 考 文 献

- 1 W.M.Mckeeman, J.J.Horning, D.B.Wortman
A Compiler Generator
PRENTICE-HALL INC.

(図2-2)は上記文献より引用させていただいた。

- 2 W.M.Mckeeman, J.J.Horning, E.C.Nelson,
D.B Wortman
The XPL compile Generator system
FJOC 1968 PP617~635

(図2-1)は上記文献より引用させていただいた。

- 3 W.M.Mckeeman
An approach to Computer language design
Stanford Univ. Technical Report NO.CS48, 1966

2.2 MPL/I

MPL/I は作成したい processor の仕様を記述する部分と、それを実行形式プログラムに変換する部分より構成されている。

MPL/I の Semantics 記述は PL/I のサブセットによるプロシージャルな方法によっている。一方 Syntax の記述は、BNF に似てはいるが、次の点が違っている。

- ① 生成規則の右辺が全く同じものは、1つの式にまとめその選択はユーザの定義する procedure による。

例

```
<array name> ::= <identifier>;
```

```
<function name> ::= <identifier>;
```

なる構文があった時

```
FN(<array name>|<function name>) ::= <identifier>;
```

```
FN: PROC;
```

```
.....
```

```
RETURN(N);
```

```
END;
```

procedure FN では、<identifier> を調べ <array name> であったならば 1, <function name> ならば 2 を返す。

- ② reference variable の使用

<\$> はすべての記号を代表するもの、<n> (n は自然数) は生成規則の右辺の記号で左から順に 1 より数えるものとした時 <\$> は右辺 <n> を左辺に使用できる。

例

```
<2> ::= <A><$>;
```

BNF によれば起こり得る場合すべてを記述するが、MPL/I においては、~ 以外のものという表現を許す。

例

$\langle C \rangle ::= \langle A \rangle ;$

$\langle D \rangle ::= \langle A \rangle \langle B \rangle ;$

という時これを次のように記述することができる。

$\langle C \rangle \langle 2 \rangle ::= \langle A \rangle \langle \$ - \langle B \rangle \rangle ;$

前者の意味は、 $\langle A \rangle$ のあとに $\langle B \rangle$ 以外のものがきたときにこれを

$\langle C \rangle \langle B$ 以外のもので、そのもの $\rangle$ とおきかえることを指定する。

reference variable関係は、MPL/I の構文解析アルゴリズムにおいて、あともどりが起こることを避けるためにある。

つぎにMPL/Iによる例を示す。

```
/* CONTROL */
DECLARE KWORD(4) CHARACTER(15), TYPE(4) FIXED BINARY;

      KWORD(1) = 'TOM';      TYPE(1) = 1;
      KWORD(2) = 'Y';        TYPE(2) = 1;
      KWORD(3) = 'SMITH';    TYPE(3) = 2;
      KWORD(4) = 'LI';       TYPE(4) = 2;

READ: GET LIST(CH);

/* INITIAL PRODUCTIONS */
<L> ::= A|B|C|D|E|F|G|H|I|J|K|L|M|N|O|P|Q|R|S|T|U|V|W|X|Y|Z;
'' ::= '';

/* PRODUCTIONS & SEMANTIC ACTIONS */
1 <WD> ::= <L>'';          R(1) = V(1);
2 <PW> ::= <L><L>;          R(1) = V(1)||V(2);
3 <PW>> ::= <PW><L>;        R(1) = V(1)||V(2);
4 <WD> ::= <PW>'';          R(1) = V(1);
5 KW(<PN>|<LN>) ::= <WD>;   R(1) = V(1);
6 <2> ::= ''<$>;           R(1) = V(2);
7 <S> ::= <PN><LN>;         PUT EDIT (V(1),V(2))(A(15),A(15));

/* PROCEDURES */
KW: PROCEDURE;
      DO I = 1 TO N; IF KWORD(I) = V(1) THEN RETURN(TYPE(I));
      END;
      GO TO READ;
END KW;
```

図 2-4 MPL/I の例

MPL/I に提案されたものの中で考慮するに値するのは、同一右辺を持つた生成規則のとりあつかい方である。これは順位文法において制限されていることであるが、実際にALGOLのSyntax 定義等にもちいられており、この手法をもちいることにより、とりあつかいができることである。

参 考 文 献

1. E.N.Ferentzy, J.R.Gabura

A syntax directed processor writing system

FJOC 1968 PP.637~647

(図2-4)は上記文献より引用させていたゞいた。

2.3 TREE META

TREE META は、ユタ大学において1969年に開発されたMETA 系列のコンパイラ・コンパイラシステムであって、ユーザが定義した Syntax 及び Semantics を読み込んで、その言語のコンパイラをFORTRANでコーディングしたものとして出力する。

このシステムを使用して作り出されるコンパイラは、原則として2 Pass となる。すなわち、第1 Pass で入力記号を見ながら Start 記号より始めて下向きに構文を解析し、その構文に対応する tree を作り上げる。第2 Pass では作り上げられた tree を drivenしながら object code を発生するという形をとっている。このためメタ言語の記述には、Syntax 記述及び tree 作成記述部分と object コード発生部分の2つを指定しなければならない。

Syntax 記述

このシステムの Syntax 記述は、extended BNF を修正した形でおこなう。これがBNF記述と相違する部分を取りだせば、次のとおりである。

- ① nonterminal 記号は< >でくくらない。
- ② terminal 記号は、1文字ならば▼のあとにその文字を書き、1文字以上の時は"をその前後につける。
- ③ ::= の代わりに = を使う。
- ④ | の代わりに / を使う。
- ⑤ 規則の最後に ; をつける。

これ以外に、下向き構文解析であるため、左回帰の禁止をし、その代りとして、\$() の形で()の中が0回以上くり返すことを指定できる。また nonterminal 記号的なあつかいとして basic recognizer 及び, EMPTY (null にあたる)を指定することができる。basic recognizer とは、.ID, .SR, .NUM, .LET, .CHR, .DIG の6つであり、それぞれ、identifier(英字で始まる英数字列), string(文字列を"でくくったもの),

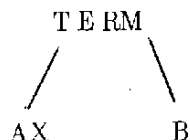
数（数字の列），英文字，文字，数字を示す。

Syntax 部分には，上記のものに：のあとに tree 作成部分を書くことができる。たとえば

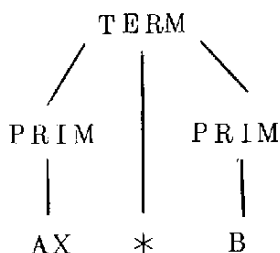
```
TERM=PRIM$(" * "PRIM:MULT{2});
```

```
PRIM=.ID/.NUM;
```

ここでMULT{2} は，node MUITから2つのnode がでていることを示す。たとえばAX*B というストリングを読み込んでtree を作れば次のようになる。



Syntax 部分に tree 作成部分を書くのは，Syntax 記述だけで作り上げられる tree



よりも冗長度が少く，tree としてとりあつかいやすいためである。

オブジェクトコード発生部分

それぞれの出力規則は，node のテストおよび出力記述より構成される。

node のテストは次のようなものがある。

① ある node からでている node の数のテスト。

ADD{一，一} はADDという node から2個の node がでているかどうかテストする。この時でている node を左から順に*1，*2，…として参照することができる。

② node の一致

ADD{—,*1}あるいはADD{*2,—} は, 第1 node と, 第2 node が等しいかどうかテストする。

③ basic recognizer とのテストは, その recognizer 名を指定する。

ADD{.ID,.NUM}は第1 node が identifier, 第2 node が number であるかどうかテストする。

④ リテラルとのテストは" リテラル #"を書く。

EQ{—,"AB"} は第2 node が"AB"であるかどうかテストする。

⑤ generateされたラベルは, 別の出力規則に持ちこされるが, このテストのために# number が使われる。

⑥ node のテストに特別の node 名を書くことができる。

STORE{—,ADD{*1,"1"}}は例えばX=X+1のように, 同一のidentifier等と, リテラルの1がADDというnodeにあるかどうかテストできる。

⑦ テストの省略は, node のあとに/をかく。

⑧ node 名を省略して同一node 名の選択子とすることができる。

ADD{—,*1}=>

{—,—}=>

これはADDの第1 node と第2 node が等しければそのあとに書かれた出力記述を実行し, そうでなければ次の{—,—}のあとの出力記述を実行する。

出力の記述は次のものがある。

① \で改行を示す(カードの終了)

② ,でタブを示す。(カラムをあわせる)

③ リテラルストリングを書くことによりそのリテラルを出力できる。

④ node : X (XはS, L, N, C)

- *1 : S *1 のリテラルを出力する。
- *1 : L *1 のストリングの長さを 10 進数で出力する。
- *1 : N もし node がストリング領域ならばそのインデックス・ポインタを, CHR ならば node に対する 10 進コードを出力する。
- *1 : C もし node が .CHR, .LET, .DIG recognizer で作られたものならばその文字を出力する。

- ⑤ # number は generated ラベルを出力する。
- ⑥ 出力記述の中にもテストを入れ, 選択子をもうけることができる。

ADD(−, −) => TNEG{*2} MINUS{*1, *2 : *1} / *1 " + "
*2 ;

これは node ADD の第 2 node が TNEG node テストで true ならば MINUS{*1, *2 : *1} を実行し, そうでなければ *1 " + " *2 を実行する。

このほか出力の記述部分には, tree の変換, エラーの処理, 他の node テストの呼び出しを書くことができる。

コード発生部は node テストと出力の記述を => でつなげて書いたものである。

例

- ① DIVIDE{−, −} => *1 " / *2 ;

は DIVIDE という名前の node から 2 つの node がでているかどうかにテストして, true ならば, 第 1 node / 第 2 node という code を発生する。

- ② ADD{MINUS{−}, −} => SUB{*2, *1 : *1}

は下図の左の tree に変換するものである。



ここで、*n:*mは現在処理中のnodeのn番目のnodeについているもののm番目のnodeをさす。この逆にtreeをさかのぼるものに↑n*mがありこれは現在処理中のnodeのn代前のnodeのm番目のnodeをさす。

例 BNFで記述した時次のような形になるものをTREE METAで記述しSemantics記述を加える。

```

<EXP> ::= <TERM> | <TERM> + <EXP> | <TERM> - <EXP>
<TERM> ::= <FACTOR> | <TERM> * <FACTOR> | <TERM> /
          <FACTOR>
<FACTOR> ::= <PRIM> | - <FACTOR>
<PRIM> ::= <identifier> | <number> | ( <EXP> )
  
```

TREE META

```

EXP = TERM ( " + " EXP : ADD { 2 } / " - " EXP : SUB { 2 } / . EMPTY );
TERM = FACTOR $ ( " * " FACTOR : MULT { 2 } / " / " FACTOR :
      DIVD { 2 } );
FACTOR = " - " FACTOR : MINUS { 1 } / PRIM;
PRIM = . ID / . NUM / ( " EXP " );
ADD { -, - } => TNEG { * 2 } SUB { * 1, * 2 : * 1 } / * 1 " + " * 2;
SUB { -, MINUS { } } => ADD { * 1, * 2 : * 1 }
      { -, - } => * 1 " - " * 2;
TNEG { MINUS { } } => . EMPTY;
MULT { -, - } => * 1 " * " * 2;
DIVD { -, - } => * 1 " / " * 2;
  
```

TREE METAにおける特徴は、Semantics 記述の簡単さと、tree のおきかえという形をとるためにObject code のOptimize まで考えることができる点にある。しかしながらTREE META には他のMETA にみられるようなアセンブラ等の出力に対する考慮が欠けているのでこの点の充実が望まれる。

参 考 文 献

1. C.S.Carr,D.A.Luther,S.Endmann

The TREE META Compiler-Compiler System:A META
Compiler System for the Univac 1108 and the General
Electric 645

TR 4-12 (RADG-TR-69-83)

2.4 META PI

このMETA PI というシステムは、オン・ラインで、インタラクティブなコンパイラ・コンパイラとして、1966年からプリンストンのRCA研究所で開発に着手したものである。このシステムは、基本的な言語構成としては、名前からも考えられるように、D.V. SchorreとそのUCLA computing facilityにおける協力者たちによって作られたMETAシリーズを用いており、したがって厳密にはsyntax-directed Symbol Processorの範ちゅうに属するものである。さらに、METAタイプのコンパイラと同じく、基本的な解析のアルゴリズムとしては、top-down, left to right, deterministicなものを採用している。

META PIが、他のMETAシリーズものと異なる最大の点は、インタラクティブであるということである。すなわち、このシステムは同研究所の開発したBTSS (Basic Time Sharing System Version II) というオペレーティングシステムの下で動くように設計されており、開発当初の目的としては、FORTAN IVに基づくインタラクティブな言語FORTRAN PI というコンパイラを作成するにあたって、先ずこのコンパイラ・コンパイラMETA PIを作り、できうるかぎりMETA PIを用いて、FORTRAN PIを書くことであつた。これに使用した機種はSPECTRA 70/45であつて同シリーズの70/46でもオペレートできるようである。なお始めの目的であつたFORTRAN PIは既に1968年以前にでき上り、そのあとのMETA PIを用いてDartmouth BASICをインプリメントしたことが報告されている。

META PIのステートメントは、他のMETAコンパイラと同じく、BNFに手を加えたextended BNFを用いている。このextended BNFはBNFとは記述方法の細かい点でも多少異なるが、大きな相違点は次のようなものである。

1. extended BNFによるSyntaxの記述の中にはSemantics (すなわちcode generation)を含めることができ、従つてソース・ステートメン

トの scan を実行中に、オブジェクト・コードを generate でき、scan の終了と同時に code generation を終了できる。

2. くりかえしやくくり出しを、表現するために \$ および () を用いる。
\$ の意味としては () の中にある任意の記号が 0 回以上続くということ
をあらわす。

たとえば BNF ステートメント

$$\langle A \rangle ::= B | \langle A \rangle C | \langle A \rangle D$$

を META PI で表現すると

$$A ::= B : \$ (: C : / : D :)$$

のようになる。ここで BNF の $::=$ は、META PI では $::=$ 、BNF の
| は META PI では /、BNF で $\langle \rangle$ で囲んで示す nonterminal
は META PI では $\langle \rangle$ をつけずに表わし、terminal を
示すためには META PI では、: (コロン) で両側を囲んで表現する。

META PI	BNF
$::=$	$::=$
/	
ABC	$\langle ABC \rangle$
: ABC :	ABC

3. META PI の創成するコンパイラがソース・ステートメントの scan
を行なう場合、特定のコマンドを用いると scan のある点から code
generation を逆戻りすることができる。但し、META PI そのも
のは deterministic な言語であるので、META PI でコンパイラを
作成する場合には、逆戻りは必要ないのである。

META PI における extended BNF と通常の BNF とを比較するために

Dartmouth BASIC の2,3のステートメントをふたつの表現で書いてみると次のようになる。但し、この例では、META PIからSemanticsに関するオペレーション部分は削除してある。

BNFを用いた例：

```
<READ statement> ::= READ <read list>
<read list> ::= <variable> | <read list>,
               <variable>
<FOR statement> ::= FOR <simple variable>
                  <expression> <OPTEX>
<OPTEX> ::= STEP <expression> | <EMPTY>
```

同じ例をMETA PI で書いた場合：

```
READST ::= READ: READLST
READLST := VAR$( : , : VAR)
FORST := FOR: SIMVAR ::= EXP : TO:
        EXP( : STEP: EXP / . EMPTY)
```

META PIのステートメントは次の3つのタイプの要素を含んでいる。

1. Syntactic な要素：これは、ユーザーのコンパイラ内のコードにコンパイルされ、ユーザーのコンパイラに対するソースの中のSyntactic な要素と比較されるものである。これはユーザーのコンパイラの syntax checker などに用いられる。
2. Semanticsの要素：ユーザーのコンパイラ内のコードにコンパイルされ object code を generate するのに用いられる。
3. META Syntactic 要素：これもユーザーのコンパイラ内のコードにコンパイルされ、逆戻りの機能を用いて、入力ソース・ステートメントをもう一度定義し直す時に生ずる矛盾を解決するのに有効なものである。

ユーザーはこれらの3つのエレメントを組合わせることによって独自のコンパイラを作成するようなMETA PIのステートメントを書くことができる。META PIステートメントの一般型は次の如くである。

LABEL:=expression

ここで、左辺は右辺の式を参照する時に用いる一意的なidentifierである。たとえば、数字をMETA PIのステートメントで定義すると

DIGIT:=:0:/:1:/:2:/:3:/:4:/:5:/:6:/:7:/:
:8:/:9:

となる。右辺にある式は、ユーザーのコンパイラの中のコードにコンパイルされるものでこの式は自分自身に対する間接的あるいは直接的な参照を含みうる回帰的なものであってもよい。

META PIは、その式をユーザーのコンパイラの中のcodeにコンパイルする際その式が後で呼び出された時に次の3つの結果のいずれかとなるようにgenerateするのである。

1. True :インプットをscanした結果がその式を満足する場合で、呼び出されたルーチンはtruth indicatorがtrueを示すようセットしinput pointerを正しくscanできたデータの後に移してリターンする。
2. False :インプットがその式を満足しないときでinput pointerは変らない。truth indicatorはfalseを示すようにセットされる。
3. Error: 例えばGO TO 20.3というステートメントのようにその式の最初のGO TOの部分は正しく識別されるのに、途中から誤りであることがわかる場合、これが生じるとエラー・ルーチンが呼ばれ、input pointerは、一部更新され、正しくscanされた最後の文字のあとに?を挿入する。

以上の3つの condition は、META PI の式によってユーザーのコンパイラ内に generate されるコードの働きについて述べたものである。次にこれらの expression を構成する要素をいくつか紹介することにする。

Syntactic な要素

：XXX…X： Xは任意のストリングを表わす。この Syntactic な要素はそのストリングと現在のインプットを比較するためにユーザーのコンパイラ中にコードを generate するものである。

例えば、前述の DIGIT ステートメントの中では、0 か、1 か 2 か、…をテストするためのコードが generate される。

A B C これは、この名前（ここでは A B C）を持つルーチンを呼び出すようなコードをユーザーのコンパイラの中に作り出すものである。このルーチンはユーザーによって META PI で既に書かれているものと仮定する。INUM:=DIGIT\$DIGIT のような使い方をしてもよい。また、このルーチン名は、現在存在している FORTRAN PI のルーチンの名前であってもよい。この使い方は META PI の中のサブルーチンを link する方法にはもうひとつピリオドを先頭につけるやり方もある。このピリオドを先頭につける場合はそのルーチンが回帰的でなく、truth indicator を返さねばならないことを意味している。ピリオドのない時は回帰的に呼び出すことも可能である。いずれの場合にも、truth indicator をセットする。

ID これは identifier に対するテストであり、ID ルーチンに対して link されるコードが generate される。先頭にピリオドがついているので自分自身を link することはできない。

. EMPTY これは truth indicator を強制的に true にセットする特殊な要素である。

- .INT FORTRAN の整数についてのテストである。
- .NUM これは、次のようなMETA PI ステートメントによって規定される数値に対するテストである。
- NUM := \$DIGIT(: : / . EMPTY) \$DIGIT(
- : E : (: + : / : - . EMPTY) DIGIT(DIGIT
- / . EMPTY) / . EMPTY)
- .LKUP これは、LKUPルーチンにlink するようなコードをgenerate するものである。このLKUPルーチンは、検出された最近のinput に対してlabel テーブルをscan する。label テーブルの項目は文番号あるいは変数名である。各項目には、このほかに適当な制御情報を含んでいる。このルーチンからは次の3つの結果の1つを返す。
- ① そのインプットがlabel テーブル中にあり、代入する。
memory location を定義する。
 - ② そのインプットが、label テーブル内にない。
 - ③ label は見つかったがそのmemory location はまだ定義されていない。これはまだ出現していない文番号を指すGO TO ステートメントなどである。
- .TYPE(: NNY :) このルーチンは、うけ渡されたインプットをlabel テーブル中で調べ、そのタイプを示すバイトが実引数NNY によって識別されるクラスにあるかをテストする。このルーチンはmixed mode error のチェックなどに用いる。

Semantics の要素

META PI ではSemantics は、META PI のステートメント中に埋めこまれ、このSemantics は、code generation に用いられるものである。Semantics は、

1. Semantic commands
2. Semantic operations

の2つの部分から成り、このうち Semantic operation は Semantic command の中に含まれ、一般型は、Semantic — command (Semantic — operations) である。

Semantic command

各 Semantic command はコンパイラによって generate される code に直接効果を及ぼす。META PI は入力ステートメント中の Semantic command に会おうとユーザーのコンパイラ中に、それらが、object program を generate するのに必要な object code を generate する。Semantic command は基本的なものとしては5つある。

.OUT (...) このコマンドは output area (そのユーザーのコンパイラによってコードが創成される一時的なエリア)の現在の内容を内部的な形に変換してユーザーの code area に入れるものである。code area は computer が実行する object code が入るエリアである。その output そのものは括弧内で指定される Semantic operation によって作られるものである。このコマンドはその中に含まれる Semantic operation の構造に従って3つの action をとりうる。

- ① 第1文字目が英字、数字のいずれでもなければ、さいごの：(コロン)の対までの文字をそのまま code area にコピーする。
- ② 第4文字目がピリオドまたはスペースならば、その output がインデックス・レジタスを用いる命令であり、第4文字目以降に記号番地をもつものと仮定する。そして、この記号番地を

labelテーブルから引いて、そこに含まれる情報からマシン・アドレスをgenerateする。

③ 上記以外の場合、その文字ストリングは機械コードであると考えてそのまま変換してcode areaへ入れる。

.LABEL(…) これはoutput areaの現在の内容を取ってきて、labelテーブルに入れるものである。このlabelが定義済みであれば、エラーとし、このlabelはcode areaのロケーション・カウンタの現在の値に関連している。

.ING(…) output areaの内容を無視(削除)するもので、これは或るSemantic operationがcodeをgenerateする時、レジスタも解放しようとするような副作用を生み出すので有用である。

.NOP(…) このコマンドは他には何もせず、Semantic operation効果を生ずるだけである。そのSemantic operationの結果はoutput areaに残っている。

.DO(…) これは特殊なコマンドでMETA PIにすぐに()内の命令を実行させるものである。

Semantic operation

Semantic operationはoutput areaにcodeをgenerateするのに用いられるもので、これらのcodeは記号的(シンボリック)な形をしているので、そのまま実行することはできない。また、これらのoperationは通常はinput pointerやtruth indicatorの変更は行わない。output area内の次の有効なlocationを指すためにはひとつのポインタを用いていて、このポインタは各Semantic operationのあとで更新する。以下semantic operationを列挙する。

:CCO...C: コロンにはさまれているストリングをそのままoutput areaの続きに入れる。syntacticな要素と同じ形であるが,seman-

tic command の内側か外側かで、いずれであるかは容易に判断できる。

- * output area の内容に現在のインプットを続けて入れる。これは普通は有効な .I D のテストと共に用いられるものである。
- S push down list の中に out area 中の現在の内容のコピーを退避して、リストを push する。
- R push-down list の一番上のものを output area に返し、リストを pop up する。
- I push-down list の一番上のエレメントを無視 (pop up) する。
- X push-down list の上の二つのエレメントを入れかえる。
- * 1 井で始まるグローバルに一意的な 4 バイトの文字ストリングを generate する。このストリングはローカルには一定で generate された code 内のロケーションに label をつけたり、これを参照したりするのに便利である。

output code 内で Semantic routine のあるものを用いればレジスタを使用することもできる。汎用レジスタと浮動小数点レジスタに対する push-down list のタイプは、実行時に保たれている。

このような Semantic operation に対するレジスタは、汎用レジスタが 6 個、浮動小数点用が 4 個である (SPECTRA 70 シリーズ)。これ以上のレジスタを要する場合には、コーディングによって自動的にレジスタの退避や restore をインプリメントするように generate でき、それには次のような Semantic routine を用いる。

- O F 現在の汎用レジスタを output する。
- O 現在の浮動小数点レジスタを output する。
- + 次の空いている汎用レジスタを output し、それを現在のものと

- する。
- + 2 次の空いている浮動小数点レジスタを output し、それを現在のものとする。
 - 2 個の汎用レジスタを output する。第 1 のものはひとつ前のレジスタで、第 2 のものは現在のレジスタである。この operation が終了すると、ひとつ前のレジスタを現在のものとする。この output は、常に数値の対であり、そのフォーマットは SPECTR-A 70 におけるレジスタ同志の operation に都合のよいようになっている。
 - 2 浮動小数点レジスタであること以外は—に同じである。

MEAT Syntactic command

以下に述べるコマンドは、ユーザーのコンパイラにおける逆もどりの機能をゆるすために設けるものである。

- .LATCH(name) これは () 内の名前のルーチンを呼び出すようなコードをユーザーのコンパイラ内に generate するものである。さらに、もし latch されたルーチンあるいはこれに続いて呼び出される任意のルーチンがエラールーチンに行きついたならば逆もどりを行なう。
- C このコマンドは Semantic operator の書けるところには、どこでも置くことができ、呼び出し側のルーチン内の .LATCH の発生をおさえるようなコードをユーザーのコンパイラ中に generate する。このコマンドはふつうは、subexpression 内の初期的なあいまいさが解決したときに用いるものである。たとえば FORTRAN の DO ステートメントで、が見つかったときや論理 IF 文で論理演算子が見つかったときである。このときは、エラー

が生じても逆もどりは起らずエラーの location を指適する。

.CLAMP このコマンドはCを書けるところならばどこでも書くことができ
先行 .LATCH で有効なものをすべて suppress してしまうこと
をコンパイラに指示するものである。

次に, META PI の例を少し挙げることにする。

[例1] FORTRAN PI では各ステートメントの区切りのセミコロンのあ
とにコメントをつけることができるが, この特徴をやめて, その代わり,
1 行にセミコロンの区切ることによって複数のステートメントを書ける
ようにするには, 次のような META PI のコマンドを書けばよい。

```
USERCC:=LABST.NOP(.CLAMP)$LABST
```

ここで LABST は FORTRAN PI のラベル付け可能なステートメン
トの定義を指すものである。

[例2] 次に BASIC の Read ステートメントを META PI で書くと
次のようになる。

```
READ:= :READ: RID$(:::RID)::;
```

このとき, META PI はこのステートメントを左から右へと scan し
次のようなコードを generate する。

- (1) 語 READ に対するテスト
- (2) RID の定義とのリンク。この定義は BASIC に関する META
PI の記述中に含まれるものである。
- (3) read の identifier のあとの, に対するテストがくりかえし行
なわれるような命令
- (4) 行の終了を示す ; に対するテスト

このMETA PI ステートメントは SyntacticなものだけとなっているがこのREADステートメントに関するSemanticsはRIDの定義の中に行なわれているのである。

前にも述べたようにMETA PIは開発当初、FORTRAN PIの作成を目的としたものであったがこのMETA PIを他のコンパイラをインプリメントするのにも用いようという試みがなされた。その結果BTSS II用のBASICがMETA PIで記述して作成された。この作成期間は1968年4月15日～同年6月6日で、

Total man hours 96時間

Console hours 33時間

Processor hours(CPU) 1時間以下

の時間を要したとのことである。

でき上ったFORTRAN PIやBTSS II用 BASICの処理効率がどの位のものであるかは明らかではないが、META PIのようなon-lineのコンパイラのインプリメンテーションにも、また個々のユーザーの必要に応じてコンパイラを修正したりするのにも有力であると思われる。

参 考 文 献

1. J.T.O'Neil Jr.

META-PI An on-line interactive compiler-compiler

FJCC 1968 PP201~218

2.5 VITAL

このシステムは、1967年初頭にMITのLincoln Laboratoryにおいて開発されたコンパイラ・コンパイラで、TX-2というコンピュータに対してインプリメントされ、すでに実用に供されているものである。システムの設計思想については、J. A. Feldman や J. E. Curry の Production Language (PL) および Formal Semantic Language (FSL) をとり入れ、このシステムの中でも Syntax を記述する言語を PL, Semantics を記述するものを FSL と呼んでいる。このように VITAL システムは、基本的な方向性としては、PL や FSL の思想を持っているがタイムシェアリングシステムのもとで動いていることがもともとの FSL との最大の相違点である。

我が国においても日立中央研究所で PL/IW という言語を用いて HITAC 5020 に対して TSS 用 TITAL システムに開発したことが 1971 年に発表されている。

VITAL という名前は (Variably Initialized Translator of Algorithmic Language) の略で、コンパイラ作成の際の細かな仕事を自動的にこなすことを目的としている。

このシステムの構成は次のようになっている。

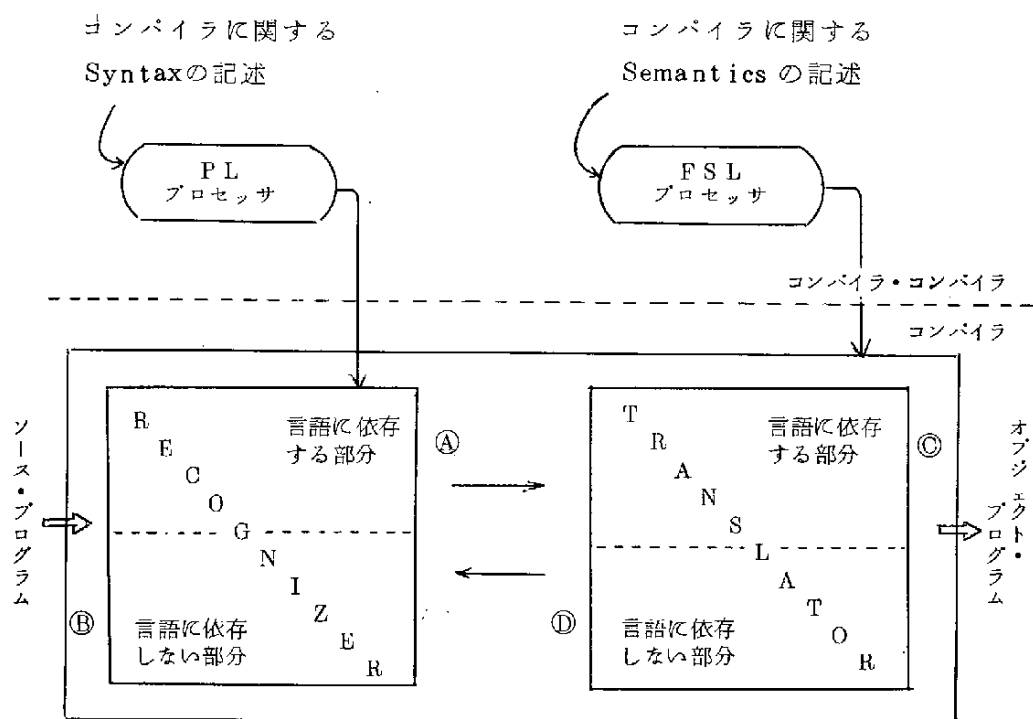


図 2-5 コンパイラ構成図

すなわち、VITAL システムでは、コンパイラ・コンパイラは PL プロセッサと FSL プロセッサの 2 つから成り、さらにその作り出すコンパイラは RECOGNIZER と TRANSLATOR の 2 つの部分に分けることができる。そして、その相互の関係は次の如くである。

PL プロセッサは、現在作成しようとしているコンパイラに関する PL (Production Language) で書かれた Syntax の記述を受け取って production table を作り出す。これは、コンパイラ言語に依存するもので (図 2-5) の ① に当り、これにコンパイラ言語の体系には依存しない production interpreter ② がつけ加えられて RECOGNIZER ができ上る。この RECOGNIZER は、コンパイラにインプットされるソース・プログラムの Syntactic な性質を認識するものである。

これに対して、FSL プロセッサは、そのコンパイラに関する FSL (Formal Semantic Language) で書かれた Semantics の記述を受け付けて、Semantic routine を作り出す。Semantic routine は、コード作成 (code generator) 用のサブルーチンの呼び出しの形であり、コンパイラ言語に依存するもので図の⑩である。これに実際にコードを generate させるためのサブルーチン群⑪が加えられて、TRANSLATOR が作られる。このサブルーチン群は、すべてのコンパイラ言語に共通なものであり、従って言語に依存しないものである。そして RECOGNIZER と TRANSLATOR は相互に働きかけ、最終的には TRANSLATOR によってアウトプットとしてのオブジェクト・プログラムのコードが作られるのである。

なお、PL プロセッサと FSL プロセッサはコンパイラとしては、前の図のように同じ構造を持っており、これらはハンド・コーディングとブート・ストラッピングによって創成された。その手順は、先ず言語に依存しない code-generator と production interpreter とがハンド・コーディングによって作られ、次に PL プロセッサの Semantics と productions がハンド・コーディングで作られ、最後に FSL プロセッサの Semantics がハンド・コーディングされ、その production を PL で書くことによって作成されたとのことである。

次に、コンパイラのオペレーションに関して少し説明し、そのあと VITAL における PL と FSL について簡単にまとめることにする。

コンパイラのオペレーション

Main スタック

VITAL の RECOGNIZER と TRANSLATOR は、Main スタックと呼ばれる 1 個の push-down スタックを用いて処理を行なう。この Main スタックは、インプット・テキストおよびそれを処理した結果としてできる解析列 (parse) を push down したり、pop up したりするものである。このスタ

ックのエントリは3つの連続したマシンワードから成っており、1ワードが syntax word, 残りの2ワードが semantic word, として用いられる。

TRANSLATORとRECOGNIZERの間の情報の受けわたしは、この Main スタックを通じてのみ行なわれる。

インプット・テキストの処理

VITAL の作り出すコンパイラでは、インプット・テキストをその中に含まれる symbol を認識しながら1度だけ scan する。これらの symbol は production によって reserved word として定められたもの（例えば FORTRAN の DO や + など）とその他の symbol（すなわち identifier）の2種類に大別できる。

reserved word が見つかるとその symbol に対応する表現形式（representation）が main スタックの syntax word に push down され、identifier であることが認められた時には、その identifier の代わりに“identifier”であることを示す項目が push down されるのである。この時点では、push down した syntax word に対応する semantic word は、関係ないものとして何も入れない。

production の一部として expected stack configuration のテーブル即ちスタックの中の予想される状態のテーブルが在り、上記の scan のあと main スタックの上の方にある2,3のエントリの syntax word をこのテーブルと比較する。一致するものが見つかった場合には、この一致した configuration に対応しているいくつかのルーチンが実行されるのである。

このルーチンは

- (1) main スタックの一部の修正
- (2) セマンティック・ルーチンへの jump
- (3) エラーメッセージの記録

- (4) さらに一つ以上の symbol の scan
- (5) main スタックと比較すべき次の configuration の指定のうちのいずれかを行なうことができる。

この scan と照合という過程は、入力プログラムのすべての symbol を scan してしまいか、あるいは production か Semantics 内の終了コマンドが現われるかのどちらかが起るまで続けられるのである。

PL (Production Language)

Production Language は、VITAL コンパイラ・コンパイラ システムにおいて、コンパイラの Syntax を指定するための production を記述するのに用いる言語である。

そしてこの production は、(1)インプット・テキストを1度だけ scan して main スタックの継続的な状態を作り出し、(2)スタックの中にある syntactic な要素に対応する Semantic routine にコントロールを移す働きがある。

PL プログラムの構造は、初めに宣言を書き、そのあとに production を続けて書いていけばよく、プログラムの終りには END という記号をつける。

PL における宣言文の中から二、三のものを挙げてみることにする。

RES <argument>

この argument としては VITAL の symbol をスペースで区切って書くことができる。この宣言は argument として与えた symbol が今作ろうとしているコンパイラのソース言語において reserved word として用いられることを示すものである。

INT <argument>

この argument としては VITAL の identifier をスペースで区切って書くことができる。そしてこれは internal symbol を定義するもので、ここで指定した internal symbol は今作り出そうとしているコンパイラに対するソー

ス言語の一部でも、PLの一部でもなく、そのコンパイラがスタックの Syntactic な部分に用いる名称である。

<argument 1> = <argument 2 >

argument 1 としては、任意の VITAL の identifier, argument 2 としては、reserved word を並べたものを書くことができる。これは、記述するのに便利のようにクラス名を定義するもので、クラス名を含む production を1つだけ書くことによって、対応する reserved word を含む production を汎用書き並べる代わりにすることができる。

次に production については例を用いて説明することにする。

B E 3	I F E T H E N	→	I C L	EXEC 1 1 0 S 0
⏟	⏟			⏟
		↑	↑	
(a)	(b)	(c)	(d)	(e)

- (1) B E 3 はラベルである。任意の production に英数字から成るラベルをつけることができ、ラベルの終りは || で示す。
- (2) 次に MAIN スタックの内容を調べる。上の例では、スタックの上から3つ分の内容が、上から順に THEN, E, IF であれば、スタックからこの3つのエントリを pop up し、その代りに、矢印→の右側に出現している I C L をスタック上に push down する。
矢印の左側に S G という symbol が出現した場合には、スタックの当該エントリの内容が何であってもこの S G と一致したとみなすことができる。また、I も特別な意味を持っており、identifier が scan された場合にスタック上に push down される。
- (3) スタックの内容がある production の矢印の左側と一致している場合は、変換を行ない、次の部分 — ここでは、EXEC 110 と S 0 を実行する。この部分は actions とされるもので、EXEC 110

は Semantic routine 110を呼び出すことを示し、 $\Rightarrow$ S 0 は S 0 という
ルの付いた production にコントロールを移すことを示している。action
にはこのほかに次のようなものがある。

UNSTK m : スタックの上からm個のエントリを取り除き、これを
storage に退避する。

STK n : UNSTK によって取り出された上からn番目のエントリ
をスタック上に push down する。

STAK<argument> : argument をスタック上に push down する。

SCAN : インプット・シンボルを1個 scan する。

SCAN n : インプット・シンボルをn個 scan する。

ERROR n : 出力バッファにエラーメッセージnをストアする。

HALT n : 現在実行中の production の実行を中止し、出力バッファ
にnをストアし、システムにコントロールを返す。

NEXT<argument> : $\Rightarrow$ <argument>に同じ

TEST<argument> : SIGNALが真の場合に限り、 $\Rightarrow$ <argument>
と同じ働きをする。

RETURN : 最後にCALLを行なった production へ、コントロール
を移す。

(4) スタックの内容が、矢印の左側の部分と一致しても、内容を変化させる必要のない場合には、矢印および変換後のスタックの内容に当る部分を省略すればよい。

(5) スタックの内容がある production の左側と一致しない時には、物理的に次の production にコントロールが移され、再び新しい production の左側の部分との比較が行なわれる。このように構文解析は、直接的に行なわれるから逆もどりを避けるために、production を書く際にユーザーが注意しなければならない。例えば

$P \rightarrow T$, $T * P \rightarrow T$ がある場合には、 $T * P \rightarrow P$ を先に置くようにしなけれ

ばならない。

ここで production に関してもう一つ例をあげて考えてみる。PL プログラムを書く場合、実際には、そのプログラミング言語の文法を直接 PL で記述すればよいわけであるが、ここでは、それを先ず Backus Normal Form で表わして、その次に PL で書いてみることにする。

BNF による式の表現：

$$\langle \text{式} \rangle ::= \langle \text{項} \rangle \mid - \langle \text{項} \rangle \mid \langle \text{式} \rangle - \langle \text{項} \rangle$$
$$\langle \text{項} \rangle ::= \langle 1 \text{ 次子} \rangle \mid \langle \text{項} \rangle * \langle 1 \text{ 次子} \rangle$$
$$\langle 1 \text{ 次子} \rangle ::= \langle \text{identifier} \rangle \mid (\langle \text{式} \rangle)$$

同じものを PL で記述したもの：

$$\left. \begin{array}{ll} E_1 \parallel & T \rightarrow E \\ & - T \rightarrow E \\ & E - T \rightarrow E \\ T_1 \parallel & P \rightarrow T \\ & T * P \rightarrow T \\ P_1 \parallel & I \rightarrow P \\ & (E) \rightarrow P \end{array} \right\} \quad (1)$$

ここで P, T, E は実際のプログラム上では internal symbol として定義すべきものであり; *, -,), (は reserved word として宣言しているはずである。

上記(1)の PL プログラムは、PL 語の記述方法については正しいが、実際的なプログラムとして用いるためにはもう少し工夫を加えなければならない。第 1 に単純な構成のものの方から記述すること、および第 2 に例えば E₁ というラベルのところでは T, -T, E-T と並んでいて最後の symbol は全部 T である。このような場合長いストリングのものからサーチしないと誤った解析を行なうことになるので注意することである。

これを考慮して(イ)を書き換えると次のようになる。

$P_1 \parallel$	$I \rightarrow P$	}	(ロ)
	$(E) \rightarrow P$		
$T_1 \parallel$	$T * P \rightarrow T$		
	$P \rightarrow T$		
$E_1 \parallel$	$E - T \rightarrow E$		
	$-T \rightarrow E$		
	$T \rightarrow E$		

さて、productionの実行順序は、スタックの内容と一致するものが見つかった時には、無条件に、あるいはその次に在る symbol を scanすることによって、それに従ってジャンプするわけで、シーケンシャルになってはいない。また、一致するものが見つかった場合には、適当な FSL ルーチン呼び出しで、main スタックの Semantic な要素を更新したり、ストレージを修正したり、code を作り出したりするはずである。さらに production は、インプットステートメント中のシンタックス・エラーも見つけねばならない。このことを(ロ)の PL プログラムにつけ加えると次のようになる。

$B_1 \parallel$	—		SCAN	☞	P_1
$P_1 \parallel$	(		SCAN	☞	B_1
	$I \rightarrow P$	EXEC 1		☞	T_1
$P_2 \parallel$	$(E) \rightarrow P$	EXEC 2		☞	T_1
	SG	ERROR 1			
$T_1 \parallel$	$T * P \rightarrow T$	EXEC 3	SCAN	☞	TBR
	$P \rightarrow T$		SCAN	☞	TBR
TBR	$T *$		SCAN	☞	P_1
$E_1 \parallel$	$E - T$	SG → E SG	EXEC 4	☞	EBR
	$-T$	SG → E SG	EXEC 5	☞	EBR
	T	SG → E SG		☞	EBR

EBR #	E)		$\Rightarrow P_2$
	E -	SCAN	$\Rightarrow P_1$
SG	ERROR 2		

($-A * B - C$) という式を考え、これが上の production でどのように実行されるかを追いかけてみることにする。

表 2-6 VITAL の例

ラベル	Production	スタック ([initial state])
P_1	SCAN $\Rightarrow B_1$	(-
B_1	- SCAN $\Rightarrow P_1$	(- I
$P_1 + 1$	I P EXEC 1 $\Rightarrow T_1$	(- P
$T_1 + 1$	P T SCAN $\Rightarrow TBR$	(- T (- T *
TBR	T* SCAN $\Rightarrow P_1$	(- T * I
$P_1 + 1$	I P	(- T * P

	EXEC 1	
	T_1	
T_1	$T * P$	
	T	(-T
	EXEC 3	
	SCAN	(-T-
	TBR	
$E_1 + 1$	-TSG	
	ESG	(E-
	EXEC 5	
	EBR	
$EBR + 1$	E-	
	SCAN	(E-I
	P_1	
$P_1 + 1$	I	
	P	(E-P
	EXEC 1	
	T_1	
$T_1 + 1$	P	
	T	(E-T
	SCAN	(E-T)
	TBR	
E_1	E-TSG	
	ESG	(E)
	EXEC 4	
	EBR	
EBR	E)	

	P_2	
P_2	(E)	
	P	P
	EXEC 2	
	T_1	
$T_1 + 1$	P	
	T	T
	SCAN	この例の終り

表 2-6 VITAL の例

この時のスタックの状態はこの式の解析列を形成していることになる。次にこの状態を tree であらわして、PL についての説明を終ることにする。この tree では最初に scan された文字が一番下にあり、横の線はスタックの状態を示すものである。

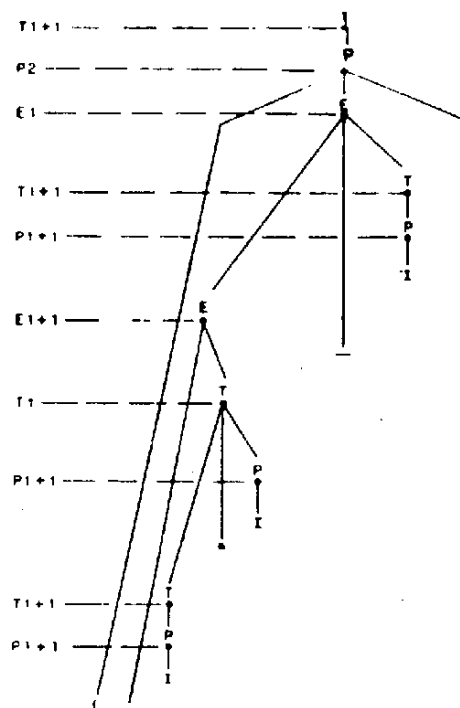


図 2-7 Syntax tree

F S L (Formal Semantic Language)

F S Lは、これから作ろうとしているコンパイラに対するソース言語の Semantics を記述するのに用いる言語である。V I T A L の特徴の一つは、Semantics を処理する T R A N S L A T O R が、言語に依存する部分と依存しない部分とに分かれていることであり、このため言語に依存しない code generator のサブルーチンさえ備わっていれば、F S L を用いて記述する Semantic routine は、特定のコンピュータを想定しなくてよい点である。したがって、F S L 自体はセマンティック・メタ言語であるといえることができる。

F S L プログラムの基本的な単位は、Semantic routine と呼ばれる番号を付けた1個のステートメント、あるいは一連のステートメントであって、production の action の部分に、EXEC n (n は Semantic routine の番号)と書くことによって呼び出され、これによりコードをジェネレートさせたり、T R A N S L A T O R の状態を変化させたりするものである。

F S L は、T R A N S L A T O R の働きを記述するものであるので、コンパイルするプログラムの状態についての情報として cell やテーブル、スタックなどを用いることができる。これらの cell やテーブルなどは、F S L の基本的な構成要素で、これらは1次子(primary)(論理的要素であれば論理1次子)と呼ばれるものであり、1次子の中には、ユーザーやシステムによって値の与えられる変数も含まれているのである。以下に1次子の中の代表的なものについて簡単に説明を述べることにする。

まず、cell はユーザの管理できる最も簡単な変数で1個のマシン・ワードに対応し、コンパイラ中の単一の storage location を指すものである。そして cell は、T R A N S L A T O R のワーク・エリアやカウンタとして用いることができ、通常は cell 宣言(cell declaration)によって定義しなければならない。

cell には、system cell と呼ばれるシステムで定義している cell があ

る。その主なものは次のとおりである。

CODELOC……コンパイルされたコードの入り得る次の空き番地を指し
分岐命令やサブルーチンの呼び出し等が行なわれる時に用
いられる。

STORLOC… ソース・プログラムのデータの入り得る次の空き番地を指
している。

PERSLOC …コンパイル・コードにおいて、システムの用い得る最初の
ワーク・エリアを指している。

ACCUM………実行時に用いるアキュムレータの使用状態を示す。

MAIN………Main スタックの一番上のエントリを指している。

前にも述べたように、system cell 以外の cell を用いる場合は、それを使
用する以前に、次のような cell 宣言で定義しなければならない。

CELL <argument>

ここで argument としては、cell identifier をカンマで区切って並べる。
cell の名前には任意の VITAL の identifier を用いることができる。

FSL では、コンパイル時に操作するテーブルを次のような形で宣言した上
で用いることができる。

TABLE <argument 1> <argument 2> <argument 3>

ここで argument 1 は、そのテーブル名を示すもので任意の VITAL の
identifier を使うことができ、argument 2 はそのテーブルに含まれる行
(row) の数を示す整数、argument 3 はテーブルの 1 行に含まれる列(column)
の名前として VITAL の identifier をスペースで区切って並べたものであ
る。

例えば TABLE SYMBOL 400 ID SEMANT EXTRA は、
ID, SEMANT, EXTRA という 3 つの列を持ち、サイズが 400 行の
SYMBOL という名のテーブルである。

テーブルの中に新しいエントリを設ける場合には、ENTER ステートメン

トを用いて行ごとに作成していく、たとえば前記の SYMBOL というテーブルの第1カラムには NAME、第2カラムに REAL、第3カラムに XX を入れるのであれば、1 # ENTER {SYMBOL, NAME, REAL, XX} と書く。

F S L で記述する時、テーブルの参照にはテーブル・オペランドという1次子を用いる。

これは <argument 1> {<argument 2>, <argument 3>, <argument 4>} という一般形を持つもので、

argument 1 は テーブル名

argument 2 は 1次子もしくは0

argument 3 は 列の名(column name)

argument 4 は 0 または 1 (モードとして使用)

である。テーブル・オペランドの値は、次のようにして見つけ出されるテーブルのエントリである。argument 2 が1次子ならば argument 1 のテーブルの第1列の中にその1次子と等しいものがあるかどうかサーチを行なって、見つければその行の中で argument 3 で指定された列の値をそのテーブル・オペランドの値とする。argument 2 が0 ならば最も最近テーブル・サーチが行なわれた行の指定列の値をセットする。テーブル・サーチしても見つからない場合には論理1次子 SIGNAL を FALSE にセットする。なお、argument 4 はサーチの方向を示すもので0 ならば forward search、1 ならば backward search を行なう。

F S L で使用可能なスタックは push - down スタックで、TRANSLATOR が用いるものは STAK; 実行時に用いるものは RSTAK 宣言を用いて定義する。これらのスタックを参照する際には、そのスタックの一番上のエントリを見ることになる。スタックの大きさは、特に指定しないかぎり (100) である。

以上述べてきた cell、テーブル・オペランド、スタックはいずれも F S L

ステートメントで用いることのできる1次子(オペランド)であるが、1次子にはこのほかに production operandやFLAD, その他のものがある。

production operandとは、ある production が実行される前、もしくは、実行された後のMainスタックのセマンティック・ワードの内容を指すものでL1, LL1, L2, LL2, ..., L7, LL7 で実行前の上から7個のエントリ, R1, RR1, ..., R4, RR4 で実行後の上から4個のエントリを指すことができる。例えば、

L1 → R2; LL1 → RR2;

L2 → R1; LL2 → RR1;

は、Main スタックの上の2個のエントリのセマンティック・ワードを互いに入れ替えるものである。

コンパイラを作成する際に、未だ出現していないステートメントへの分岐命令の処理は厄介なものであるが、FSLでは FLAD1, ..., FLAD4を用いてこれを解決している。FLADはそれぞれスタックであって、この種の分岐命令の度にPUSH FLADnを行ない、ASSIGN ステートメントが実行されるとその時のCODELOC の値がセットされる。

同じような動きを持つものにCHAIN があり、FLAD が前向きの chain を作るのに対して、CHAIN は前後両方向への chain を作ることが可能である。

さらに1次子としては定数や、他のオペランドの内容やアドレスを参照する時に用いるCONTENTS やLOCATION, 絶対値をとり抜かうABSOLUTE VALUE, Main スタックの内のエントリの属性を調べるTYPE, ビット単位の処理に用いるTAG, 算術式をカッコで囲んだものなどがある。

1次子とならんでオペランドとなるものに論理演算用の論理1次子(Boolean primary)がある。この中に含まれるものにはTRUE とFALSE からなる論理定数(Boolean constant)をはじめ、指定された1次子の指定されたビットにタグがつけられているかをテストするTEST, 指定された1次子が定数であるか否かを調べるCONSTANT TEST, 1次子どうしの間の大

小関係をみる RELATION, および特別な BOOLEAN CELL と呼ばれる OK と SIGNAL, それに論理式をカッコで囲んだものがある。テストを行なうような論理 1 次子は, 与えられた条件が満たされれば TRUE であり, それ以外の場合には FALSE となる。cell OK および SIGNAL は, それぞれ, ユーザーあるいはシステムによって TRUE か FALSE の値を与えられるものである。もし cell OK が FALSE にセットされると, インプットはひきつづいて scan されるがこれ以降のコードはジェネレートされない。

cell OK はこれ以上オブジェクト・コードを作り出したとしてもでき上ったプログラムが動かせないとか, 致命的なエラーがあるものになってしまうような重症のエラーが生じた場合に用いるように設計されたもので, したがって一度 FALSE になるとこれを reset することはできない。

このシステムはあるルーチンが終了した時, それが正常に終わったかどうかに応じて, SIGNAL に TRUE あるいは FALSE の値を与える。例えば table search を行なった時, 該当するエントリがあれば SIGNAL は TRUE, なければ FALSE となる。

次に FSL で使用可能な式には, 算術式と論理式の 2 種があり, 算術式は 1 次子および単項演算子 +, -, 二項演算子 +, -, *, / およびべき乗の演算子から成るものである。論理式は論理 1 次子および ~, ^, V (inclusive or), ⊕ (exclusive or) という演算子から成る。演算の順序は一般の優先順位に従がい, 括弧があれば, 括弧内を優先する。

また, 1 次子および論理 1 次子はそれ自体, それぞれ算術式, 論理式とみなすことになっている。

FSL における宣言には, 前述の CELL, TABLE, STAK, RSTAK のほかに INDEX および DATA に関するものがある。INDEX 宣言というのは, INDEX identifier で定義するもので, これは通常の方法で, 添字あるいは算術式としてコンパイラが用いるものである。このシステムでは, 8 個のインデックスが使用可能である。

DATA 宣言は、コンパイル時にビット単位に処理を行なうオペランドにタグをつけたり、テストを行なったりすることができるように DATA identifier を定義するものである。VITAL では BOOL, INTGR, FRACT, REAL という 4 つの data identifier をシステムが定義しているので、これらは宣言しないで用いることができる。

FSL 中で用いることのできるステートメントとしては、先ず非条件ステートメント (unconditional statement) として、スタック類をとり扱かう PUSH ステートメントと POP ステートメントがあり、これらはスタックにエントリを入れたり、とり出したりするのに使われる。テーブル類のとり扱かうには、前述の ENTER を使用する。

左辺の内容を、右辺にストアする場合には STORE ステートメントを用いこれは

$$\langle \text{argument 1} \rangle \rightarrow \langle \text{argument 2} \rangle$$

(argument 1 は 1 次子, argument 2 は 1 次子あるいは Boolean cell) のように矢印で表わす。

ある 1 次子に算術式の値を加えるのには TALLY ステートメントを用いてもよい。これは

$$\text{TALLY } \langle \text{argument 1} \rangle$$

あるいは

$$\text{TALLY } \{ \langle \text{argument 1} \rangle, \langle \text{argument 2} \rangle \}$$

の形で argument 1 の 1 次子に argument 2 の算術式の値を加えることを意味し、argument 2 を指定しないと 1 を加えるものである。

このほか、前述の FLAD や CHAIN に値を与えるための ASSIGN ステートメント、ストレージをとりあつかう LOAD, SUBLOAD ステートメント、サブルーチンを定義するステートメントの 、これを呼び出すための CALL ステートメント、無条件に分岐することを示す JUMP, その他 RETURN, STOP などのステートメントがある。

次に条件ステートメント (conditional statement) としては、IF
<argument 1> THEN <argument 2> Ⅱ, IF NOT <argument 1>
THEN <argument 2> Ⅱ, IF <argument 1> THEN <argument 2>
ELSE <argument 3> Ⅱの3つの形の IF ステートメントがあり、いずれも
argument 1 の論理式の値に従って argument 2 のステートメント群
(セミコロンで区切って複数のステートメントを書くことができる) を実行
するか否かを決定するステートメントである。

これらの IF ステートメントは入れ子 (nesting) を許している。

FSLプログラムの構造は、PLの場合と同様、初めに宣言を書き、そのあ
とに Semantic routine をいくつも続けて書いていく形を採っている。そし
て Semantic routine は次のような項目から成り立っているものである。

- ① #のあとに整数をつけたもの、(この整数は production の EXEC
n の n に当るものである。)
- ② 空白 (スペース)
- ③ ステートメント (セミコロンで区切った条件あるいは非条件ステートメ
ント)
- ④ #

このほか、プログラムの最初と最後は BEGIN と END でなければならない。
またコメントは2個のアスタリスク**で始まり、キャリッジ・リターン
で終わるようにすれば、プログラムのどこに置いても構わない。

FSLにおいては、コード・ブラケットと呼ばれる特別な括弧⌈, ⌋を用い
て、FSLで書かれたオペレーションが、コンパイル時に行なわれるものであ
るか、実行時に行なわれるものであるかを区別している。コード・ブラケット
内に書かれたオペレーションは実行時のもので、従ってコンパイル時にオブジェ
クト・コードが作り出される。一方、コード・ブラケットに囲まれていないオ

ペレーションは、コンパイル時にTRANSLATORが実行するものである。例えば

$$L4 + L2 \rightarrow R2$$

の場合、コンパイル時にL4とL2の内容を加えてR2にストアするのに対して

$$\boxed{L4 + L2 \rightarrow R2} \boxed{}$$

の場合には、コンパイル時にはL4にL2を加えてR2にストアする演算のオブジェクト・コードをジェネレートし、実際にこの演算が行なわれるのはオブジェクト・プログラムの実行時である。

また、コンパイルの際に必要なデータの属性は、セマンティック・ワードにセットされるもので、このセマンティック・ワードの第1ワード目は、データの型やインデックスの番号などを含んでいて、システム側で定めたフォーマットを持っているが、第2ワード目は、このほかの必要な属性などを示すためにユーザーが自由に用いることができる。これらのとり扱いかいについては、前述の1次子やステートメントの働きで比較的自由に行なうことができるようになっているのも、FSLの特徴となっている。

このVITALシステムは、コンパイラ・コンパイラとしてかなり力のあるものである。このシステムは既に実用に供されていてこれによって作られたコンパイラもいくつもあるようである。

参 照 文 献

1. L. F. Mondschein

VITAL Compiler-Compiler System reference manual

Lincoln Laboratory Technical Note 1967-12

(図2-5, 6, 7)は上記文献より引用させていた。Sいた。

2. J. A. Feldman

A Formal Semantics for Computer Languages
and its Application for Computer Languages

C. ACM Vol. 9 No. 1 1966 PP 3~9

3. 二村良彦他

PL/IWによるTSS用コンパイラ・コンパイラの作成

第12回プログラミング—シンポジウム報告集 1971 PP 43~53

2.6 CC (the Compiler-Compiler)

このシステムは、マンチェスター大学で研究開発されてきたもので、Brooker や Morris, Rohl などの手によるものである。この研究は、かなり古くから行なわれてきたものと思われ、1960年には、Brookerが、その基本的なアイディアを既に発表している。このシステムは、きわだって効率がよいとは思われないが、いくつものコンピュータ・システムでインプリメントされており、CCで書かれた ALGOL, FORTRAN II, そのほか Atlas autocode, Mercury Autocode などが作り出されていることも報告されている。

ひとつのプロセッサというものを考える場合、そのプロセッサの使用目的によって、とり扱かうデータなどの種類が異なってくる。例えば、ALGOL や FORTRAN であれば、浮動小数点数の配列などを多く扱っていると考えられるわけであり、これに対して、このCCは tree 構造のデータを操作できるよう設計されていると言える。CCで記述するプログラムの構造はFORTRAN や ALGOL のように、メインルーチンとサブルーチンから成り、そのサブルーチンのあるものは、作り出そうとしているコンパイラ言語のいろいろなソース・ステートメントを処理するものであり、またあるものは、コンパイラの処理を容易にするために導入する補助ステートメント (auxiliary statement) に対するものである。メインルーチンは、CCシステムによって作られていて、ユーザーの作り出すコンパイラ言語には関係なく、次のような定まった働きをするものである。すなわち、1度に1行ずつソース・ステートメント (コンパイラに対するソース・ステートメント) を読み、これをユーザーの書いた phrase と format の定義を用いて、その言語の構文規則に従って解析する。その解析方法は、top-down であり、ソース・プログラムの行がそれに対して対応しているような format を識別し、セマンティック・フェーズで用いられる完全な syntax tree を得る。そして、そのあと、メインルーチンは、この format に関連する routine を呼び出して、認識されるインストラクションを変換する。このあと、次のソース・ステートメントを読むように

コントロールは、メインルーチンに戻されるのである。

シンタックス・フェーズへのインプットの形は、拡張 BNF (extended Backus Normal Form) を用いている。これは、普通の BNF に、くりかえしを示す * とオプションを表わす ? を加えたものである。例えば

$$[A1] = [B*][I], \quad [B] = a, \quad [I] = i$$

であれば、A1 は b を 1 個以上続けて書いたあとに i をつけたもの (bi, bbbi, …… etc.) であり、

$$[A2] = [B?][I]$$

であれば、A2 は bi または i である。

CC で採用している BNF の記述方法は、正統的な nonterminal をカッコで囲んで書き、terminal symbol はそのまま書くやり方をとっているものと、nonterminal をそのまま書き、terminal を引用符で囲むものなど、CC をインプリメントしているコンピュータ・システムによって多少異なっているが、いずれの場合も ALGOL などですでに示すメタ言語の or はカンマ(,) で置き換え、 $\therefore$ は = で表わしている。

シンタックス・フェーズへ、この BNF で書いた Syntax をインプットする場合、その並べ方の順序が問題になるのは FSL などの場合と同様である。すなわち、

$$[X] = [Y], \quad [Y][\pm][Z]$$

のような場合、上の書き方では $[Y][\pm][Z]$ に対応するソース・ステートメントも $[Y]$ であると認識してしまっていて、正しく解析できないので

$$[X] = [Y][\pm][Z], [Y]$$

のように書かねばならない。

さらに、 $[FD] = [BD][OD][OD][OD], [OD][OD][OD]$

$$([BD] = 0, 1 \quad \text{および} \quad [OD] = 0, 1, 2, 3, 4, 5, 6, 7)$$

のような場合、これは 2 進数 1 桁の後に 8 進数を 3 桁付けたものであるが、これを

$$[FD] = [BD?][OD][OD][OD]$$

と書くとまちがいである。その理由は、例えば(157)₈のような0または1で始まる3桁の8進数は後の定義によると、FDに入らないことになってしまうからである。

ここで FORTRAN コンパイラからカッコを含まない算術式を例にとって、コンパイラ・コンパイラCCについて考えることにする。

$$\text{FORMAT}[SS] = [\text{VARIABLE}] = [\text{EXPR}] \dots\dots\dots(1)$$

ここでメタ変数は、次のように定義する。

$$\text{PHRASE}[\text{EXPR}] = [\text{PRECEDING}\pm][\text{TERM}][\text{TERMS}]$$

$$\text{PHRASE}[\text{PRECEDING } \pm] = [\pm], \text{NIL}$$

$$\text{PHRASE}[\pm] = +, -$$

$$\text{PHRASE}[\text{TERMS}] = [\pm][\text{TERM}][\text{TERMS}], \text{NIL}$$

$$\text{PHRASE}[\text{TERM}] = [\text{FACTOR}][\text{FACTORS}]$$

$$\text{PHRASE}[\text{FACTORS}] = *[\text{FACTOR}][\text{FACTORS}],$$

$$/[\text{FACTOR}][\text{FACTORS}], \text{NIL}$$

図2-8 CCの例

(1)は format definition, (図2-8)は phrase definition と呼ばれ、SSは、source statement を表わす format class を示している。(1)は簡単に言えば、 $[\text{VARIABLE}] = [\text{EXPR}]$ という Syntax がこの言語の source statement として含まれることを表わしているのである。 phrase definition は format definition, あるいは、他の phrase definition を補うためのものである。

さらに補助ステートメント (auxiliary statement) を導入する。

$$\text{FORMAT}[AS] = \text{LOAD}[\text{PRECEDING } \pm][\text{TERM}]$$

```

FORMAT[AS]=ACC=ACC[+](TERM)
FORMAT[AS]=DIV BY [FACTOR]
FORMAT[AS]=MULT BY [FACTOR]
FORMAT[AS]→LOAD(PRECEDING ±)(FACTOR)
FORMAT[AS]=DUMP ACC IN [VARIABLE]

```

図 2-9 CC の例

実際に、コンパイラに対するソース・ステートメントを読んで、それに対応する tree を format definition を用いて作り出すと、次にコンパイラは、その tree に対して、オブジェクト・コードをジェネレートするとか、宣言に対するリストを作るとかなどの処理を行なうことになる。その処理ルーチンを CC では format routine と呼んでいる。例えば (1), および (図 2-9) の第 1 番目の format に対する format routine は次のようになる。

```

ROUTINE(SS)≡[VARIABLE]=[EXPR]
    LET(EXPR)≡[PRECEDING ±](TERM)(TERMS)
    LOAD(PRECEDING ±)(TERM)
2) → 1 UNLESS(TERMS)≡[±](TERM)(TERMS)
    ACC=ACC[±](TERM)
    → 2
1) DUMP ACC IN [VARIABLE]
    END

ROUTINE(AS)≡LOAD(PRECEDING ±)(TERM)
    LET(TERM)≡[FACTOR](FACTORS)
    LOAD(PRECEDING ±)(FACTOR)
2) → 1 UNLESS (FACTOR)≡/[FACTOR](FACTORS)

```

DIV BY [FACTOR]

→ 2

1) → 3 UNLESS [FACTOR] ≡ *[FACTOR][FACTORS]
MULT BY [FACTOR]

→ 2

3) END

図 2-10 FORMAT ルーチン

format routine の第 1 行目は heading と呼ばれ、"≡" の右側は format definition の符号の右側と同じ形になっている。これによって format definition に対応した format routine が呼び出されることになり、また heading の右側の部分は、format definition から format routine へのパラメタを渡すのに用いられている。ユーザーは後述する BS クラスのシステムに組み込まれている routine や、すでに定義済の format routine を用いながら新しい format routine を定義することができる。したがってこのシステムは 拡張可能であると言える。次に format や phrase definition などに関連して、もう少し説明を補うことにする。

CC に対するプログラムの入力システムは 2 つのフェーズから成っており、第 1 のフェーズは句構造文法の定義を受け取る場所であり、第 2 のフェーズはその言語で書かれたソース・プログラムを変換するものである。

CC においては、基本的には 4 つの format class が用いられる。それらは次のとおりである。

[SS] source statement の format の class

[MP] master phrase の class (即ち、PHRASE や FORMAT など)

[BS] format routine 内で用いられるシステムに組み込まれた
ステートメントの class

[AS] ユーザーが、その format routine 内での使用のために導入

できる補助ステートメントの class

あとの3つは phrase, format および format routine すなわち, その言語の定義 (definition), Syntax, および Semantics を読み込む CC の第1フェーズでのみ用いられる。その言語に固有の format や phrase definition はコンパイラが定義され, 第2フェーズで operational になったら削除することができる。これを削除しない形で残すと, そのコンパイラはインクリメンタル・コンパイラとなり, ステートメントの定義を付け加えることができるようになるのである。

前にも述べたように, この CC は既にいろいろなコンピュータでインプリメントされており, 特に Atlas に対して作られた CC に関しては, 多くの文献が著わされている。その中には, 実際に CC を用いて作成した Atlas Autocode コンパイラと同じものを手書きで作ったものとを比較したものもあるが, その結果は, CC を用いた方が, スペースについても, 時間に関しても2倍程かかっている。これは, かなり CC を使い慣れた人が使った場合で, 普通は4~5倍の時間を要するということも述べられている。コンパイル時間の遅いことの主な理由は, 構文解析にも幾分原因があるが, それよりも算術式のとり扱いかいによるところが大きい。これは $X = 1$ のような簡単な代入ステートメントでも, 上記の算術式 $[VARIABLE] = [EXPR]$ として解析すると, 多くの不要な枝を持った tree になってしまうためである。こういう場合には演算子順位などの方法による方がずっと効率がよい結果を生じるようになるはずである。

けれども CC の製作に当たっている人々は, CC を用いてコンパイラを作る方が, 手書きでコンパイラを作成するのよりは速いと書いている。そして, 特に Syntax が複雑な場合には, 有利であるとのことである。

CC については, 巻末の TWS の翻訳の中にも説明があるので詳しくは, そちらを参照されたい。

参 照 文 献

1. R.A.Brooker, D.Morris
A general translation program for phrase structure language
J.ACM Vol.9 No.1 1962 PP 1 ~ 10
(図2-8, 9, 10)は上記文献より引用させていた。いた。
2. R.A.Brooker, D. Morris
Tree and Routines
Comput.J. Vol 5 1962 PP 33 ~ 47
3. R.A.Brooker, D.Morris, et. al.
The Compiler - Compiler
Annual Review in automatic Programming Vol 3 1963
PP 229 ~ 275
4. R.A.Brooker, D. Morris, et. al.
Compiler - Compiler Facilities in Atlas Autocode
Comput. J. Vol 9 1967 PP 350 ~ 352
5. R.A.Brooker, D. Morris, et, al.
Experience with the Compiler - Compiler
Comput. J. Vol 9 1967 PP 345 ~ 349

2.7 CABAL

CABAL システムは、Feldmann や Curry の Production Language (PL) と Formal Semantics Language (FSL) の思想を取り入れたコンパイラ・コンパイラで Carnegie - Mellon 大学で3年がかりで開発されたことが、1968年に発表されている。このため、同じく PL, FSL の流れをくむ MIT の VITAL システムと酷似している部分もあるが、CABAL は、全体として ALGOL-like な言語であることに付随して、言語の表現はかなり、異なっている。VITAL と同じように CABAL もインタラクティブなシステムである。

CABAL プログラムの構成

前述のように、CABAL は PL, FSL の思想を取り入れてはいるが、CABAL プログラムの記述の際には、reduction 関係のステートメントと、code generation および一般の CABAL ステートメントを混ぜて書くことができるようになっている。すなわち、CABAL におけるプログラム単位は、通常の ALGOL procedure に類するものと CO-routine と呼ばれるものから成り、これらのプログラム単位内では、ALGOL と同じく BEGIN, END に囲まれたブロック構造や identifier のスコープを持ち、宣言文およびステートメントから成っている。

CO-routine というのは、サブルーチンであるがタスクの概念を備えたもので、前にコントロールが中断された点から再び処理を続行できる機能を持っている。CO-routine では、procedure の先頭に普通の宣言文に先立って COMMON ステートメントを置いて、これによって処理の再開される場所を示すこと以外は、通常の ALGOL procedure と同じ形を取っている。CO-routine が本来の CO-routine として呼び出された時には、呼び出し側とその CO-routine とには主従関係は存在しないが、これを普通の procedure として用いることも可能である。

CABAL は VITAL のように PL プログラムと FSL プログラムを、記

述の際に区分する必要はないが、実行の上では、reduction ステートメントを順に実行し、その中から他のステートメント、またはブロックを呼び出していくようになっている。次に先ず、CABAL における宣言文、reduction ステートメント、一般のステートメント、Code - Generation に関するステートメントについて述べる。

CABAL の宣言文

CABAL における宣言文は、どちらかという PL / I 的なもので、この言語で使用可能な 5 種類のデータ構造 scalar, array, list, stack, queue を 1 つの形で宣言でき、その機能の違いは宣言につけ加えるパラメタの有無によって決められる。

宣言文の Syntax は次のとおりである。

```

<declaration> ::= DECLARE <declaration list>
<declaration list> ::= <name list>
                                (<structure>)
| <declaration list>, <name list>
                                (<structure>)
<structure> ::= <nature> <element linkage>
<nature> ::= NUMBER, <bound>
                                <number option>
| LOGIO, <bound>
| STRING, <bound>
| NAME, <bound>
<bound> ::= <number expression> ! ? ! ?
                                <number expression>
<number option> ::= , FIXED | <empty>
<element linkage> ::= , [ <bound pair list> ]
                                <pop linkage>

```



```

|<empty>
<pop linkage> ::=, LIFO|, FIFO|, RANDOM|
<empty>
<bound pair list> ::= <bound> : <bound>
|<bound pair list>, <bound> : <bound>
<name list> ::= <name> | <name list>,
<name>
<name> ::= <letter> | <name> <letter> |
<name> <digit>
<empty> ::= ( the null string of characters )
<number expression> ::= ( expression of type
number )

```

図 2 - 11 C A B A L の 例

例えば，次の例では，A，BはNUMBER属性を持つ有効桁 6 の scalar ，CはLOGIC で 有効桁 4 でエレメント数 10×10 の 2 次元の array ，DはSTRING属性を持ち，有効桁 10 でエレメント数が 15 コ位の last - in - first - out のスタック，EはSTRING で 有効桁が不定で エレメント数も不定の list を示している。

```

DECLARE A, B (NUMBER, 6),
        C (LOGIC, 4, [0:10, 0:10]),
        D (STRING, 10, [1:?15], LIFO),
        E (STRING, ?, [1:?], RANDOM);

```

ここでいう有効桁は，データ構造のとりうる 4 つの属性に対応し，その単位は

```

NUMBER (decimal) ..... digit
STRING (alphanumeric) ..... character
LOGIC ..... bit

```

NAME その name 全体

のようになっている。

宣言の中で？が出て来た場合には、ダイナミック・アロケーションが行なわれる。

Reduction ステートメント

Syntax の認識のためには、reduction ステートメントが用いられる。

これは例えば、

L: | 'WHILE', EXP | → STA | WHCODE ; RETURN |

のようなもので、スタックの内容を左半分の | | に囲まれた部分と比較して一致していれば "WHILE" と EXP とをスタックから pop - up し、代りに STA を入れる。

スタックの内容を示すのに用いるエレメントは、端記号と超変数を区別して、端記号は " " で囲んで表現する。さらに

| A, B, A | → A, B | RETURN |

のような場合は、右辺の A が左辺のどちらの A に対応するか不明なので、これを避けるため、左辺の A に右から順に添字を付け、

| A, B, A | → A(1), B | RETURN |

のようにする。

また C A B A L では、端記号を DEFINE ステートメントを用いて前以ってまとめて超変数であらわすことができる。

DEFINE UNOP = " - " | " + "

BNOP = UNOP | " / " | " * " | " ^ " | " 0 " | " ε " ;

スタックの内容が何であっても、置き換えを実行させる場合には、VITAL などの場合と同じような Sigma function を用いる。

さてスタックの内容が左辺と一致した時には、スタックの置き換えのあと、コントロールが右側の | 以降のステートメントに移される。この部分には、

すべてのCABALステートメントを書くことができる。但し reduction ステートメントは BEGIN, END で囲まれたブロックの中でなければ使うことができない。

一致しない場合は、コントロールは次の reduction ステートメントに移される。

ステートメント

CABALにおける reduction と code generation 以外の主なステートメントについて述べるがその前に、式について特異な点を示す。

CABALでは、データ・属性と同じく式も4つの属性をもつが、これはオペラントではなく、オペレーションによって決められる。但し式がオペラント1個のみから成る時はオペラントに依る。例えば、+, -, *, /, ↑, などの演算は、オペラントが何であっても演算の結果は NUMBERタイプとなる。このほか、modを取る演算 およびトランケーションは NUMBERタイプ、論理和、論理積、論理否定、関係演算は LOGICタイプ、concatenationは STRINGタイプの演算結果をとる。このような方法を用いて、式を算定するのは変数中の情報を自由に部分的にアクセスできるようにするためである。

CABALの主なステートメントは、代入文、条件文、繰り返し、push - popの4種で前3者は、いくらかSyntaxが異なるが、ALGOLとほぼ同じである。

push - pop ステートメントは、LIFO, FIFO, RANDOMの指定された順にリンクされた stack, queueやlistをpushしたり、popしたりするものであって Syntaxは次のとおりである。

$$\begin{aligned} \langle \text{push-pop statement} \rangle &::= \langle \text{stack operator} \rangle \\ &\quad \langle \text{name} \rangle \\ &\quad | \langle \text{stack operator} \rangle \langle \text{group reference} \rangle \\ \langle \text{stack operator} \rangle &::= \Delta | \langle \text{stack operator} \rangle \Delta \end{aligned}$$

```

|▽|<stack operator>▽
<group reference> ::= <name>
                        [<group definition>]
<group definition> ::= <index range list>
                        <number expression list>
<index range list> ::= <empty> | *, |
                        <index range list> *,

```

Code Generation

コード発生のためのステートメントは itemステートメントと code ステートメントを組にして用いる。

例えば

```

ITEM(SMT[4], SMT[1], #A, #B);
CODE(FOR#1←#4 TO #3 BY #2 DO);

```

では、#1はSMT[4]に、#4は#Bに、のように itemステートメントの内容には左から順に番号をつける。item ステートメント内に#が出てきたときは、それが定数であることを示している。code ステートメントは VITAL のコード・ブラケットと同じように、このCODE() の中には、任意の CABAL ステートメントを書くことができ、それらのコードを発生させるものである。この例では、

```
FOR SMT[4]←#B TO #A BY SMT[1] DO
```

に対応するコードが作られる。

未だ出現していないラベルへの ジャンプ命令に対しても、VITALと内じように CHAINとASSIGN function を用いて、処理することができる。

CABALでは、GO TO <ラベル名> NOW を用いると、いまコンパイラ・コンパイラが作り出したコード内のラベル名で指定されるステートメントへコントロールを移すことができる。そして generate されたコード内に

ある RETURN NOW ステートメントを実行すると、コンパイラ・コンパイラに戻ってくるのである。

ま と め

1968年に発表されたCABALについての文献によれば、このシステムの設計思想は、

1. トランスレータの適応性。
2. 作り上げることのできるトランスレータの範囲が多様であること。
3. 種々のエンバイロンメントに対して、柔軟性を持つこと。
4. モジュール化できること。
5. 専門のシステム・プログラマに対して、使いやすく簡潔であること。
6. 専門的でないユーザーに対して、平明であること。
7. 読みやすいこと。
8. 機種に依存しないこと。(machine independence)
9. 信頼性の高いこと。
10. エンバイロンメントとの適合性。

を基調とするものであって、実際にインプリメントした結果は、これらの要求事項をほぼ満足するものであると述べている。例えば、コンパイラ・コンパイラの適応性については、CO-routineの使用による並行処理、ストリング処理用の演算子、構文解析、master symbol tableがcode generator システム・ルーチンとユーザーの書いたセマンティック・ルーチンのいずれからのアクセスできること、高レベルの言語でcode generatorが書けることによる程度最適化したコードが作り出せることなどがある。

また、コンパイラ・コンパイラ CABAL の作り出すトランスレータは、インタープリタであっても、会話型のコンパイラ、あるいは、多重の言語システムであっても構まわない。さらにエンバイロンメントへの柔軟性にも富んでいてマクロの機能をつけ加えることも可能である。

そのほか、モジュール化や読みやすさ、信頼性も高いが、難点としては code generation に時間のかかることがある。これは、CODE ステートメントの書き方で、効率をよくすることができるであろうと書いてあるが、それには、このシステムでコンパイラを作成する人々の熟練度に依存する面が多いと解釈できる。

CABALは、コンパイラ・コンパイラとしてはかなり優秀なものであるらしいが、実際にこのシステムによって作り出されたコンパイラ、もしくはトランスレータがこの時点では数少ないことと、それ以降の文献が不足しているため、どの程度実用的なものであるのかという点は不明である。

参 照 文 献

1. P.K.Dove

Design highlights of CABAL — a compiler — compiler

FJCC 1968 PP 1321 ~ 1328

(図2-11)は上記文献より引用させていた。いた。

2. R.A.Brooker, D.Morris, et.al

Compiler — compiler facilities in Atlas autocode

Comput.J., Vol 9 1967 PP 350 ~ 352

3. R.W.Floyd

A description language for symbol manipulation

J.ACM Vol. 8 No10 1961 PP 579 ~ 584

2.8 G U L P

G U L P は、ケンブリッジ大学において、大形コンピュータ T I T A N につながっている P D P - 7 (8 k W / 1 W = 18 bit) 及びライトペン等のついたグラフィック・ディスプレイ (C R T) のシステムに対して開発されたものである。

G U L P の考え方は、Brooker & Morris の開発した C C から発想しているが、Syntax および Semantics の記述方法はかなりの違いがみられる。すなわち、C C においては P H R A S E , F O R M A T , R O U T I N E の 3 つで構成されていたが、G U L P においては (図 2 - 12) にみられるように同一ステートメントとして表現するようになっている。特に Semantics の記述に関しては、ユーザがルーチンを別で作って指定することができるが、ほとんどの場合 (表 2 - 1) に示した 21 個の標準ルーチンにより記述ができるとしている。

(図 2 - 12) に示すのは、加算のみを許す代入文及び出力を指定できるユーザ言語の定義である。

ここで . P L U S . のように 前後にピリオドがあるのが nonterminal 記号を示し . A S T A T ! . のようにエクスクラメーション記号がついているのが complete statement 記号を示す。phrase の指定のあとに : がありそれに続くのが Semantics 記述である。

```

1) Language definition
CHARACTERS;
'CR'; TTYPE; 215;           (Define carriage return)
'RUB'; TTYPE; 377;         (Define a rub-out character)

DELIMITERS;
space;

ERASERS;
'RUB';

TERMINATORS;
'CR';

MNEMONICS;
LAC=200000;                 (Load accumulator from store)
DAC=40000;                  (Deposit accumulator in store)
ADD=300000;                 (Add from store to accumulator)

CONVERSIONS;                (Octal conversion, numbered 1)
1;
0,1,2,3,4,5,6,7;

SYNTAX AND SEMANTICS;
0,1,2,3,4,5,6,7 → .DIGIT.;
A,B,C,D,E,F,G,H,I,J,K → .LETTER.;

.LETTER..LETTER★?.,.DIGIT★? → .NAME.;           (name of a variable)
==→ .EQUALS.;
) → .KET.;
+ → .PLUS.;
.PLUS..NAME → .TERM..INSTR(ADD,SYMADD(C2));       (positive term, to be added to sum)
.PLUS?..NAME → .FTERM..INSTR(LAC,SYMADD(C2));     (first term with optional plus sign. This term is
                                                    the first loaded in the accumulator.)
.NAME..EQUALS..FTERM..TERM★? → .ASTAT!..INSTR(DAC,SYMADD(C1));
                                                    (Complete arithmetic assignment statement, se-
                                                    mantics to store result)
.NAME..KET..DIGIT★ → .CONDEF!..INSTR(LAC,INPUT(C3,1));INSTR(DAC,DEFSYM(C1,1));
                                                    (Statement to define a symbol and its initial
                                                    value)
OUTPUT → .OUT.;
.OUT..NAME → .PRINT!..OUTPUT(SYMADD(C2),1);      (Defines a compound character)
                                                    (Statement for printing value of a variable)

```

図 2 - 1 2 G U L P の例

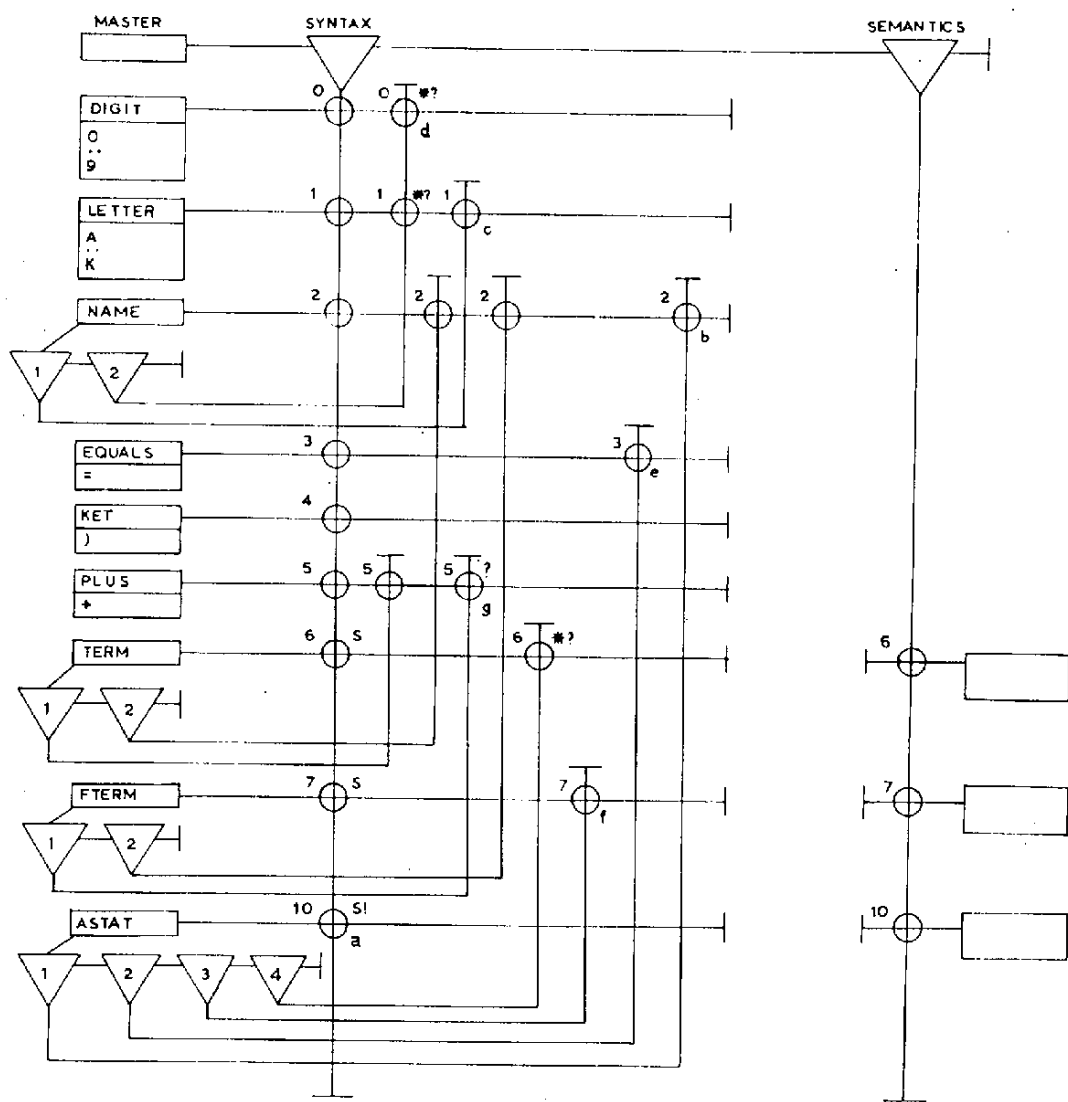


図 2-13 GULPによって使用されている ring structure

入力されたユーザの Syntax および Semantics の記述は, A S P - 7 (associative structure package) によりまず (図 2 - 13) に示すようなリング・ストラクチャーの形に変換される。

ユーザ言語のコンパイルは, 入力ストリングにより complete statement 記号から始めて, 構文解析のためにテストされた規則を入れるもの及び, 構文解析の結果でき上った tree を入れる 2 つのスタックを利用して, リング・ストラクチャーをおっかけることによりおこなわれる。

このシステムは, C R T を端末として考えているために, ライトペンの割り込み, ファンクションキーの割り込みを, 1 つの文字に変換することにより通常の入力記号と同じようにあつまっている点に特徴がある。

同じようにグラフィックを対象としたシステムに A E D があると報告されているが, G U L P が参照している A E D 関係文献ではその特徴が不明である。

表 2 - 1 G U L P の Semantics ルーチン

名 前	argument の 数	routine の 結 果	定 義
I N S T R	変 数		すべての argument の値を加え, instruction か data word として, STORE 内にそれを置く。 dummy location counter に 1 を加える。
DEFSYM	2	reserve された location の address	1 つの SYMBOL を定義する。 ARG. 1 は SYMBOL 名を構成する文字列を示す。 ARG. 2 は reserve される文字数を示す。 もし, SYMBOL がすでに定義されていたら ERROR である。
SYMADD	1	SYMBOL の address	SYMBOL TABLE を look up する。 argument は SYMBOL を示す。

名 前	argument の 数	routine の 結 果	定 義
INPUT	2	data word の address	input string を data word に変換 する。 ARG. 1 は input string を示す。 ARG. 2 は変換数を示す。
OUTPUT	2		data word を string に変換し、それ を output する。 ARG. 1 は data を示す。 ARG. 2 は 変換数を示す。(もしなけれ ば変換はされない)。
SET	2		Semantic variable (ARG. 1) を setする。 ARG. 2 はなくてもよい。
INTENS	1		表示の intensity (強さ) を setする。
DSPLAY	1		ARG. が ON なら picture を出し、 ARG. が OFF なら出さない。
PEN	1		ARG. が ON なら LIGHT-PEN 使 用可である。 ARG. が OFF なら不可である。
POINT	2		座標へ POINT (ARG. 1, ARG. 2) を画 く。
LINE	2		X, Y increment を使って、line を引 く (ARG. 1, ARG. 2 別々に)
ALINE	4		POINT (ARG. 1, ARG. 2) から POINT (ARG. 3, ARG. 4) へ line を引く。
SETDIS	4		DISPLAY の set up ARG. 1 は intensity (強さ) を示す。 ARG. 2 は scale を示す。 ARG. 3 は PEN に対応する。 ARG. 4 は DSPLAY に対応する。

名 前	argument の 数	routine の 結 果	定 義
CROSON	0		picture に tracking cross を付加する
SETX	2		座標 (ARG.1, ARG.2) へ tracking cross を置く。
CROSS	2		tracking cross がある座標を ARG.1 と ARG.2 へ読み込む。
DEFCHR	2		ENDCHR call があるまで、 表示 character を定義する。 character については、ARG.1 で名前 が示され、ARG.2 で示される array 長が reserve される。
ENDCHR	0		DEFCHR の end を示す。
STORE	変 数		名前 (ARG.1) で連想される array に data を set する。 ARG.2 は使用される最初の element を 示す。 ARG.3 は array に置かれるべき連続した data を先に与える。
LOCK	3		ARG.1 によって指定される line にもっ とも近い point に tracking cross を置き ARG.2 と ARG.3 にその座標を置く。
DELETE	1		(DEFCHR によって前もって定義されて いるとする) 表示 character (ARG.1) を system から消去する。

参 照 文 献

1. P.J.Pankhurst
GULP- A compiler - compiler for verbal and graphic languages
Proc. ACM 1968 PP 405~421
(図2-12, 13), (表2-1)は上記文献より引用させていただいた。
2. R.A.Brooker, D.Morris
A General Translation Program for Phrase Structure
Languages
J.ACM 1961, 7. PP 1~10
3. R.A.Brooker, I.R. Maccallum, D.Morris, J.S.Rohl
The Compiler Compiler
Annual Review in Automatic Programmig Vol 3, 1963
PP 229 ~ 275 Pergamon Press
4. D.T.Ross
The AED approach to generalized computer - aided design
Proc. ACM 1967. PP 367 ~ 385

2.9 INSCAN

INSCANはRADC(Rome Air Development Center)において開発されたDM-1(Generalized data management System)のためにインプリメントされたSyntax-directed language processorである。

このシステムの概略は、Action Graphの導入から始まる。Action Graphは入力ストリングに対しておこなう処理をコントロールするものと考え、各要素はそれぞれの命令とみなすことができる。

ここでその例を出してみよう。

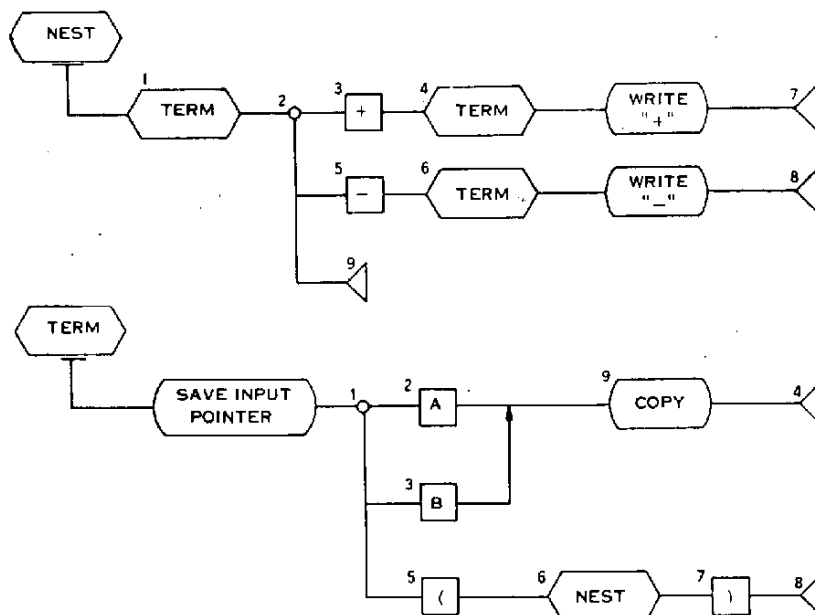


図 2-14 Action Graph の例

これは、加減算記号及びカッコより成る式を逆ポーランドの形になおす Action Graph である。

NEST, TERM はそれぞれのグラフ名をさす。

NEST から実行が始まり、1によりTERMが呼び出される。TERMに

においては、SAVE INPUT POINTER により、入力記号のバッファポイントがSave される。グラフ 'TERM' の中の1は分枝を示している。すなわち、2において入力記号が "A" ならば、2, 9の順に実行して呼んだグラフにもどる。そうでなければ、3において入力記号が "B" かどうか調べられ、等しければ3, 9の順に実行、等しくなければ5の "(" との比較がなされる。

グラフ 'TERM' における9のCOPY は入力記号をSAVE INPUT POINTER においてSave した点から現在まで処理した点までを入力することを指定する。

すなわち、グラフ 'TERM' においては、オペランドであればこれを出力し "(" であればNEST を呼び出し、そのあとに ")" があるかどうか調べるものである。TERM からコントロールが帰ってくると、2の分析により "+" あるいは "-" と入力記号との比較がなされ、等しければTERM が呼び出されその後、記号 "+" あるいは "-" が出力される。

Action Graph は計算機に対する入力としては不適當であるので、グラフをセンテンスの形にして入力する。

```
NEST: EXECUTE TERM;CHOICE(3, 5, 9);  
      3: "+" ;EXECUTE TERM; WRITE "+" ;GOOD;  
      5: "-" ;EXECUTE TERM; WRITE "-" ;  
      9: GOOD.  
  
TERM: SAVE INPUT POINTER; CHOICE(2, 3, 5);  
      2: "A" ;9: COPY; GOOD;  
      3: "B" ; GOTO 9;  
      5: "(" ; EXECUTE NEST; ")" ; GOOD.
```

図 2 - 15 INSCAN の例

INSCANには、例に現われたセンテンス以外に次のものがある。

RECURSE : 自分自身のグラフを呼び出す。

CALL : 外部サブルーチンを呼び出す。

また、アクショングラフによる記述の中で次のものが禁止されている。

1. 左 回 帰

すなわち

```
A : CHOICE ( 1 , 2 ) ;  
  1 : RECURSE ; " A " ; GOOD ;  
  2 : " A " ; GOOD .
```

という表現は禁止されており、

```
A : " A " ; EXECUTE B ; GOOD .  
B : CHOICE ( 1 , 2 ) ;  
  1 : EXECUTE A ;  
  2 : GOOD .
```

というように書きなおさなければならない。

2. 入力記号を含めたあともどり

すなわち、

```
Q : CHOICE ( 1 , 2 ) ;  
  1 : " B " ; EXECUTE R ; GOOD ;  
  2 : " B " ; EXECUTE S ; GOOD .
```

という表現は禁止されており、

```
Q : " B " ; CHOICE ( 1 , 2 ) ;  
  1 : EXECUTE R ; GOOD ;  
  2 : EXECUTE S ; GOOD .
```

というように書かなければならない。これは前者の表現では1において入力記号と"B"と比較され等しければRを実行し、ここでfalseの状態の時は、すでに1グラフ要素を実行しているために分枝点にもどれないからである。

3. 分枝指定において null を最初にもってくること。

分枝は書かれた順にテストするため null は最後にもっていかなければならない。

INSCAN における action Graph の使用は、flow chart 的なものとしてわかりやすいと思われるが論文の中で Semantics の記述に関しては、何も述べていないので断定はできないが、複雑な処理は CALL によって外部サブルーチンの呼び出しをおこなうのが普通の方法であると考えられる。もしそうであるならば、Semantics 記述の手段をもう少し考えておいてほしいと思う。また、DM-1 システム中の何の役割をはたすかについても不明である。

参 照 文 献

1. M. Resnick, J. Sable

INSCAN: a Syntax-directed language Processor

Proc. ACM 1968 PP 423 ~ 432

(図 2-14, 15) は上記文献より引用させていただいた。

2. P. J. Dixon, J. Sable

DM-1-a generalized data management System

SJCC 1967 PP 185 ~ 198

2.10 Canonic Translator

Canonic Translator はマサチューセッツ工科大学 (MIT) において Project MAC の一環として作成されたものである。

このシステムの最大特徴は、通常のコパイラ・コンパイラが Context Free 文法をあつかうのに対し、Context Sensitive 文法をあつかえるように作られている点にある。

このシステムでは、規則は Canon と呼ばれたとえば次のように記述する。

- ① $a_1 \text{ set } A_1 \phi a_2 \text{ set } A_2 \phi \dots \phi a_n \text{ set } A_n \vdash b \text{ set } B$

これは $a_1, a_2, \dots, a_n$ がそれぞれ $\text{set } A_1, \text{set } A_2, \dots, \text{set } A_n$ の要素の時 b は $\text{set } B$ の要素であることを示す。

例

$v \text{ letter } \phi x \text{ expression } \vdash v = x \text{ assignment}$

これは、BNF で書けば

$\langle \text{assignment} \rangle ::= \langle \text{letter} \rangle = \langle \text{expression} \rangle$

の意味となる。

- ② $a_1 a_2 \dots a_n \text{ set } B$

これは、 $a_1, a_2, \dots, a_n$ が $\text{set } B$ の要素であることを示す。

例

$0 \ 1 \ 2 \ 3 \ \dots \ 8 \ 9 \ \text{digit}$

これは、BNF で書けば

$\langle \text{digit} \rangle ::= 0 \mid 1 \mid 2 \mid 3 \mid \dots \mid 8 \mid 9$

の意味となる。

- ③ $\langle a_1 a_2 \dots a_n \rangle$

これは、 $a_1 a_2 \dots a_n$ の順にならんでいることを意味する。

例

$x \text{ identifier } \phi y \text{ letter } \phi z \text{ digit } \vdash \langle x \langle y, z \rangle$

identifier

これは、BNF で書けば

$$\langle \text{identifier} \rangle ::= \langle \text{identifier} \rangle \langle \text{letter} \rangle | \langle \text{identifier} \rangle \langle \text{digit} \rangle$$

この記述方法をもちいて、Syntax および Semantics を記述するわけである。

ここで PL/I-like な GO TO ステートメントを考えてみよう。

DCL 変数 LABEL; ラベル変数の定義

GO TO ラベル変数; L 1, ラベル変数

BR 1

の命令の発生。

GO TO ラベル定数; B ラベル定数

の命令の発生

このような場合

ℓ label-const $\vdash$ GO TO ℓ ; goto stm with ref label ℓ

label-var $\wedge^{*1}$ assembler stms

B * BRANCH TO ℓ

ℓ label-var $\vdash$ GO TO ℓ ; goto stm with ref label $\wedge$

label-var ℓ assembler stms

L 1, ℓ

BR 1

v label-var $\vdash$ DCL V LABEL; dcl stm with dcl

label-var

というように記述する。goto stm with ref label は、他の場所においてこのラベルが定義されているかどうかのチェックのために使用する。

*1 $\wedge$ は null (空) を意味する。

このシステムは、Context sensitive な文法を受け入れるために、理論的には、他のシステムにくらべると適用範囲ははるかに大きい、次のような欠点を持っている。

- ① プログラムのスピードは普通のコンパイラよりうchwにみて1000倍遅い。
 - ② ソースプログラムが多量の時に適応できない。
 - ③ エラーが発生した時，ソースプログラムの中で最後に調べられた文字を示すことができるがコンパイルの続行ができない。
- このため実用化には問題点が多いといえる。

参 照 文 献

- 1. J.J. Donovan, H.F. Ledgard
A formal system for the specification of the syntax and
translation of computer languages
FJCC 1967 PP 553~569
- 2. J.W. Alsop
A canonic translator
MAC-TR-46 1967

3. 構文解析手法

構文解析手法の開発は、TWS に掲載されているものは省略した。また、Mckeeman による Mixed strategy precedence については、2.1 XPL で紹介したのでここでは省略した。

3.1 Weak precedence 文法および右順位文法

weak precedence 文法は、J. D. Ichbiah 及び S. P. Morse により、右順位文法は井上によりほぼ同時期に定義されたものであり、文法の範囲は若干右順位文法の方が広いが、基本思想は全く同じものである。

Weak precedence 文法

次のような条件を満たすとき、文法を weak precedence 文法という。

- ① S を start 記号としたとき。

$$S \Rightarrow^* \alpha AB\beta$$

なる 2 つの記号 A, B の間に、一義的な weak precedence 関係が成立する。

- ② 右辺が全く同じである生成規則はない。

- ③ $U_1 \rightarrow x S y$

$$U_2 \rightarrow y$$

のとき、S と U_1 の間に weak precedence 関係がない。

weak precedence 関係

- ① $U \rightarrow \alpha U_1 B\beta$ なる生成規則があるとき

$$U_1 \Rightarrow^* z A \text{ なる } A \text{ に対して}$$

$A \cdot > B$ である。

- ② $U \rightarrow \alpha U_1 U_2 \beta$ なる生成規則があるとき

$$U_1 \Rightarrow^* z A, U_2 \Rightarrow^* B \omega \text{ なる } A, B \text{ に対して}$$

$A \cdot > B$ である。

③ $U \rightarrow \alpha AB\beta$ なる生成規則があるとき

$A < \cdot B$ である。

④ $U \rightarrow \alpha AU_1\beta$ なる生成規則があるとき

$U_1 \xRightarrow{*} B\omega$ なる B に対して

$A < \cdot B$ である。

右 順 位 文 法

次のような条件を満たすとき，文法を右順位文法という。

① S を start 記号としたとき

$$S \xRightarrow{*} \alpha AB\beta$$

なる 2 つの記号 A, B の間に一義的な右順位関係が成立する。

② 右辺が全く同じ生成規則はない。

③ $U_1 \rightarrow \alpha\beta$

$$U_2 \rightarrow \beta$$

のとき

$$C \rightarrow \gamma H_i(\alpha) D \delta, D \xRightarrow{*} T_{n-i}(\alpha) B \delta'$$

は存在しない。

右順位関係

ここで次のものを定義する

$$L(A) = \{ F \mid A \xRightarrow{*} F\alpha, A \in V_N, F \in V, \alpha \in V^* \}$$

$$R(A) = \{ F \mid A \xRightarrow{*} \alpha F, A \in V_N, F \in V, \alpha \in V^* \}$$

$$LT(A) = \{ F \mid A \xRightarrow{*} F\alpha, F \in V_T, A \in V_N \} \cup \{ A \mid A \in V_T \}$$

① $U \rightarrow \alpha U_1 U_2 \beta$ なる生成規則があるとき

$A \in R(U_1), B \in LT(U_2)$ なる A, B に対して

$$A \cdot > B$$

② $U \rightarrow \alpha A U_1 \beta$ なる生成規則があるとき

$B \in LT(U_1)$ なる B に対して

$$A \leq \cdot B$$

構文解析

weak precedence 文法, 右順位文法のいずれにおいても構文解析手法は全く同じである。すなわち, スタックの頭の記号 (a) と入力記号 (b) の順位関係を調べる。 $a < \cdot b$ ならばスタックに push down し, $a \cdot > b$ ならばスタックと生成規則の右辺が一致するものの中で最長の規則を適用し, 新しい記号をスタックの頭に入れて解析を進めてゆく, $a < \cdot b$ でも $a \cdot > b$ でもなければエラーである。

参 照 文 献

1. J. D. Ichbiah, S. P. Morse
A Technique for Generating Almost Optimal Floyd-Evans
Production for Precedence Grammars
C. ACM Vol 13 No 8 1970 PP 501~508
2. 井上 謙 蔵
順位言語の右順位解析
情報処理 Vol 11 No 4 1970 PP 231~233
3. 井上 謙 蔵
右順位文法
情報処理 Vol 11 No 8 1970 PP 449~456
4. 井上 謙 蔵 他
右順位文法に対する構文解析プログラムの発生
第12回 プログラミング・シンポジウム 1971 PP 64~76

3.2 Total Precedence 文法

この文法は, Alain Colmerauer によって定義されたものである。

Binary Relations

Set E 上の binary relation を Cartesian 積 $E \times E$ の subset ρ として定義する。

$$a \rho b \quad \text{.....} \quad (a, b) \in \rho$$

binary relation はギリシア小文字, 又は特殊記号で示す。

$E \times E$ に関する ρ の complement $\bar{\rho}$ は

$$\bar{\rho} = E \times E - \rho$$

binary relation $\exists \rho$, on set E

$\Rightarrow \rho \sigma$ は E 上の binary relation で

$$a \rho \sigma b \stackrel{\text{def}}{=} \exists c \in E \quad \text{s.t.} \quad G \rho c, \quad c \sigma b$$

この積が結合的であり, 結合に関して左右に分配できることが容易にわかる。結合性によって, binary relation ρ についての (非負整数) 累乗が定義できる。

$$\rho^i = \rho \rho^{i-1}, \quad a \rho^0 b \equiv (a = b)$$

transitive closure ρ^+ は次のように定義できる。

$$\rho^+ = \bigcup_{i=1}^{\infty} \rho^i$$

E が n 元からなる有限集合である時

$$\text{対 } a, b \in E \quad \text{s.t.} \quad a \rho^0 b, \quad k > n$$

$$\Rightarrow \ell \leq n \quad \text{s.t.} \quad a \rho^\ell b$$

$$\text{従って} \quad \rho^+ = \bigcup_{i=1}^{\infty} \rho^i = \bigcup_{i=1}^n \rho^i$$

Relation α, λ, ρ の定義

文法 G の $V_T \cup V_N$ 上の 3 つの binary relation を定義することが必要で

ある。ここで α は "adjacent" , λ は "left" , ρ は "right" を指すことにする。

$\forall A, B \in V_T \cup V_N$ に対して

$$A \alpha B \equiv [\exists U \in V_N, \exists x, y \in (V_T \cup V_N)^* \text{ s.t. } U \rightarrow xAB y]$$

$$A \lambda B \equiv [\exists y \in (V_T \cup V_N)^* \text{ s.t. } A \rightarrow By]$$

$$A \rho B \equiv [\exists x \in (V_T \cup V_N)^* \text{ s.t. } B \rightarrow xA]$$

あいまいでない total precedence relation の存在に関する必要十分条件は

$$(\rho^+ \alpha \cap \alpha) \cup (\alpha \lambda^+ \cap \alpha) \cup (\rho^+ \alpha \lambda \cap \alpha) \cup (\rho^+ \alpha \cap \alpha \lambda^+) = \emptyset \text{ である。}$$

3つの binary relation $<\cdot, \cdot, \cdot>$ が (文法 G 上の $V_T \cup V_N$ 中の vocabulary についての) total precedence relation であり, これに関連する analyzer が $V_T \cup V_N$ 上のすべての String を解析することのできるのは, 次の包含関係を満たすときであり, かつ, この時に限る。

$$\alpha \subset \cdot, \alpha \lambda^+ \subset <\cdot, \rho^+ \alpha \subset \cdot >, \\ \rho^+ \alpha \lambda^+ \subset <\cdot \cup \cdot >$$

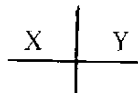
Relation $<\cdot, \cdot, \cdot>$ の計算

Boolean matrix によって α, λ, ρ の各 relation を表現することができ $<\cdot, \cdot, \cdot>$ を表現する Boolean matrix を得ることができる。これは, コンピュータ上での最も簡単な手続きである。

これを手で行なうためには, 次のようにすればよい。

$$(x, y \in (V_T \cup V_N)^*):$$

- ① $U \rightarrow xXYy$ の形の規則が存在するような X, Y の対の各々に対して 十字を書き, X, Y を書きこむ。



- ② $U \rightarrow xX$ の形の規則が存在し, かつ U が十字の左半分にあらわれており, かつ X が同じ十字の左下にあらわれていない U と X の対の各々に対して左下の隅に X をつけ加える。これ以上の変更が行なわれないところ

まで step 2 をくりかえす。

- ③ $U \rightarrow Yy$ の形の規則が存在し、かつ U が十字の左にあらわれていて、かつ、 Y が同じ十字の右下にあらわれていない U と Y の記号対の各々に対して、右下に Y を加える。これ以上の変更が行なわれないところまで step 3 をくりかえす。

さて、記号 X, Y の対の各々に対して、次のことが容易にわかるであろう。

- ① configuration $\begin{array}{c|c} X & Y \end{array}$ は $X \alpha Y$ に対応している。
 ② configuration $\begin{array}{c|c} X & Y \\ \hline & Y \end{array}$ は $X \alpha \lambda + Y$ に対応している。
 ③ configuration $\begin{array}{c|c} & Y \\ \hline X & \end{array}$ は $X \rho + \alpha Y$ に対応している。
 ④ configuration $\begin{array}{c|c} & Y \\ \hline X & Y \end{array}$ は $X \rho + \alpha \lambda + Y$ に対応している。

そして、relation $<, =, >$ はただちに導かれる。

構文解析

total precedence 文法による構文解析は、 S_1, S_2 という 2 つのスタックをもちいる。

ここで $(\alpha, \beta) \rightarrow (\alpha', \beta')$ とは、スタック S_1 に頭から順に α というものがはいつており、同様に S_2 に β があつた時、スタック S_1 に α の代り α' 、スタック S_2 に β の代りに β' を入れることを意味する。

構文解析は次の手順によっておこなわれる。

初期状態として、スタック S_1 には $\#$ 、スタック S_2 には解析したい記号列 ω 及び $\#$ を入れておく。

- ① $A < B$ あるいは $(A = \# \text{ かつ } B \neq \#)$ ならば
 $(A, B) \longrightarrow (A(B, \wedge)$
 ② $A = B$ ならば
 $(A, B) \longrightarrow (AB, \wedge)$
 ③ $A > B$ ならば

$(A, B) \longrightarrow (A], B)$

④ $([u], \wedge) \longrightarrow (\wedge, U)$

これは、 $U \rightarrow u$ なる生成規則の適用を示す。

(注) $\wedge$ は空を示す。

この手順により、スタック S_1 に $\#$ 、スタック S_2 に $S\#$ (S はStart記号を示す) になった時解析は終了する。

参 照 文 献

1. A. Colmerauer

Total Brecedence Relation

J. ACM Vol. 17 No.1 1970 PP 14~30

3.3 SLR(k)

SLR(k) は F. L. DeRemer により定義されたものである。

Context Free 文法が $s+1$ 個の生成規則を持ち、 $\#_0, \#_1, \dots, \#_s$ を文法の中で使用していない特殊記号とする。

P 番目の生成規則が $A \rightarrow \omega$ で、 $\alpha' = \rho A \beta$ と $\alpha = \rho \omega \beta$ が Start 記号 S から $S \xrightarrow{*} \alpha' \rightarrow \alpha$ という canonical derivation される canonical form であるとき、 $\rho \omega \#_p$ は α の characteristic string であるという。

この characteristic string は regular 言語なのでこれを受け付ける Finite State Machine を作ることができる。これを Characteristic Finite State Machine (CFSM) と呼ぶ。

例として次のような生成規則を持つ文法の CFSM を作り上げると (図 3-1) になる。

$$S \longrightarrow \vdash E \vdash$$
$$E \longrightarrow E + T$$
$$E \longrightarrow T$$
$$T \longrightarrow P \uparrow T$$
$$T \longrightarrow P$$
$$P \longrightarrow (E)$$
$$P \longrightarrow i$$

CFSM において $\#_0$ の遷移と他の記号の遷移がある状態を inadequate 状態という。((図 3-1) においては 7 の状態)。

この inadequate 状態において、 k 個の記号の look ahead により状態遷移を一意的に決定できる時、この文法を Simple LR(k) (SLR(k)) と言う。

例においては、 $\#_4$ の 1 記号 look ahead の set は $\{ \vdash, +,) \}$ となり分離できるので、SLR(1) 文法である。

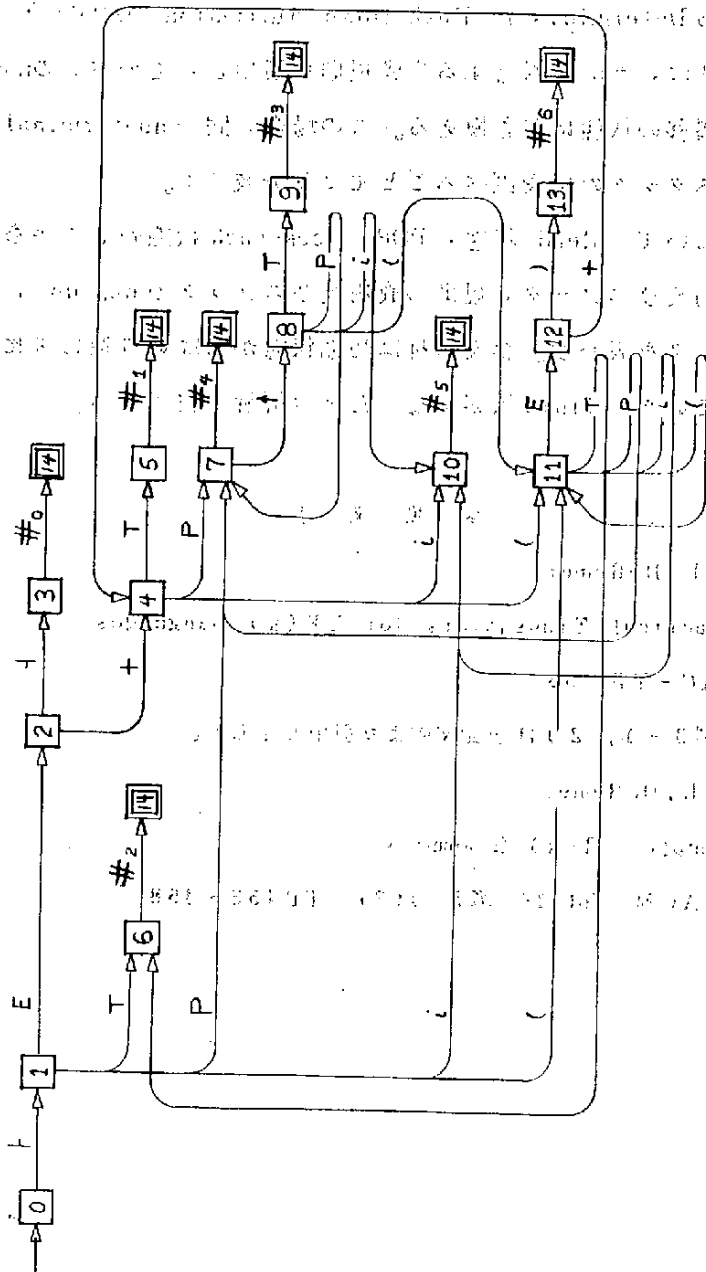


图 3-1 CFSM

構文解析

CFSM から Deterministic Push Down Automaton (DPDA) を作り上げる。この方法は、 $\#p$ に対応する生成規則を適用し、その左辺の nonterminal 記号に対する遷移の状態に置き換える。この場合、同一 nonterminal 記号が 2 個以上あればスタックの中を調べることにより分岐する。

DPDA において、Read 状態、POP、look-back 状態の 3 つがあり、それぞれ記号の入力及びスタック、規則の適用及びスタックの pop up、スタックの内容のチェックを表わす。構文解析は初期状態からはいり遷移図にしたがって状態を移していき、final 状態になったとき解析が終了する。

参 照 文 献

1. F. L. DeRemer

Practical Translators for $LP(k)$ Languages

MAC-TR-65

(図 3-1, 2) は上記文献より引用しました。

2. F. L. DeRemer

Simple $LR(k)$ Grammars

C. ACM Vol 14 No 7 1971 PP 453 ~ 458

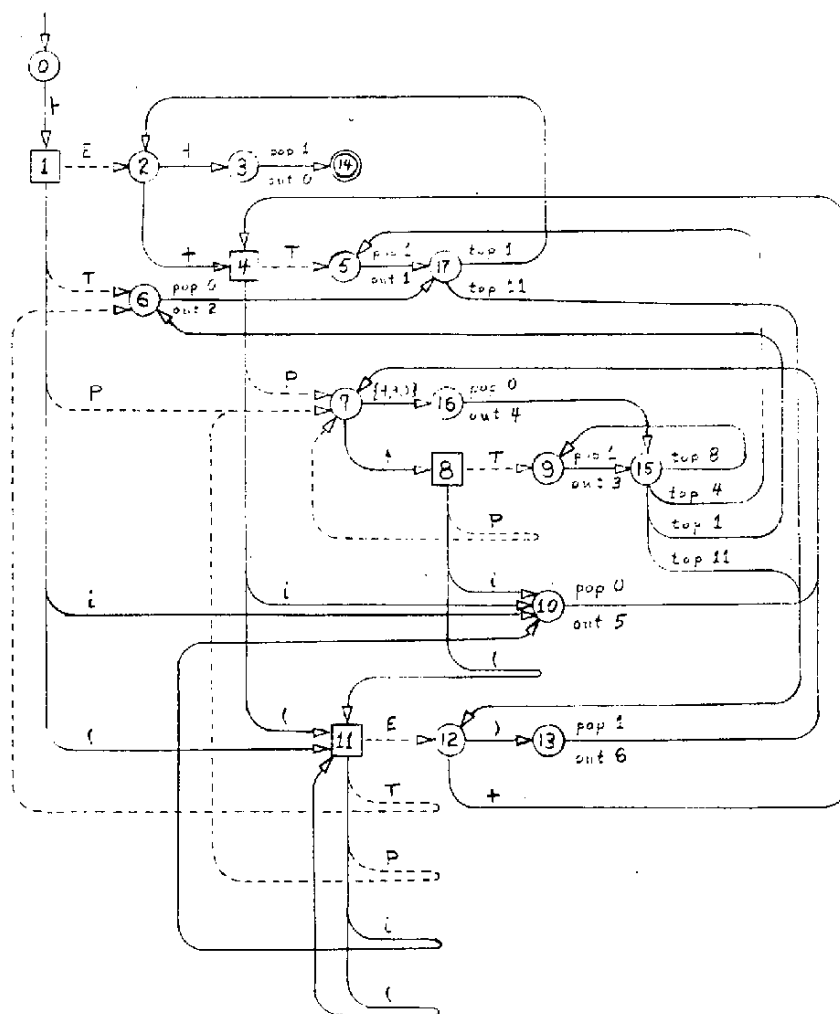
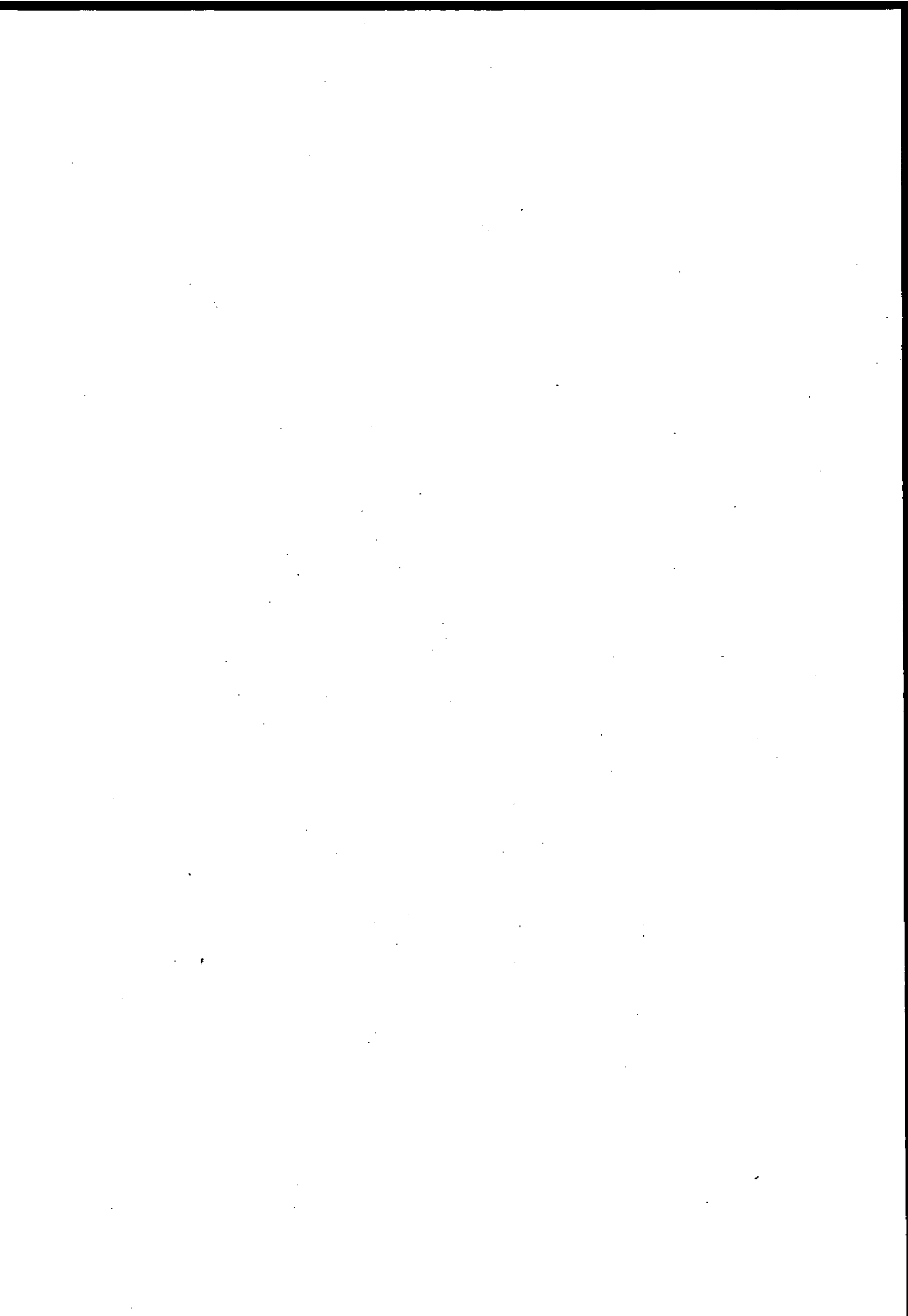


图 3-2 DPDA



第 2 部

コンパイラ・ジェネレータの作成

第 1 編

コンパイラ・ジェネレータ J P A C

目 次

第 I 編 コンパイラ・ジェネレータ JPAC

1. JPAC の概要	99
2. JPAC 動作概略	101
3. JPAC 言語仕様	105
3.1 SYNTAX DEFINITION	106
3.2 SEMANTICS DEFINITION	112
3.3 Lexical Analyzer	117
3.4 エラー処理ルーチン	123
4. STAGE 内部処理	124
4.1 処理概略	124
4.2 生成規則グラフ作成	126
4.3 nonterminal 記号つなぎあわせ処理	130
4.4 ϵ 遷移処理	136
4.5 LOOK-AHEAD 処理	139
4.6 $L(m)R(k)$ 処理	143
4.7 遷移図作成処理	154
5. 遷移テーブル類	157
6. 中間テーブル類	163
6.1 遷移図作成用ワークテーブル	163
6.2 遷移図出力用ワークテーブル	167

1. J P A C の 概 要

J P A C は F A C O M 230-60 の Monitor V の管理下で動作するコンパイラ・ジェネレータであり，ジェネレートされるコンパイラも F A C O M 230-60 用のものである。

このシステムは，目的とするコンパイラ言語の文法の Syntax および Semantics を入力することにより所要のコンパイラを作り出すことができる。

現在までに開発された内外のコンパイラ・ジェネレータは研究的要素の特に強いものと，ある程度実用性を持ったものに分けられるが，J P A C はどちらかと言えば後者に属し，特に次のような点を目標とした。

- ① 使いやすさ … 必ずしもコンパイラ作成の経験者でなくともコンパイラ言語の文法を B N F で記述すること及びプログラミングの力があれば使えることを目標とする。
- ② 実 用 性 … 今回作成するものは F O R T R A N クラスの言語のコンパイラが作成でき，コンパイラの効率（処理速度，メモリ占有量など）が実用上問題とならない程度のものであること。
- ③ 拡 張 性 … J P A C の仕様は，この報告書に記載されている範囲で終わるものではない。すなわちこれを使用した結果，その仕様を変更する必要性がでてきた場合には修正，変更が簡単にできなければならない。

J P A C で使用している構文解析手法，Syntax および Semantics 記述方法の特徴は次の点である。

① De Remer の手法の使用

De Remer は Knuth の作り出した L R (k) 文法の解析が膨大なテーブルを必要とするため，これを小さくすることに努力した。その結果，S L R (k)，L (m) R (k)，L R (k) の順に，順次大きな文法を定義してそれらに対し統一的な解析手法を作り出した。この手法によれば小さな文法の範囲におさまるものは S L R (k)

を，大きな文法の範囲になれば $L^{(m)}R^{(k)}$ を使用するという方法がとり得る。

② LSD 言語の開発

LSD は第 1 部の概説で述べたコンパイラ記述言語に必要となる機能をすべて持っている言語である。この言語によりコンパイラの Semantics をプロシージャルに記述できる。

LSD は Semantics 記述をするに十分な機能を持っているが，今後 Semantics 記述に必要となる標準ルーチンをまとめ，ライブラリ化することによりさらに強力な機能を持たせることができる。

LSD はコンパイラ記述言語として独立に使用することもでき，Syntax 解析プログラムも LSD を使用して作成した。LSD の外部仕様書はこれだけを独立して利用できるため第 II 編にまとめた。

③ SYNTAX の記述

Syntax の記述は，規則が適用された時に制御をわたすセマンティックス・ルーチンを BNF 記述のあとに指定するという形でおこなわれる。このセマンティックス・ルーチンは LSD で記述した Procedure あるいは標準ルーチンでなければならない。さらに標準ルーチンを整備すれば LSD によるプロシージャルな記述を使うことなく Semantics を簡単なルーチン・コールだけで済ますことができる。

2. JPACの動作概略

JPACでは、作成したいコンパイラの文法を入力してシンタックス・テーブル等を作り出すSTAGEとSemanticsを記述するためのシステム記述言語のコンパイラLSDの2つのステップより構成されている。

STAGEはバックス記法(BNF)で書かれたコンパイラの文法を入力して、遷移テーブル等のSyntax関係のテーブルを出力する。

STAGEが受けつける文法はSLR(k) ($k=0, 1, 2$)あるいはL(m)R(k) ($k=0, 1, m=1, 2$)である。

ユーザの記述したSyntaxは、まずSLR(0)文法が適用され得るか否かが調べられ適用できなかったならば、SLR(1), SLR(2)の順で調べ適用できるレベルで遷移テーブル等が出力される。しかしSLR(k)自体が適用不可能ならば今度はL(1)R(0), L(1)R(1), L(2)R(0), L(2)R(1)の順に調べ、適用可能な文法における遷移テーブル等が出力される。この遷移テーブル等はアセンブラの形で出力される。SLR(k), L(m)R(k)文法のいずれも適用不可能な場合は、どの規則が原因となって不可能であるかを表示して処理を中止する。

LSDはコンパイラ作成の手法として一般に使用する機能すなわちテーブル、スタック、ビット、文字のとりあつかいが簡単にできるように考慮したシステム記述言語であって、Syntax記述部分で定義した生成規則の適用によりどのようなアクションをとるかを指示する部分である。LSDもオブジェクト・プログラムとしてアセンブラを出力する。

STAGE, LSDの2つのステップで出力され、アセンブラ言語で書かれたオブジェクトは、FASPステップによりリロケートブルになる。そして必要とあればLIBEによりユーザ・ライブラリの作成もできる。

コンパイラ作成のためにはリロケートブルおよびシステム・ライブラリ、必要ならばユーザ・ライブラリをつなげるLIEDにより実行形式を作ればよい。

おのおのステップは全く独立に実行することができる。すなわち文法の定義

だけを変える必要があれば STAGE-FASP-LIED のステップを使えばよい。
 またコンパイラ作成途中において Semantics 記述部分に虫があった場合には
 LSD-FASP-LIED ステップにより新コンパイラを作りあげればよい。こ
 の手順は (図 2-1) JPAC の概略に示しているとおりである。

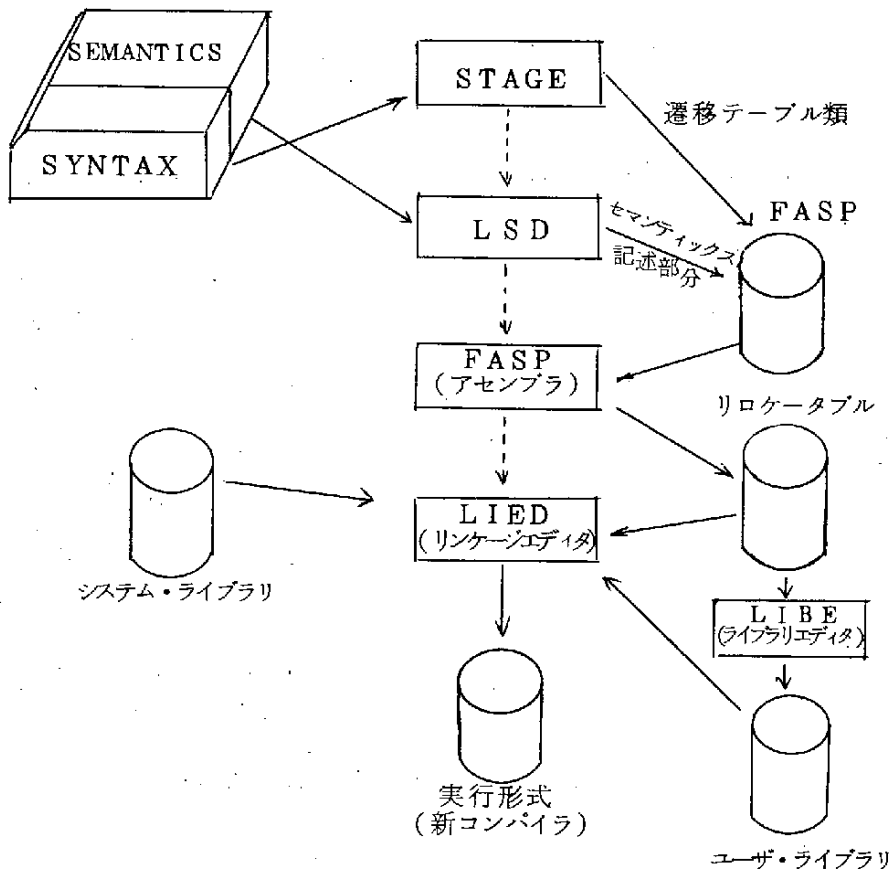


図 2-1 JPAC の概略

JPAC を使用することにより作り出されるコンパイラは (図 2-2) の形をしている。この中で遷移テーブル類およびセマンティックス・ルーチンはそれぞれ STAGE および LSD ステップによって作り出されるものである。

コントローラはまづ最初にユーザがOWNコーディングした Initiator にコントロールをわたす。Initiator においてはセマンティックス・ルーチンの初

期設定等をおこなう。再びコントロールが Initiator からもどってくると Lexical Analyzer を通してソース・プログラムを入力し、遷移テーブルにしたがって構文解析を進める。

遷移テーブルの中で、セマンテックス・ルーチンの呼び出し指定があると(図2-3)に示したCG-STACK 等に構文に関する情報をセットしてセマンテックス・ルーチンにコントロールをわたす。セマンティックス・ルーチンにおいては、構文に対応する処理をユーザの責任においておこなう。この場合、オブジェクト・プログラムをアセンブラ(FASP)で出力するとすれば、システム・ライブラリの中にあるコード発生ルーチンを使用できる。

コントローラあるいはセマンティックス・ルーチンにおいてエラーが検出された場合に、エラー・メッセージの出力、ソースプログラムの読みとばし(不必要な部分の)を行なうのがエラー処理ルーチンの役割である。エラーがあった場合には、コントローラはリカバリー可能な点までさかのぼって遷移テーブルのドライブを始める。

ソースプログラムを読み終わり、構文解析によって start 記号に還元されると、コントローラは制御を Terminator にわたす。Terminator はユーザOWNコードルーチンであって、最終処理をおこなう。

Terminator からコントローラを呼び出して、再び新しいプログラムのコンパイルも可能であり、また任意のユーザ・ルーチンにコントロールをわたしてもよい。

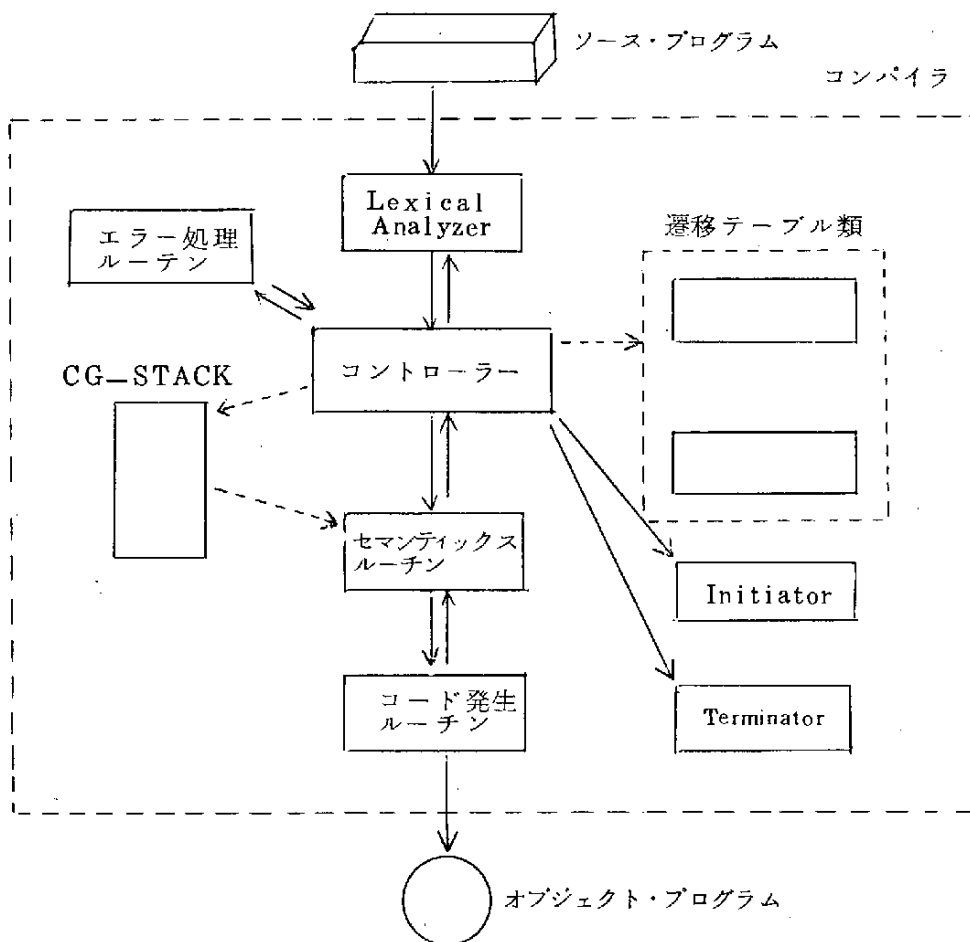
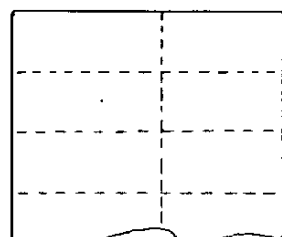


図 2-2 コンパイラ・ジェネレータによって作成されたコンパイラ概略

CG - STACK



属性

属性

numeric

ユーザーオプション

ただし, ADDRESS, POINTER, 数値, 文字 (4 文字), BIT の
スカラー変数に限る。

図 2-3

3 JPAC の言語仕様

JPAC は次の 2 つの定義部より構成される。

SYNTAX DEFINITION;

SEMANTICS DEFINITION;

以下に JPAC の文法を BNF で記述する。

$\langle \text{Alphabet} \rangle ::= A | B | C | D | \dots\dots\dots | X | Y | Z$

$\langle \text{Numeric} \rangle ::= 0 | 1 | 2 | 3 | \dots\dots\dots | 7 | 8 | 9$

$\langle \text{Alpha numeric} \rangle ::= \langle \text{Alphabet} \rangle | \langle \text{Numeric} \rangle |$

$\langle \text{Name} \rangle ::= \langle \text{Alphabet} \rangle | \langle \text{Name} \rangle | \langle \text{Alpha numeric} \rangle |$

$\langle \text{Name} \rangle _ \langle \text{Alpha numeric} \rangle$

(注) 名前は 256 字以内で、かつ最初の 12 文字は同じものであっては
いけない。

また、BNF 記述中次の 2 文字以上の記号の連なりは、右に書いた記号とみなす。

これ等は BNF 記述のための記号と terminal 記号を区別するためである。

$\nabla \nabla \rightarrow \nabla$

$\nabla \langle \rightarrow \langle$

$\nabla \rangle \rightarrow \rangle$

$\nabla | \rightarrow |$

$\nabla ::= \rightarrow ::$

3.1 SYNTAX DEFINITION

この定義部においてはコンパイラ・ジェネレータで作成しようとしている言語の文法を BNF で記述し、その規則が適用された時にどのセマンティックス・ルーチンにコントロールをわたすかを指定するものである。

SYNTAX DEFINITION は次の3つの節より構成される。

BASIC SYMBOL;

BASIC STATEMENT;

SYNTAX RULE;

① BASIC SYMBOL

文法を BNF で記述する場合 terminal 記号は文字そのものを、nonterminal 記号は記号名を $\langle \quad \rangle$ でくくってそれを示している。

例

$$\langle \text{NUMERIC} \rangle ::= 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9$$
$$\langle \text{NUMBER} \rangle ::= \langle \text{NUMERIC} \rangle | \langle \text{NUMBER} \rangle \langle \text{NUMERIC} \rangle$$

通常 $\langle \text{identifier} \rangle$, $\langle \text{NUMERIC} \rangle$, $\langle \text{NUMBER} \rangle$ 等の nonterminal 記号はプログラムにほぼ共通であり、SYNTAX 定義部において定義するよりは Lexical Analyzer を own-coding した方が処理も速く有意義なことが多い。しかしながら Lexical Analyzer を通してソースプログラムを入力する場合には、その最小単位は単語というレベルになるので単語を terminal 記号とみなす必要がある。

BASIC SYMBOL 節は Lexical Analyzer で解釈する単語はこれこれであるということを定義する部分である。

BASIC SYMBOL 節で定義された名前には定義された順に 1 から連番がふられる。これを BASIC SYMBOL 番号と呼び、Lexical Analyzer で Basic Symbol を解析した時に、この番号を呼んだルーチンに返さなければならない。

BASIC SYMBOL を定義しない場合にはコンパイラ・ジェネレータが

用意しているユーティリティ・ルーチンを使うことができるが、このルーチンはソースプログラムを1文字ずつとってくる役割のみを持っているため出来る上がるコンパイラのスピードはかなり遅くなる。

一般形

BASIC SYMBOL

<Basic Symbol definition>;

<Basic Symbol definition> ::= <空> | <Basic Symbol list>

<Basic Symbol list> ::= <Basic Symbol> |

<Basic Symbol list>, <Basic Symbol>

<Basic Symbol> ::= ▽<<Name>>▽

② BASIC STATEMENT

構文解析中にエラーが起った場合、すなわちソース・プログラムの中に文法に合わない文字列がでてきたときには、コンパイラはエラー・コンディションを発生する。

この場合に、コンパイラが処理を続行するためには、ソース・プログラムの一部を無視しなければならない。しかしながらソース・プログラムのどこからどこまでを無視するかということは非常に大きな問題である。

JPAC においては、この基準としてBasic Statementとして定義されたnonterminal 記号の直後の記号からユーザが指定する記号までを無視することとしている。

FORTRAN を例にとって考えてみよう。

<Program> ::= ⊢<Statement list> END ⊣

<Statement list> ::= <Statement> |

<Statement list><Statement>

<Statement> ::= <non executable> | <logical if> | <do> |

<sigma executable>

<logical if> ::= IF (<logical exp>) <sigma executable>

$\langle \text{sigma executable} \rangle ::= \langle \text{Assign} \rangle | \langle \text{GO TO} \rangle | \langle \text{arith if} \rangle | \dots\dots\dots$
 というように定義されている時に $\langle \text{statement} \rangle$ を BASIC STATEMENT
 と宣言し, Lexical Analyzer においてステートメントのエンド (開始行
 のみの時はその行の最後, 継続行がある時には最終継続行の最後) までを
 読みとばせるようにしておけば, エラーが起った場合にそのステートメン
 トのみを無視することが可能である。

③ SYNTAX RULE

この節は作成したいコンパイラの文法を BNF で記述し, その規則が適用
 された場合に制御をわたすべきセマンティックスルーチンを指定するもので
 ある。

この節で定義されている生成規則からは (図 2-2) で示されているもの
 中の遷移テーブル類が作成される。この遷移テーブル類は, コンパイラが動
 作する時にコントローラーが解読し生成規則が適用されると, CG-STACK
 の上から n 個の要素を生成規則の右辺として処理するように指定されたセ
 マンティックス・ルーチンにコントロールをわたす。セマンティックス・ル
 ーチンから帰ると, CG-STR に還元された nonterminal 記号をセットする,
 そして CG-STACK にスタックし再び遷移テーブル類の解読にはいる。

セマンティックス・ルーチンの指定がない場合には, コントローラーは
 [生成規則の右辺の項の数 - 1 項] を pop up して, スタックの最上位に還
 元された nonterminal 記号をセットして遷移テーブル類の解読をつづけ
 る。

一般形

$$\begin{aligned} \langle \text{Production} \rangle &::= \langle \text{Production rule} \rangle | \\ &\quad \langle \text{nonterminal} \rangle^{\nabla} ::= \langle \text{Production element list} \rangle \\ \langle \text{Production element list} \rangle &::= \langle \text{Production element} \rangle | \\ &\quad \langle \text{Production element list} \rangle | \langle \text{Production element} \rangle \\ \langle \text{Production element} \rangle &::= \langle \text{Right Part} \rangle : \langle \text{Semantics exec Part} \rangle \end{aligned}$$

$\langle \text{Semantics exec part} \rangle ::= \langle \text{Semantics call} \rangle$

$\langle \text{Semantics exec part} \rangle, \langle \text{Semantics call} \rangle$

$\langle \text{Semantics call} \rangle ::= \langle \text{Name} \rangle | \langle \text{Name} \rangle (\langle \text{Parameter list} \rangle)$

$\langle \text{Production rule} \rangle ::= \langle \text{nonterminal} \rangle \nabla ::= \langle \text{Right Part} \rangle$

$\langle \text{Right Part} \rangle ::= \langle \text{Symbol} \rangle | \langle \text{Right Part} \rangle \langle \text{Symbol} \rangle$

(注) $\langle \text{Semantics call} \rangle$ の中で使用する名前は 8 文字以内で、セマンティックス・ルーチン名でなければならない。 $\langle \text{Parameter list} \rangle$ はスタック名あるいは定数あるいは $\langle \text{Common Variable list} \rangle$ 中の変数でなければならない。

$\langle \text{Production List} \rangle ::= \langle \text{Production} \rangle ; |$

$\langle \text{Production List} \rangle \langle \text{Production} \rangle ;$

$\langle \text{Syntax rule} \rangle ::= \text{SYNTAX RULE} ; \langle \text{Production List} \rangle$

$\text{SYNTAX RULE END} ;$

(注) $\langle \text{Production list} \rangle$ の中で最初の $\langle \text{Production} \rangle$ は Start 記号の記述でなければならない。また、一つの nonterminal 記号に対しては一度しか $\langle \text{Production} \rangle$ を定義できない。

SYNTAX RULE は $\langle \text{Production} \rangle$; を必要な個数だけ並びたものであるということができ、 $\langle \text{Production} \rangle$ はおおざっぱに考えれば次の 2 つの形のうちのいずれかである。

$\langle E \rangle ::= \langle E \rangle + \langle T \rangle : \text{EXP} ;$

$\langle P \rangle ::= \langle \text{id} \rangle ;$

前者は、 $\langle E \rangle + \langle T \rangle$ という形が現われた時、EXP という名前のセマンティックス・ルーチンにコントロールをわたすことを指示している。セマンティックス・ルーチンの処理が終ると、 $\langle E \rangle + \langle T \rangle$ は新しく $\langle E \rangle$ になる。後者は $\langle \text{id} \rangle$ が現われた時 $\langle P \rangle$ に変えることのみを指定し、セマンティックス・ルーチンにはコントロールをわたさない。

$\langle \text{Semantics exec part} \rangle$ すなわちセマンティックス・ルーチンの実行指

定は，セマンティックス・ルーチン名及びコンパイラ・ジェネレータ標準ルーチンを(,)でくぎって任意の個数指定することができる。ここで指定されたものは，指定された順序にしたがって実行されてゆく。

一つの nonterminal 記号に対して 2 個以上の生成規則がある場合，たとえば，

$$\langle E \rangle ::= \langle E \rangle + \langle T \rangle : \text{EXP};$$

$$\langle E \rangle ::= \langle T \rangle;$$

という時には，これを(1)で区切って書かなければならない。すなわち

$$\langle E \rangle ::= \langle E \rangle + \langle T \rangle : \text{EXP} \mid \langle T \rangle;$$

という形で指定しなければならない。

例 3-1 生成規則が次の 8 個より成り立っているものとする。

表 3-1

```

S → | A |
A → identifier = E;
E → E + T
E → T
T → P ↑ T
T → P
P → ( E )
P → identifier

```

これを JPAC の言語仕様にしたがって，SYNTAX 定義部分のみを書けば，次のようになる。

表 3-2

```

SYNTAX  DEFINITION;
BASIC  SYMBOL;
<IDENTIFIER>;
SYNTAX  RULE;

```

```

<S> ::=  $\vdash$  <A>  $\vdash$  ;
<A> ::= <IDENTIFIER> = <E> : ASS ;
<E> ::= <E> + <T> : EXP | <T> ;
<T> ::= <P>  $\uparrow$  <T> : TERM | <P> ;
<P> ::= (<E>) | <IDENTIFIER> : PRIM ;
SYNTAX RULE END ;

```

3.2 SEMANTICS DEFINITION

SYNTAX DEFINITION 定義部においては，作成したいコンパイラ言語の文法の定義および，制御をわたすセマンティックス・ルーチンの指定をおこなった。すなわち，ある生成規則の適用がおこなわれた場合にセマンティックス・ルーチンの指定があれば，そのルーチンに一時的にコントロールがわたされる。

SEMANTICS DEFINITIONは，このセマンティックス・ルーチンのプログラムの集合とすることができる。

セマンティックス・ルーチンの作成のためには（図3-1）に示した概略を正しく理解する必要がある。

$\langle E \rangle ::= \langle E \rangle + \langle T \rangle : \text{EXP};$

という生成規則があって，EXP ルーチンにコントロールがわたった場合に，コントローラからの連絡情報は，CG-STACK にあり，次のような形をしている。

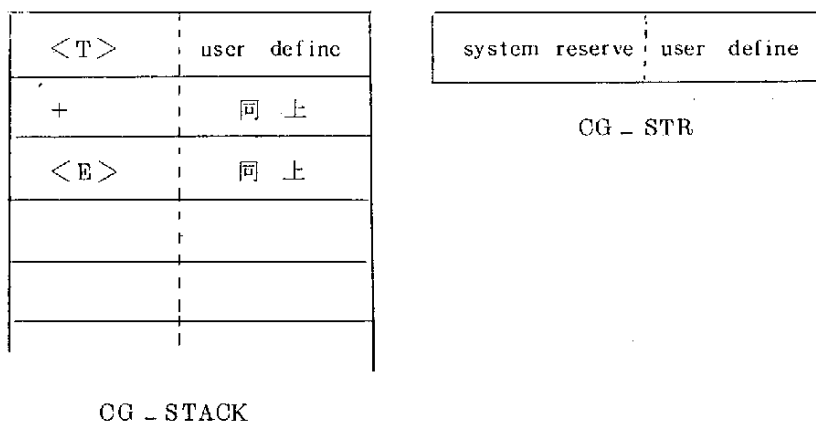


図 3 - 1

セマンティックス・ルーチンにおいては，必要な処理をおこない $\langle T \rangle$ ，+， $\langle E \rangle$ を pop up して，CG-STR の第2要素に還元された $\langle E \rangle$ に関する情報をセットしてコントローラに帰る。こうするとコントローラでは，CG-STACK の頭に push down して構文解析を続行する。

セマンティックスの記述方法はLSDの外部仕様にしたがうが、メイン・ルーチンとして使用してはいけない。

SEMANTICS DEFINITIONの一般形

SEMANTICS DEFINITION;

<Common Variable list>

SEMANTICS BODY;

<Semantics Procedure list>

//END

<Common Variable list> ::= <Common Variable> |

<Common Variable list> <Common Variable>

<Common Variable> ::= <DCL statement>

<DCL statement>はLSDの外部仕様書にあるDCLステートメントである。

<Common Variable list>として定義されたものについては<Semantics Procedure>の中で宣言してはいけない。すなわち<Common Variable list>は<Semantics Procedure list>中にあるすべてのProcedureの中にDCLステートメントとしてさし込みがおこなわれる。

<Semantics Procedure list> ::= <Semantics Procedure> |

<Semantics Procedure list> <Semantics Procedure>

<Semantics Procedure> ::= //LSD <External Procedure>

<External Procedure>はLSD外部仕様書にあるexternal procedureであり、このprocedure名がセマンティックス・ルーチン名に対応している。

例3-2

(表3-2)で定義した文法に関するものをJPACで記述してみる。セマンティックス記述は、コンパイラの言語仕様により、あるいはLexical Analyzerの機能によりいちじるしく異なるが、ここではLexical Analyzerはidentifierを認識した時にname tableに登録し、そのテーブルのポインタがスタック

クにのせられているものとする。

ここで示したのは、オブジェクト・コードをアセンブラ言語で直接出力する例であり、この他に種々の方法が考えられる。

***** JPAC (JIPDEC AUTOMATIC COMPILER GENERATOR) LIST *****

```
1          SYNTAX DEFINITION;
2          BASIC SYMBOL;
3          <IDENTIFIER>,<FACT>;
4          SYNTAX RULE;
5          <S> ::= "<A>";
6          <A> ::= <IDENTIFIER>=<E>;ASSI;
7          <E> ::= <E>+<T>;EXPI<T>;
8          <T> ::= <P><FACT><I>;IERML<P>;
9          <P> ::= (<E>)<IDENTIFIER>;PRIM;
10         SYNTAX RULE END;
```

***** JPAC (JIPDEC AUTOMATIC COMPILER GENERATOR) LIST *****

```
1          SEMANTICS DEFINITION;
2          DCL OPTBL(50) TABLE CONTROLLED(OPT);
3          * OPNAME CHAR(6);
4          * OPDISP;
5          * OPADD CHAR(4);
6          * OPATTR BIT;
7          DCL OPSTR STRUCTURE EXTERNAL;
8          * STRNAME CHAR(6);
9          * STRDISP;
10         * STRADD CHAR(4);
11         * STRATTR BIT;
12         DCL NAMETBL(200) TABLE CONTROLLED(TBP);
13         * TBNAME CHAR(6);
14         * TBATTR BIT;
15         * TBADD;
16         DCL CG_STACK(100) STACK;
17         * CG_SYS;
18         * CG_SPNT POINTER;
19         DCL CG_STR STRUCTURE EXTERNAL;
20         * CG_SSYS;
21         * CG_SSPT POINTER;
22         DCL WORKCNT EXTERNAL INIT(0);
```

***** JPAC (JIPDEC AUTOMATIC COMPILER GENERATOR) LIST *****

```

23      SEMANTICS BODY:
24      //LSD
25      PRIM : PROC:
26          POP CG_STACK INTO CG_STR;
27          /* MOVE PRIMARY INFORMATION TO OPTBL */
28          STRNAME=CG_SSPNT.TBNAME;
29          STRATTR=CG_SSPNT.TBATTR;
30          STRDISP=0;
31          STRADD='7';
32          SET_OPSTR INTO OPTBL;
33          CG_SSPNT=OPT;
34          RETURN;
35      END;
36      //LSD
37      ASS : PROC:
38          /* GENERATE      LA      <T> */
39          CALL PUT(' ','LA',CG_SSPNT.OPTBL);
40          POP 2;
41          /* GENERATE      STA      <IDENTIFIER> */
42          CALL PRIM;
43          CALL PUT(' ','STA',OPTBL);
44          RETURN;
45      END;
46      //LSD
47      EXP : PROC:
48          /* GENERATE      LA      <T> */
49          CALL PUT(' ','LA',CG_SSPNT.OPTBL);
50          POP 2;
51          /* GENERATE      A      <E> */
52          CALL PUT(' ','A',CG_SSPNT.OPTBL);
53          /* GENERATE      STA      <E> */
54          CALL WORKGET;
55          CALL PUT(' ','STA',OPTBL);
56          POP 1;
57          RETURN;
58      END;
59      //LSD
60      TERM : PROC:
61          /* GENERATE      SXJ,7      C.FACT */
62          STRNAME='C.FACT';
63          STRDISP=0;
64          STRADD='7';
65          STRATTR='B'0';
66          CALL PUT(' ','SXJ',OPSTR);
67          POP CG_STACK INTO CG_STR;
68          POP 1;
69          CALL PUT(' ','NOP',CG_SSPNT.OPTBL);
70          CALL PUT(' ','NOP',CG_SSPNT.OPTBL);
71          CALL WORKGET;
72          CALL PUT(' ','STA',OPTBL);

```

***** JPAC (JIPDEC AUTOMATIC COMPILER GENERATOR) LIST *****

```
73             POP 1;
74             RETURN;
75         END;
76     //LSD
77     WORKGET : PROC
78         /* WORK AREA OPERAND SET */
79         STRNAME='L.WK';
80         STRDISP=WORKCNT;
81         STRADD=' ';
82         STRATTR=B'0';
83         SET OPSTR INTO OPTBL;
84         WORKCNT=WORKCNT+1;
85         CG_SSPNT=OPT;
86         RETURN;
87     END;
88 //END
```


3.3 Lexical Analyzer

コンパイラ・ジェネレータを使ってコンパイラを作成する場合に，Basic Symbol の定義をすることができることは，すでに述べたとおりである。Basic Symbol を使用した時，ユーザは自分の責任でLexical Analyzer は（図 2-2）に示されているように，出来上ったコンパイラのソースプログラムを読み込み，単語の列に変換し，コントローラに情報を渡す役割を持っているのでユーザは次のような情報を指定された変数にセットするように作らなければならない。

なお Basic Symbol を使用しない場合には，コンパイラ・ジェネレータのユーティリティ・ルーチンとして用意してある Lexical Analyzer があるので，ユーザはこれを使えばよい。しかしながら，出来上ったコンパイラの大きさ，コンパイルスピードに関してはいちじるしく悪くなることもある。

Lexical Analyzer のルーチン名

CG-LEX

情報の受け渡し

CG-BF: Terminal Symbol あるいは Basic Symbol をセットする場所である。

属 性 : CHARACTER 256文字, EXTERNAL

CG-SORT: Terminal Symbol の場合には 0, Basic Symbol の場合には, Basic Symbol 番号を入れる。

属 性 : NUMERIC, EXTERNAL

CG-CC: CG-BFに Basic Symbol がはいつている場合に, CG-STACKの user define areaにセットしてほしいものを入れておく。

属 性 : 1 word, EXTERNAL

その他

LSDによって CG-LEXを作成する場合は問題はないが, アセンブ

ラで作成する場合はLSDとアセンブラのリンクの章をよく読んでつな
がるようにすること。

例2-3 FORTRAN Lexical Analyzer

前提

次の文字列のあとに空白あるいは特殊文字がくる場合は、これをリザーブ
・ワードとして、文あるいは行の識別のためにもちいる。

END

COMMON

DIMENSION

EQUIVALENCE

SUBROUTINE

FUNCTION

READ

WRITE

FORMAT

GO

DO

IF

CALL

RETURN

STOP

CONTINUE

以上のものはBasic Symbolとして宣言されており、その番号は END
= 1 から順で CONTINUE = 16 となっている。

これ以外に <identifier>, <Integer-number>, <real-number> がそ
れぞれ Basic Symbol 番号 17, 18, 19 で宣言されているものとする。

```

CG_LEX:      PROC;
DCL TAB(10) TABLE CONTROLLED(P);
      *NUM1 CHAR(1) INIT('0','1','2','3','4','5','6','7',
      '8','9'),
      *NUM2 INIT(0,1,2,3,4,5,6,7,8,9);
DCL BASIC_TABLE(16) TABLE CONTROLLED(P);
      *BASIC_SY CHAR(12) INIT('END
      'COMMON
      'DIMENSION
      'EQUIVALENCE
      'SUBROUTINE
      'FUNCTION
      'READ
      'WRITE
      'FORMAT
      'GO
      'DC
      'IF
      'CALL
      'RETURN
      'STOP
      'CONTINUE
      ');
      *BASIC_NO INIT(1,2,3,4,5,6,7,8,9,10,11,12,13,14,15,
      16);
DCL CODE CHAR(1);
DCL CG_BF CHAR(256) EXTERNAL;
DCL CG_SORT EXTERNAL;
DCL CG_CC EXTERNAL;
DCL LABEL_SW BIT(1) EXTERNAL;
DCL KEY_BF CHAR(12), OR_CHAR CHAR(1), NAME_SW BIT(1),
      OR_SW BIT(1), TYPE_SW BIT(1), DIG_SW BIT(1),
      EXP_SW BIT(1), ASTER_SW BIT(1);
DCL BLANK_SW BIT(1);
/* START LEXICAL-ANALIZER */
CG_BF = ' ';
CG_SORT = 0;
CG_CC = 0;
KEY_BF = ' ';
RESETB NAME_SW.1;
CODE = OR_CHAR;
IF OR_SW.1 IS OFF THEN GO TO LEX-7;
RESETB OR_SW.1;
GO TO LEX-9;
LEX-7:      CALL CODE-GET(CODE);
IF LABEL_SW.1 IS ON THEN RETURN;
LEX-9:      IF CODE = ' ' THEN GO TO LEX-7;
      I = 1;
      IF CODE IS ALPHA THEN GO TO LEX-39;
      IF CODE IS NUMERIC THEN GO TO LEX-13;
      IF CODE = '.' THEN GO TO LEX-54;

```

```

LEX_13:      RESETB TYPE_SW.1;
             IF CODE = ':' THEN SETB TYPE_SW.1;
LEX_15:      SUBSTR(CG_BF.1) = CODE;
             I = I + 1;
             CALL CODE_GET(CODE);
             IF CODE IS NUMERIC THEN GO TO LEX_15;
             IF CODE IS ALPHA THEN GO TO LEX_26;
             IF CODE = '.' THEN GO TO LEX_24;
             SETB OR_SW.1;
             OR_CHAR = CODE;
             IF TYPE_SW.1 IS ON THEN GO TO LEX_23;
             CG_SORT = 18;
             RETURN;
LEX_23:      CG_SORT = 19;
             RETURN;
LEX_24:      IF TYPE_SW.1 IS ON THEN CALL ERROR(10);
             SETB TYPE_SW.1;
             GO TO LEX_15;
LEX_26:      IF CODE = 'E' THEN CALL ERROR(11);
             SETB TYPE_SW.1;
             RESETB DIG_SW.1;
             RESETB EXP_SW.1;
LEX_29:      SUBSTR(CG_BF.1) = CODE;
             I = I + 1;
             CALL CODE_GET(CODE);
             IF CODE IS ALPHA THEN CALL ERROR(12);
             IF CODE IS NUMERIC THEN GO TO LEX_38;
             IF CODE = '+', '-', THEN GO TO LEX_36;
LEX_35:      IF DIG_SW.1 IS OFF THEN CALL ERROR(13);
             SETB OR_SW.1;
             OR_CHAR = CODE;
             GO TO LEX_23;
LEX_36:      IF EXP_SW.1 IS OFF THEN GO TO LEX_37;
             GO TO LEX_35;
LEX_37:      SETB EXP_SW.1;
             GO TO LEX_29;
LEX_38:      SETB DIG_SW.1;
             SETB EXP_SW.1;
             GO TO LEX_29;
LEX_39:      SUBSTR(CG_BF.1) = CODE;
             IF NAME_SW.1 IS ON THEN GO TO LEX_50;
             KEY_BF = SUBSTR(CG_BF.1, I);
             SEARCH BASIC_SY      KEYED KEY_BF IF NOT GO TO LEX_50;
             I = I + 1;
             CALL CODE_GET(CODE);
             IF CODE IS ALPHA THEN GO TO LEX_49;
             IF CODE IS NUMERIC THEN GO TO LEX_49;
             CG_SORT = P.BASIC_NO;
             CG_CC = I-1;
LEX_48:      SETB OR_SW.1;

```

```

OR_CHAR = CODE;
LEX_63: RETURN;
LEX_49: SETB NAME_SW.1;
GO TO LEX_39;
LEX_50: I = I + 1;
LEX_51: CALL CODE_GET(CODE);
IF CODE = ' ' THEN GO TO LEX_53;
SETB BLANK_SW.1;
GO TO LEX_51;
LEX_53: RESETB BLANK_SW.1;
IF CODE IS ALPHANUMERIC THEN GO TO LEX_39;
GO TO LEX_48;
LEX_54: IF CODE IS ALPHANUMERIC THEN CALL ERROR(14);
IF CODE = '*' THEN SETB ASTER_SW.1;
LEX_56: SUBSTR(CG-BF.1) = CODE;
I = I + 1;
IF ASTER_SW.1 IS OFF THEN RETURN;
CALL CODE_GET(CODE);
IF CODE = '*' THEN GO TO LEX_62;
SETB OR_SW.1;
OR_CHAR = CODE;
RETURN;
LEX_62: RESETB ASTER_SW.1;
GO TO LEX_56;
/* END LEXICAL-ANALIZER */
/* START CODE-GET */
CODE_GET: PROC(G_CODE);
DCL INPBUF CHAR(80),COUNT INIT(72),SINO;
DCL G_CODE CHAR(1);
DCL JUMP_SW BIT(1);
RESETS LABEL_SW.1;
IF JUMP_SW.1 IS ON THEN GO TO GET_11;
IF COUNT < 72 THEN GO TO GET_8;
GET_2: READ SYSIN INTO INPBUF;
AT END GO TO GET_20;
WRITE SYSPT FROM INPBUF;
IF SUBSTR(INPBUF.1) = 'C' THEN GO TO GET_2;
IF SUBSTR(INPBUF.6) = ' ','0' THEN GO TO GET_10;
GET_7: COUNT = 7;
GET_8: G_CODE = SUBSTR(INPBUF,COUNT);
COUNT = COUNT + 1;
RETURN;
GET_10: IF BLANK_SW.1 IS OFF THEN GO TO GET_11;
SETB JUMP_SW.1;
GO TO LEX_63;
GET_11: RESETB JUMP_SW.1;
SINO = 0;
DO I = 1 TO 5;
G_CODE = SUBSTR(INPBUF,I);
SEARCH NUN1 KEYED G_CODE IF NOT GO TO GET_15;

```

```

        STNO = STNO * 10 + NUM2;
        GO TO GET_16;
GET_15:   IF G_CODE = ' ' THEN CALL ERROR(1);
GET_16:   END;
          IF STNO = 0 THEN GO TO GET_7;
          CG_BF = STNO;
          CG_SORT = 18;
          CG_CC = 0;
          COUNT = 7;
          SETB LABEL_SW.1;
          RETURN;
GET_20:   G_CODE = '*';
          RETURN;
          END;
/*      END      CODE-GET      */
ERROR:    PROC(NO);
          DCL MESSAGE CHAR(25) INIT(' CG-LEX ERROR      NO.      ');
          SUBSTR(MESSAGE,21,2) = NO;
          WRITE SYSPT FROM MESSAGE;
          STOP;
          END;
          END;

```

3.4 エラー処理ルーチン

JPAC で作り出されたコンパイラにおいて、コントローラは構文解析中にエラーを見つけ出した場合に、ソース・プログラム中のどこでエラーが起きたかを表示することができるが、そのエラーがどのような種類のものであるかをユーザに分り易く出力することができない。また、ソース・プログラムのどこまでを無視してエラー・リカバリーするのが適当であるかを判断することができない。

このためエラーが起きた場合に、適切なエラー・メッセージの出力とエラー・リカバリーのために、ユーザはエラー処理ルーチンをコンパイラの中に組み込むことが可能になっている。

エラー処理ルーチンは、BASIC STATEMENTを指定した場合には、ユーザがOWNコーディングしなければならないものである。

コントローラがソース・プログラムの構文解析中にエラーを検出した場合には、エラー個所をSYSPRTに出力したのちこのルーチンを呼ぶようになっている。セマンティックス・ルーチンの中でエラーが検出された場合には、ユーザは直接このルーチンを呼ぶようにしなければならない。

このルーチンではエラー・メッセージの出力とソース・プログラムの読みとばしをしてリターンすればよい。こうするとコントローラはBASIC STATEMENT に対応するところまでCG-STACKをpop upして、ソース・プログラムの構文解析を続行する。

BASIC STATEMENT の指定をしなければ、このルーチンは用意する必要はない。しかしながらそのような場合には、コントローラがエラーを検出するとエラー個所を出力したのち Terminator にコントロールをわたすことになり、エラー個所以降は全くチェックをしないことになる。

4. STAGE 内部処理

以下に STAGE の内部処理について述べるが、説明に際して Franklin Lewis De Remer の PRACTICAL TRANSLATORS FOR LR(k) LANGUAGES を一部引用した。

4.1 処理概略

STAGE は SYNTAX DEFINITION から遷移図を作成し FASP ソースの形で遷移テーブル類を出力するものである。

STAGE の処理の流れは (図 4-1) に示されるとおりである。

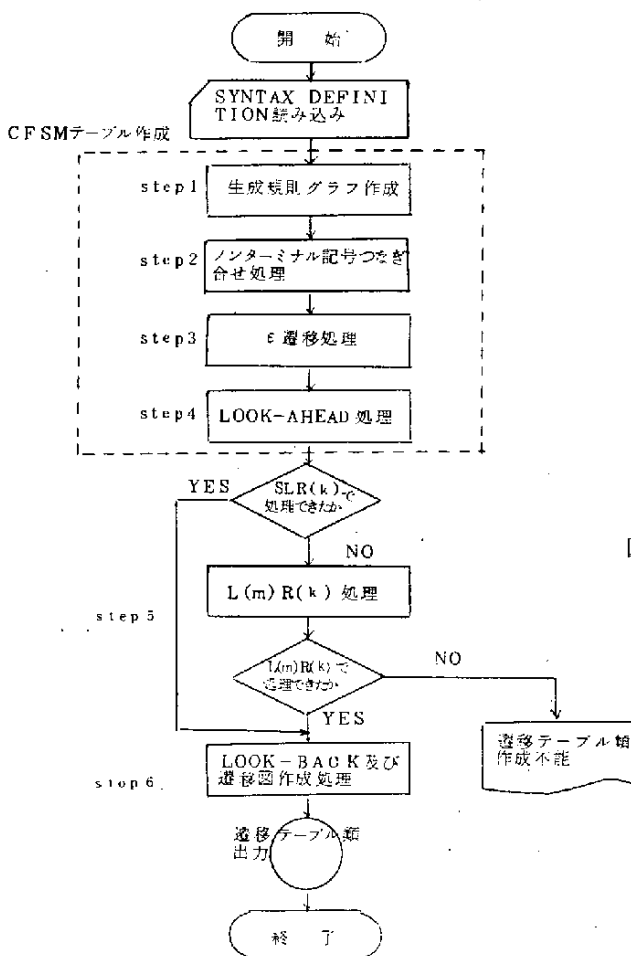


図 4-1 STAGE の処理の流れ

以下に(表 4-1) で示す SYNTAX DEFINITION が与えられた場合
を例として STAGE の処理を説明する。

表 4-1 SYNTAX DEFINITION の例

```
SYNTAX DEFINITION ;  
BASIC SYMBOL ;  
<IDENTIFIER> ;  
SYNTAX RULE ;  
<S> ::=  $\vdash$  <E>  $\vdash$  : EXEC 0 ;  
<E> ::= <E> + <T> : EXEC 1 | <T> : EXEC 2 ;  
<T> ::= <P>  $\uparrow$  <T> : EXEC 3 | <P> : EXEC 4 ;  
<P> ::= <IDENTIFIER> | (<E>) : EXEC 5 ;  
SYNTAX RULE END ;
```

4.2 生成規則グラフ作成

この処理は(図4-1)の step 1 に対応するものであり,与えられた生成規則(SYNTAX RULE 節で与えられる)から node と arrow で表現した生成規則グラフを作成する。(arrow は遷移を表わし, node は状態を表わす)生成規則グラフを作成する手順は以下の A, B の2つの部分に分かれる。

A. 生成規則から Characteristic String を作成する。

(表4-1)で示された例の場合の Characteristic String は(表4-2)で示される。

表 4-2 Characteristic String

0 :	0 :	$S' \rightarrow \vdash E \vdash$	$\#_0$ *-1)
	1 :	$S' \rightarrow \vdash E'$	
1 :	2 :	$E' \rightarrow E + T$	$\#_1$
	3 :	$E' \rightarrow E + T'$	
	4 :	$E' \rightarrow E'$	
2 :	5 :	$E' \rightarrow T$	$\#_2$
	6 :	$E' \rightarrow T'$	
3 :	7 :	$T' \rightarrow P \uparrow T$	$\#_3$
	8 :	$T' \rightarrow P \uparrow T'$	
	9 :	$T' \rightarrow P'$	
4 :	10 :	$T' \rightarrow P$	$\#_4$
	11 :	$T' \rightarrow P'$	
5 :	12 :	$P' \rightarrow i$ *-2)	$\#_5$
6 :	13 :	$P' \rightarrow (E)$	$\#_6$
	14 :	$P' \rightarrow (E')$	
		↑ Characteristic 番号	
		↑ 生成規則番号	

*-1) $\#_n$ は n 番目の生成規則に対応する特殊記号であって,この記号にであうとその生成規則に対応するセマンティックス・ルーチンの実行をする。

*-2) i : <IDENTIFIER> を意味する。

B. Aで作成された Characteristic String から生成規則グラフを作成する。

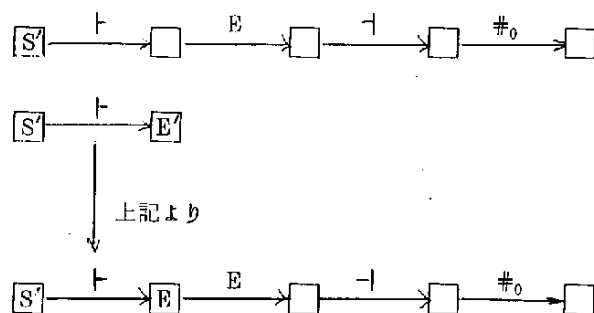


図 4-2 生成規則グラフの作成例

アルゴリズム

- ① Characteristic String 上ではじめてあらわれた記号，つまりダッシュのついている記号（例：S'，E'）は node とする。
- ② 生成規則の各 terminal，nonterminal 記号はすべて arrow に対応する。
- ③ 各 arrow は必ず node から出て node に終るものとする。
- ④ 各 Characteristic String のうち，同一記号から始るものは1つにまとめる。

（表 4-2）では Characteristic 番号の (0,1), (2,3,4,5,6) 等である。

- ⑤ node から node への arrow になんら該当する遷移記号が無いとき，その遷移記号として ϵ (無) 記号を割り当てる。

上記のアルゴリズムによって（表 4-2）の Characteristic String に対する生成規則グラフを作成すると（図 4-3）になる。

次に（図 4-3），（図 4-4）を例にして記号及び言葉の定義を行なう。

node 記号：（図 4-3）を見るといくつかの node はその中にダッシュのついている記号を持つ，この記号を node 記号という。

（図 4-3）では S'，E'，T' etc ...

SNODE : (図4-3)を見てもわかるように、生成規則グラフはいくつかのグループに分かれている(各グループには生成規則グラフ番号が1つ対応している)そして各グループの最初のnodeをSNODEという。

(図4-3)では生成規則グラフ番号が0のSNODEは S' のnode記号を持つnodeである。

N_i : i という状態にあるnodeである。

(図4-4)では $N_{E'}$ はnode 1を示す。

arrowの : 任意のarrowが N_i から出て N_j に終るとき、 N_i, N_j をそれぞれSTART node及びarrowのEND nodeという。

(図4-4)ではarrow 3のSTART nodeはnode 1でありEND nodeはnode 4である。

ARROW{S} : Sという遷移記号を持つarrowである。

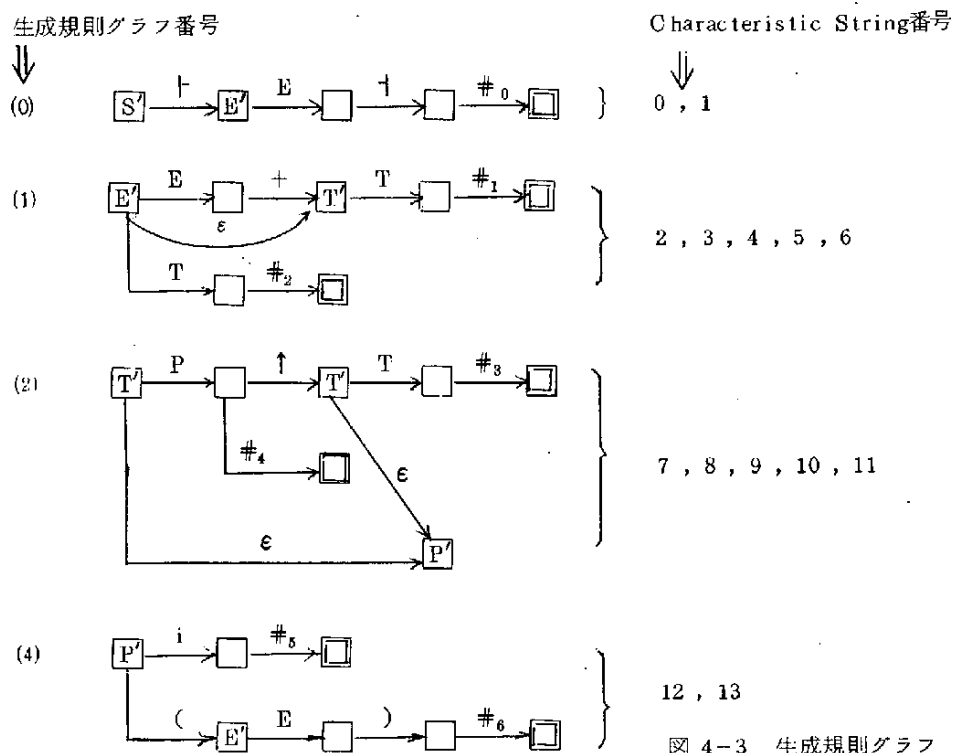


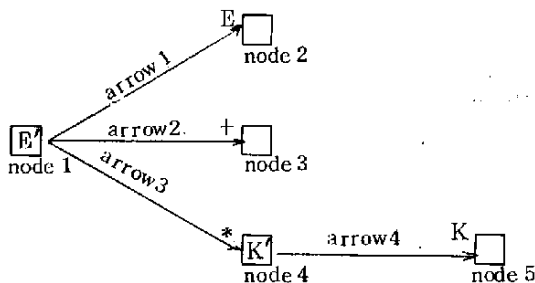
図4-3 生成規則グラフ

Niの子 arrow: Ni を START node とする arrow の集合である。

Niの子 node: Ni の子 arrow の END node の集合である。

Niの子孫 : Ni から arrow 及び node を通って到達できる arrow 及び node の集合をいう。

(図 4-4) では NE' (node 1) の子孫は arrow1~arrow4, 及び node 2~node 5 で, Nk' の子孫は arrow 4 と node 5 である。



E, +, *, K はおのおの arrow1, arrow2, arrow3, arrow4 の遷移記号である。

図 4-4

Niの子 arrow{S}: Ni の子 arrow のうち遷移記号として S をもつものである。

Niの子 node{S}: Ni の子 node のうち Ni の子 arrow {S} を経由して到達する node。

#n の帰納記号 : #n 記号により reduce される nonterminal 記号。

(表 4-1) における #n の帰納記号は E である。

#n arrow : 遷移記号として #n を持つ arrow で, arrow {#n} として参照されることもある。

4.3 Nonterminal 記号つなぎあわせ処理

この処理は(図4-1)のStep2に対応するものであり nonterminal 記号を結合してゆき, CFSM の母型となるグラフを作成する。(以下の説明ではこのグラフを母型グラフと呼ぶ)

結合処理アルゴリズム

- ① スタート・ステータスを表す node を母型グラフ上に作成する。
(例4-1)の①参照)
- ② 母型グラフ上の node 記号を持つ node で未処理のものがあれば③へ、なければ ϵ 遷移処理へ行く。
- ③ ②で見つけた未処理の node と同一の node 記号を SNODE としてもつ生成規則グラフを見つける。そして SNODE の子孫を未処理 node の子孫として追加した後②へ戻る。

結合処理の手続きは上記のように簡単なものだが、作成する母型グラフの冗長を無くすため結合処理アルゴリズム③は以下の手続きから構成されている。

- 3.1 未処理 node の子 arrow の遷移記号の集合を SUM 1 とする。SNODE の子 arrow の遷移記号の集合を SUM 2 とする。(但し ϵ 遷移のときは、その arrow の END node の node 記号を仮の遷移記号とする)
- 3.2 SUM 1, SUM 2 に含まれている記号を以下の3種に分類する。
 - 3.2.1 SUM 1 にあって SUM 2 にないもの、この集合を SUM 3 A とする。
 - 3.2.2 SUM 2 にあって SUM 1 にないもの、この集合を SUM 3 B とする。
 - 3.2.3 SUM 1, SUM 2 とともに含まれているが、その属性が異なるもの、この集合を SUM 3 C とする。
- 3.3 3.2 で分類した各集合について 未処理 node への結合処理を行う。
 - 3.3.1 $SM \in \text{SUM 3 A}$ なる記号 SM を遷移記号とする arrow は、すでに母型グラフ上に作成済みなので考慮の必要はない。
 - 3.3.2 SUM 3 B に含まれる記号については SNODE の子孫(この集合を SUM 4 とする)を母型グラフ上に複製し、その未処理 node の子孫とする。

但し該当生成規則グラフの結合が過去にされている場合は，作成すべき集合 SUM 4 はすでに母型グラフ上に作成されている。それゆえこの場合は，未処理 node から集合 SUM 4 を子孫とするような arrow を母型グラフ上に作成するだけでよい。

また SNODE，未処理 node が一時的であるとき^{*-1)} 集合 SUM 4 はすでに母型グラフ上に作成されている可能性がある。ゆえに作成済みか否かを調べ，前者の場合は未処理 node から集合 SUM 4 を子孫とするような arrow を母型グラフに作成するだけでよい。また後者の場合は集合 SUM 4 を未処理 node の子孫として母型グラフ上に作成する。

- 3.3.3 SUM 3 については未処理 node，SNODE ともに $SM \in \text{SUM}3C$ となる SM に対して子 node $\{SM\}$ を求め，それぞれ新しい未処理 node SNODE (この新しい未処理 node，SNODE は一時的なものなので一時的未処理 node，一時的 SNODE という) として 3.1 へ戻る。
- この作業は一時的未処理 node が無くなるまで繰り返される。

*-1) 3.3.3 参照

例 4-1

(表 4-1) で与えられた生成規則から，CFSM の母型となるグラフを作る手順を結合アルゴリズムの③に重点をおいて例示する。

- (1) スタート・ステータスに該当する node を持つてくる。

S'

図 4-5 (1) 終了時の母型グラフ

- (2) (図 4-5) の NS' に対して生成規則グラフを結合する。この場合の生成規則グラフは(図 4-3)の生成規則グラフ番号(0)であらわされる。

3.1 より

$$\text{SUM}1 = \{\phi\}$$

$$\text{SUM } 2 = \{ | \}$$

3.2 より

$$\text{SUM } 3 = \{ | \} \quad \text{で} \quad \text{SUM } 3A = \{ \phi \}, \text{SUM } 3B = \{ | \}$$

$$\text{SUM } 3C = \{ \phi \}$$

3.3.2 によって (図 4-6) ができる

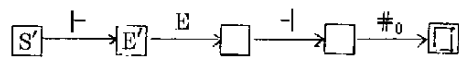


図 4-6 (2) 終了時の母型グラフ

- (3) (図 4-6) を見ると E' という node シンボルを持つ node が新しく未処理 node として出現しているこれに対して生成規則グラフを結合する場合の生成規則グラフは生成規則グラフ番号(1)であらわされる。

3.1 より

$$\text{SUM } 1 = \{ E \}$$

$$\text{SUM } 2 = \{ E, T', T \}$$

3.2 より

$$\text{SUM } 3A = \{ \phi \}, \text{SUM } 3B = \{ T', T \}, \text{SUM } 3C = \{ E \} \text{ となる。}$$

- (4) $\text{SUM } 3B$ に対して 3.3.2 を適用すると (図 4-7) となる。

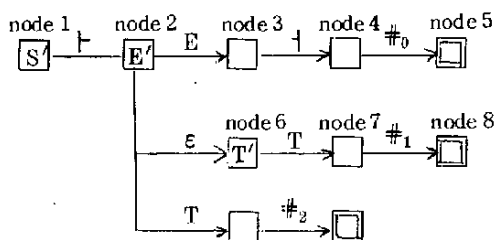


図 4-7 (4) 終了時の母型グラフ

- (5) 次に $\text{SUM } 3C$ に対して 3.3.3 を適用すると新しい一時的な未処理 node は node 3 で新しい一時的 SNode は生成規則グラフ番号(1)であらわされるグラフの SNode の子 node $\{ E \}$ である。

3.1 より

$$\text{SUM } 1 = \{-\}$$

$$\text{SUM } 2 = \{+\}$$

3.2 より

$\text{SUM } 3A = \{-\}$, $\text{SUM } 3B = \{+\}$, $\text{SUM } 3C = \{\phi\}$ となる。

SUM 3B に対して 3.3.2 を適用する, この場合 SNODE, 未処理 node は一時的ゆえ, node 3 (未処理 node) の子 node $\{+\}$ は母型グラフ上に作成済みの可能性がある。それゆえ, 調べてみると node 2 の子 node $\{\epsilon\}$ すなわち node 6 として作成済みである ((図 3-7) 参照), それゆえ node 3 から node 6 へ遷移記号 $+$ を持つ arrow を母型グラフへ追加する。結果 (図 3-8) となる。

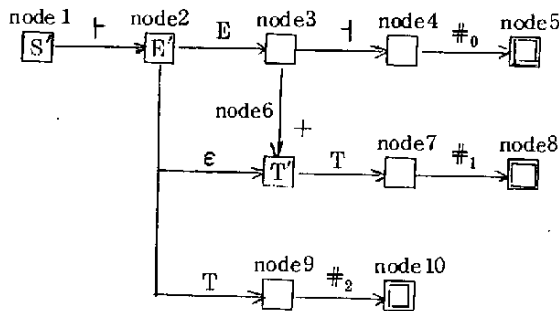


図 4-8 (5) 終了時の母型グラフ

(6) (図 4-8) の T' に対して生成規則グラフ番号(2)のグラフを結合する。

3.1 より

$$\text{SUM } 1 = \{T\}$$

$$\text{SUM } 2 = \{P, P'\}$$

3.2 より

$$\text{SUM } 3A = \{T\}, \text{SUM } 3B = \{P, P'\}, \text{SUM } 3C = \{\phi\}$$

となる。SUM 3B に対して 3.3.2 を適用すると (図 4-9) になる。

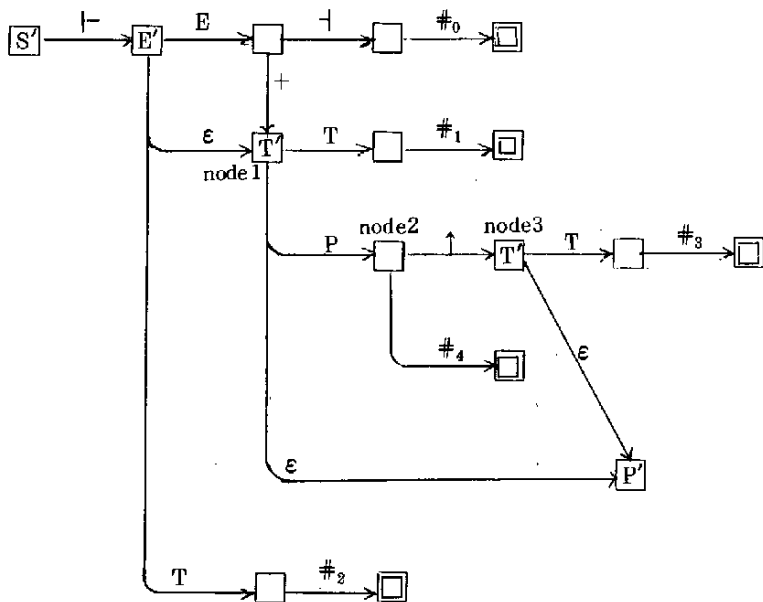


図 4-9 (6) 終了時の母型グラフ

(7) (図 4-9) をみると node シンボル T' と P' を持つ node が新しく未処理 node として示されている。まず T' に生成規則グラフ番号(2)のグラフを適用

3.1 より

$$\text{SUM } 1 = \{T, P\}$$

$$\text{SUM } 2 = \{P, P'\}$$

3.2 より

$$\text{SUM } 3A = \{T\}, \text{SUM } 3B = \{P\}, \text{SUM } 3C = \{P'\} \text{ となる。}$$

SUM 3B に 3.3.2 を適用する, しかしこのとき生成規則グラフ番号(2)のグラフは(6)で結合済みなので node 3 の子 arrow $\{P\}$ の END node は node 2 になる。結果(図 4-10) になる。

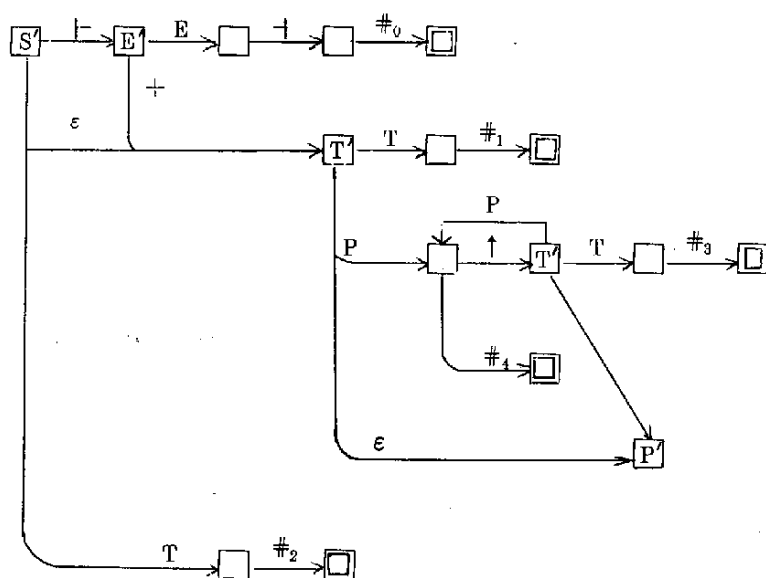


図 4-10 (7) 終了時の母型グラフ

同様の手続きを繰り返すと最終的に (図 4-11) ができる。

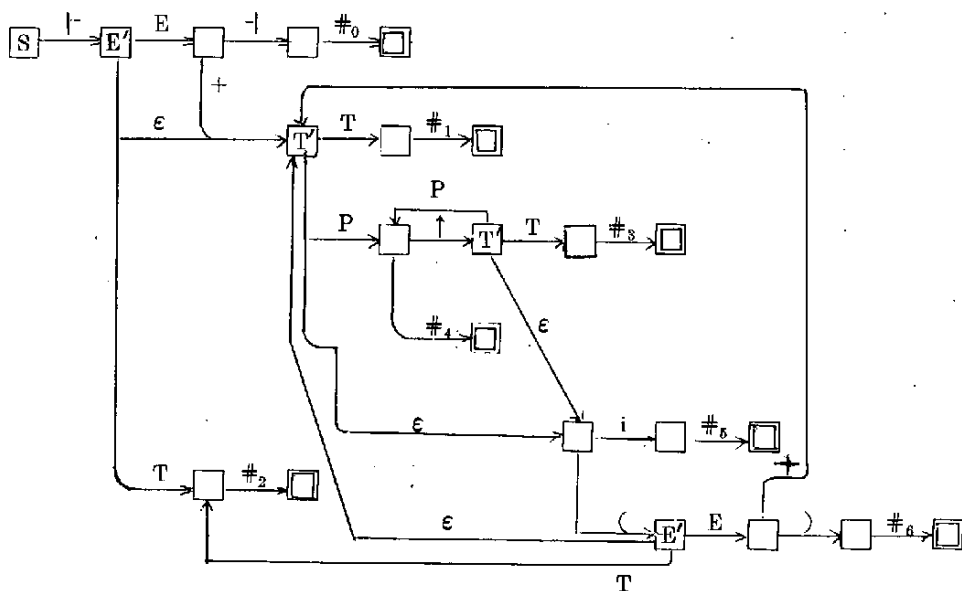


図 4-11 最終母型グラフ

4.4 ϵ 遷移処理

この処理は(図4-1)の step 3 に対応する。

ϵ による遷移を, 次の node からの遷移におきかえる。

そのアルゴリズムを以下にのべる。

ある node N_i があつてその node から ϵ 遷移があるとき, つまり N_i の子 arrow $\{\epsilon\}$ があるとき, 子 arrow $\{\epsilon\}$ の END node N_j の子孫を子 arrow $\{\epsilon\}$ のかわりに N_i に結合する。

その場合 N_i からの遷移と N_j からの遷移に同一遷移があるとき, N_i の方の遷移を優先する。そしてこの操作を全ての ϵ 遷移が無くなるまで繰り返す。

例 4-2

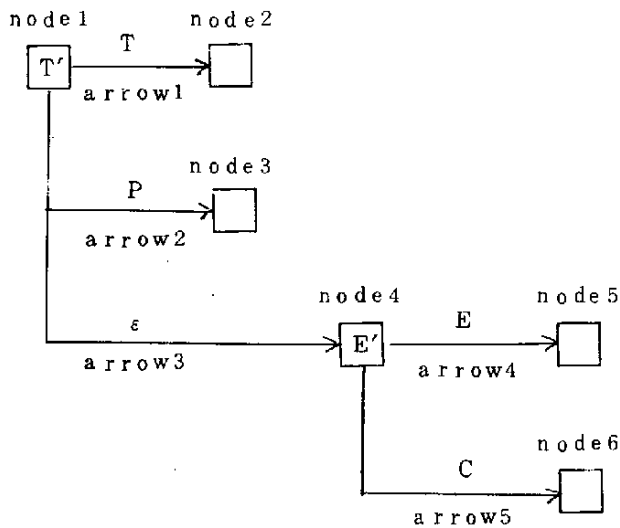


図 4-12 ϵ 遷移処理前

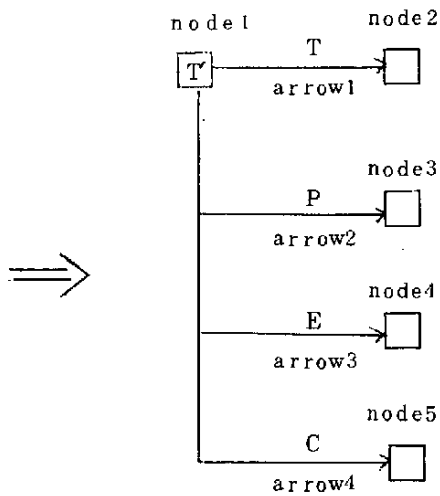


図4-13 ϵ 遷移処理後

(図4-12)を例にとって上記手続きを説明する。node1からは{T, P, ϵ }の遷移がある。ゆえ上記手続きにより ϵ 遷移に該当する。arrow3を取り去りarrow3のEND nodeであるnode4の子孫つまりarrow4とarrow5をnode1に結合する。結果(図4-13)となる。

例 4-3

(例4-1) でできた母型グラフ (図4-11) に対して上記ε処理を行ってできたものを (図4-14) に示す。これがCF SMグラフである。

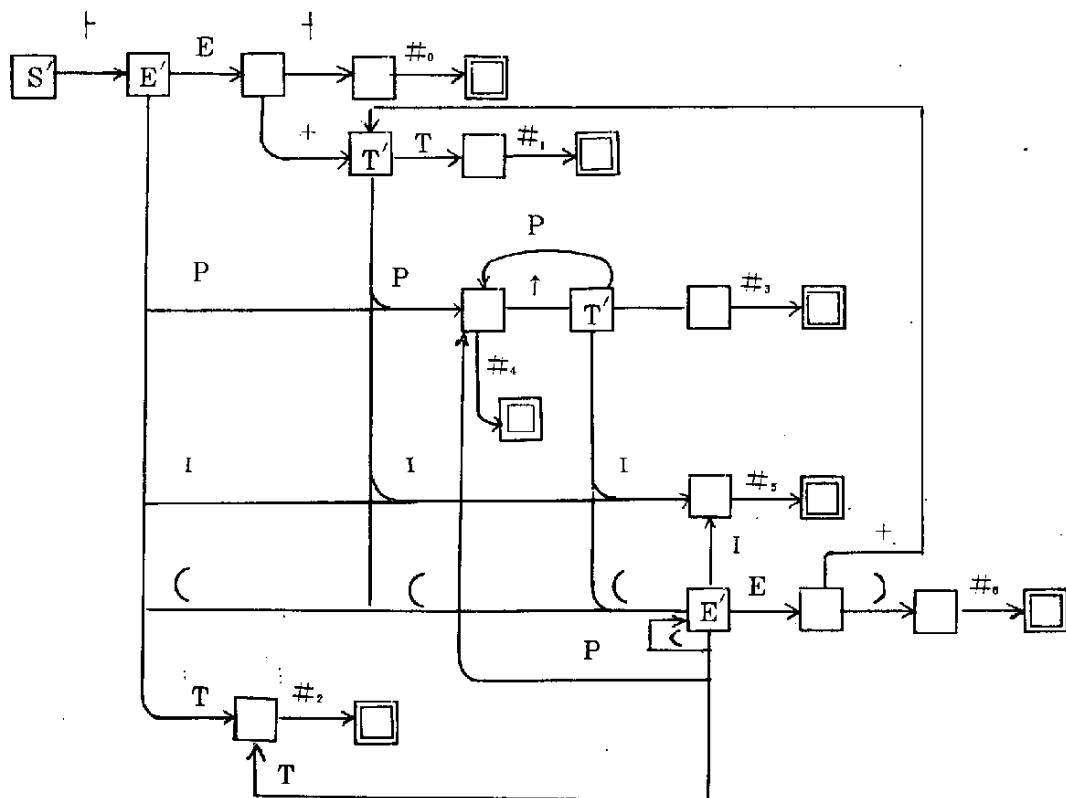


図4-14 (表4-1)に示す文法に対するCFSMグラフ

4.5 LOOK-AHEAD 処理

この処理は(図4-1)のStep4に対応し, inadequate状態に対処するものである。

a. look-ahead セットを求める。

(図4-14)をみると同一nodeから#n arrow と他の arrow がでている場合があるが, この様な状態を inadequate 状態という。

① inadequate状態がなければ遷移図作成処理へ行く。

inadequate状態がある場合には該当する#n の帰納記号をAとすると $F_T^K(A)$ と他のarrowの遷移記号と一致しないかどうか調べる。一致していれば②へ行く。一致していなければ $K=K+1$ として再び①へ行く。この場合 $K=2$ までをおこなう。それでも駄目な場合 $L(n)R(k)$ 処理へ行く。

② #n arrowに関して $F_T^K(A)$ をlook-ahead セットと称し, その状態に達した時, K個の記号を先読みして, 一致するものがあれば#nへ遷移するようにする。

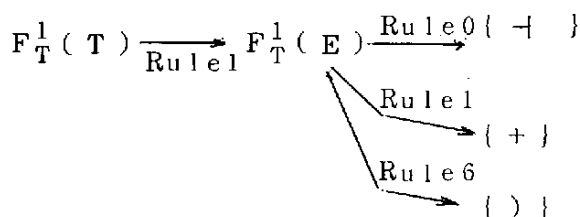
上記の $F_T^K(A)$ で求めた記号列がlook-ahead セットである。

例4-4

例として(図4-14)の#4の所のlook-ahead セットを求めてみる。

#4は(表4-1)によるとnonterminal記号Tにreduceされる記号であるので

$F_T^1(T)$ を求める



#4のlook-aheadセットは $\{ \neg, +,) \}$ となり(図4-14)は

(図4-15)となる。

b. $F_T^K(T)$ と $H_T^K(A)$

$F_T^K(T)$ は、記号 T のあとにつづく K 個の terminal 記号の列の集合を意味している。

$H_T^K(A)$ は nonterminal 記号 A が terminal 記号になったとき、その terminal 記号とそのあとに続く terminal 記号列の始めの K 個の記号列の集合を意味している。

T が $V_N \sqcup V_T$,

A, B, C, D, S が V_N (S は start 記号),

α, β, r が $V_T^* \sqcup \phi$ (ϕ は空),

ρ が V^* のとき

$$F_T^K(A) = \{ (K) : \beta \mid S \xrightarrow{*} \rho A \beta \}$$

$$H_T^K(A) = \{ (K) : \alpha \beta \mid S \rightarrow \rho A \beta, A \xrightarrow{*} \alpha \}$$

とする。

$F_T^K(T)$ の計算

$B \rightarrow \rho T r$ のように右辺に A を含む生成規則をとりだす。

$$r = \phi \text{ であれば } F_T^K(B) \in F_T^K(T)$$

$r \neq \phi$ であるとき

$r = \alpha$ であれば

$$|\alpha| \geq K \text{ ならば } (k) : \alpha \in F_T^K(T)$$

$$|\alpha| < K \text{ ならば } \alpha F_T^{K-|\alpha|}(B) \in F_T^K(T)$$

$r = \alpha C$ であれば

$$|\alpha| \geq K \text{ ならば } (k) : \alpha \in F_T^K(T)$$

$$|\alpha| < K \text{ ならば } \alpha H_T^{K-|\alpha|}(C) \in F_T^K(T)$$

$H_T^K(A)$ の計算

$A \rightarrow r C \rho$ のように左辺に A を含む生成規則をとりだす。

$A \rightarrow r$ のとき

$$|r| \geq \text{ならば} \quad (k) : r \in H_T^K(A)$$

$$|r| < \text{ならば} \quad r F_T^{K-|r|}(A) \in H_T^K(A)$$

$A \rightarrow r C \rho$ のとき

$$|r| \geq K \text{ならば} \quad (k) : r \in H_T^K(A)$$

$$|r| < K \text{ならば} \quad r H_T^{K-|r|}(C) \in H_T^K(A)$$

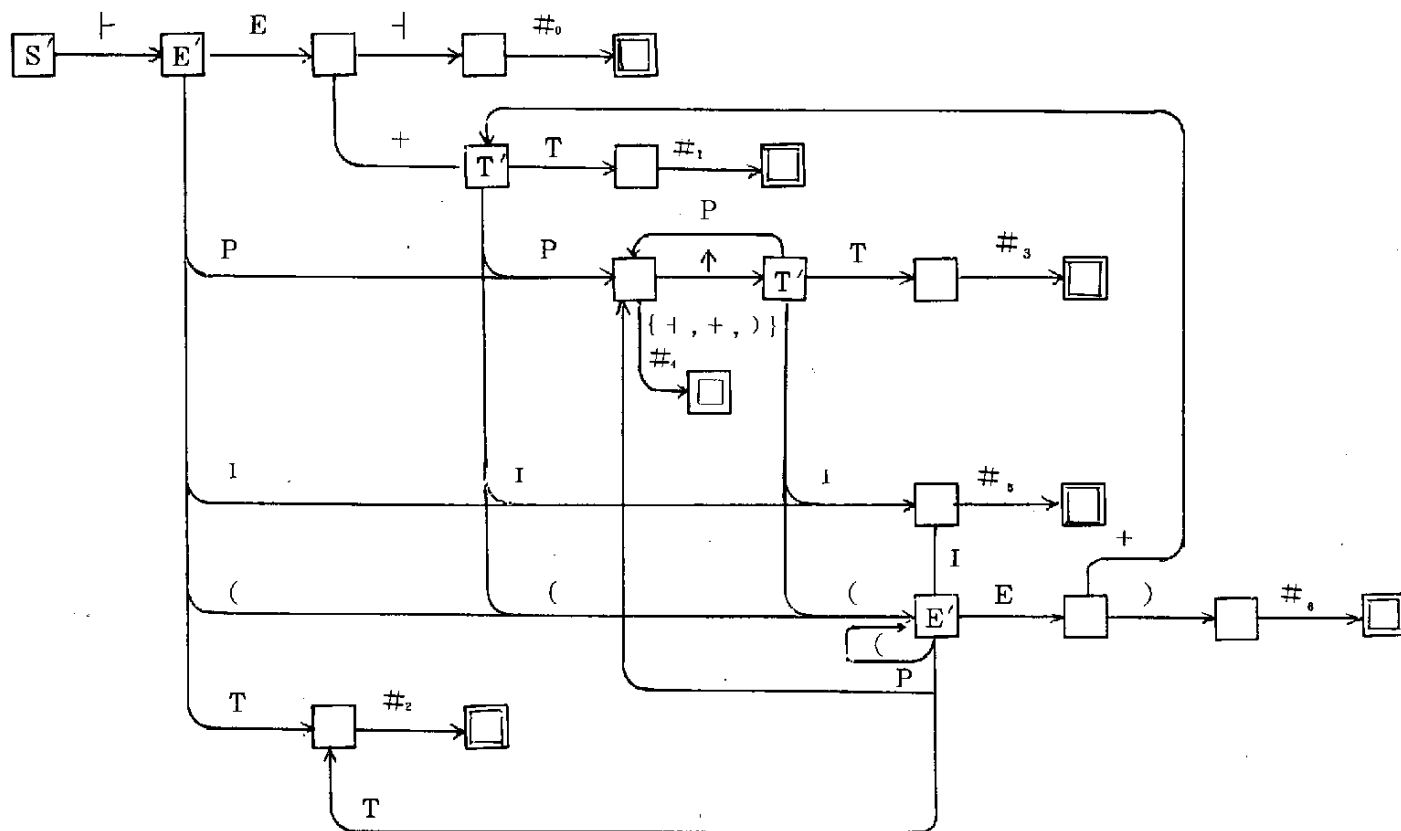


図4-13 LOOK AHEAD処理後のCFMSMグラフ

4.6 L(m)R(k) 処理

この処理は(図4-1)のstep5に対応するもので処理の概要は以下に述べるとおりである。

SLR(2)でinadequate 状態を分離できなかったとき $L(m)R(k)$ で分離できるか否かを試みる。その方法はCF SMグラフ上の該当inadequate 状態のnodeからの遷移が左のm個の記号と右のk個の記号によってユニークに決定できるか否かを調べることであり、その結果ユニークに決定できるならば該当inadequate 状態は $L(m)R(k)$ で分離できるということになる。

そして $L(m)R(k)$ で分離できる場合、該当inadequate状態のnodeをいくつかのnodeに分割することによりCF SMグラフを作りかえる。この分割操作の結果できるグラフはlook-aheadセットのみでarrow選択が可能になる。

a. 記号の定義

STAGEの $L(m)R(k)$ 処理の詳しいアルゴリズムを述べる。

前に(表4-3)で示すSYNTAX DEFINITION及び(図4-16)で示す。CF SMグラフを例として、アルゴリズムの説明で使用する記号を定義する。

(1) $\varphi : m$

φ がストリンのグの時 φ のうしろのm個のシンボル列を意味する。

(2) ${}^mL(N)$

NがCF SMグラフ上の node, φ をNの左側の文脈を構成するストリングとした時 ${}^mL(N)$ は $\varphi : m$ である。

例 (図4-16)の場合 ${}^2L(7)$ は $\{ae\}$ である。

(3) $FR^m(B)$

$C \rightarrow rB\alpha$ という生成規則があるとき、 $FR^m(B)$ は次のようにあらわされる。

$|r| \geq m$ ならば $r : m$

$|r| < m$ ならば $FR^{m-|r|}(C)r$

例 (表4-3)で示すSYNTAX DEFINITIONでは

$$FR^1(A) = \{a\}$$

$$FR^2(A) = FR^1(E) a \text{ となる。}$$

(4) $m_C^k(T)$

① T が $\#_N$ の場合

N 番目の生成規則が $A \rightarrow \omega$ であり, $B \rightarrow r A \rho$ なる生成規則があるとき $m_C^k(\#_N)$ は次のようになる。

$$|r\omega| \geq m \text{ ならば } \{(r\omega : m, F_T^K(A))\}$$

$$|r\omega| < m \text{ ならば } \{(FR^{m-1}r\omega^1(B)r\omega, F_T^K(A))\}$$

例 (表4-3)で示す構文の場合

の5番目の生成規則 $A \rightarrow e$ の場合

$${}^2C^1(\#_5) = \{(be, c), (ae, d)\} \text{ となる。}$$

② T が terminal 記号の場合

$B \rightarrow r t \rho$ なる生成規則があり, t が terminal 記号のとき $m_C^k(t)$ は次のようになる。

$$|r| \geq m \text{ ならば } \{(r : m, {}_1F_T^{K-1}(t))\}$$

$$|r| < m \text{ ならば } \{(FR^{m-1}r^1(B)r, {}_1F_T^{K-1}(T))\}$$

例 (表4-3)で示す構文の場合

$${}^2C^1(C) = \{(ae, c), (bA, c)\} \text{ となる。}$$

(5) $m_{BC}^k(T')$ (m, k) -bounded-context-pairs.

T' が node N からの遷移記号 T を持つ遷移のとき

$$m_{BC}^k(T') = \{(\sigma, \mu) \in m_C^k(T) \mid \sigma \in {}^mL(N)\}$$

となり σ, μ をそれぞれ Left, Right と呼ぶ。

例 (図4-16)のCFSMグラフの場合

① $N = \text{node } 7, T = \#_5, m = 2, k = 1$ とすると

$${}^2L(7) = \{ae\}$$

$${}^2C^1(\#_5) = \{(ae, d), (be, c)\}$$

よって

${}^2 BC'(\#_5') = \{(ae, d)\}$ となる。

② $N = \text{node } 7, T = C, m = 2, k = 1$ とすると

${}^2 L(7) = \{ae\}$

${}^2 C'(C) = \{(ae, c), (bA, c)\}$

よって

${}^2 BC'(C') = \{(ae, c)\}$ となる。

表 4-3

SYNTAX DEFINITION;

SYNTAX RULE;

$\langle S \rangle ::= \mid \langle E \rangle \mid : EXEC 0;$

$\langle E \rangle ::= a \langle A \rangle d : EXEC 1 \mid aec : EXEC 2 \mid$

$b \langle A \rangle c : EXEC 3 \mid bed : EXEC 4;$

$\langle A \rangle ::= e : EXEC 5;$

SYNTAX RULE END;

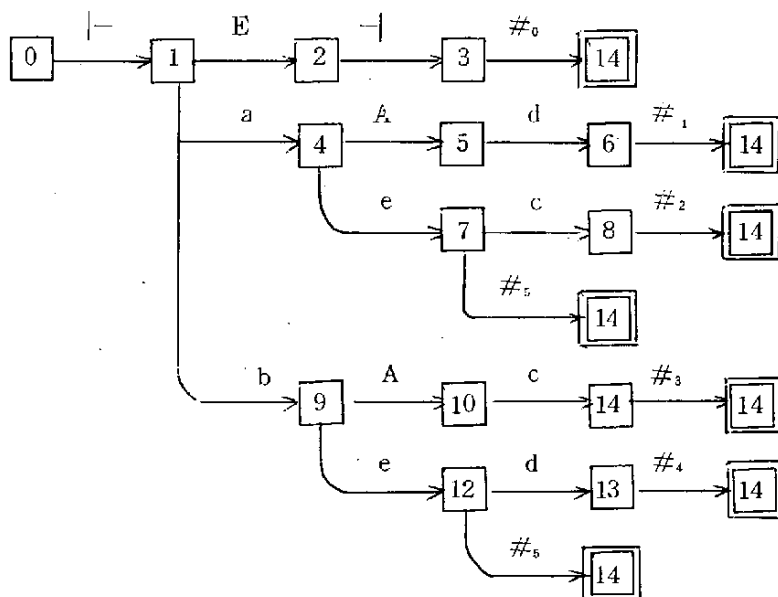


図 4-16 (表 4-3) に示す文法に対する CFSM グラフ

b. $L(m)R(k)$ 処理アルゴリズム

(1) $SLR(2)$ で分離できなかった *inadequate* 状態の *nodeN* に対し、 $L(m)R(k)$ で分離できるか否かを下記の手順により $m=2$, $k=1$ まで行なって調べる。

① *node N* からの遷移の集合を T' その要素を Ti' とするときすべての遷移 T' に対して *bounded-context-pairs*

$${}^mBC^k(T'i) = \{(\sigma_{i1}, \mu_{i1}), (\sigma_{i2}, \mu_{i2}), \dots\}$$

を求める。

② Ti' の *bounded-context-pairs* がユニークであるならば $L(m)R(k)$ で分離可能であるから(2)以降の処理を行なう。 $L(2)R(1)$ によっても分離できなかった場合は、どの生成規則が処理できなかったというメッセージを出力し処理を中止する。

(2) $L(m)R(k)$ は遷移を行う場合 *Left* の履歴を見る必要がある、これは効率が悪くなるため、*inadequate* 状態の *node* に関し、*CFSM* グラフを作りかえ *Right* すなわち *look-ahead* のみで遷移がユニークに決定できるようにする。

① (1)で求めた *bounded-context-pairs* が *Right* μ_{ij} だけでユニークになるとき、各遷移 Ti' に対する *look-ahead* セットはそれぞれ μ_i であり、この場合 *CFSM* グラフを作りかえる必要はない。^{*-1)} なぜならばこの場合は *Left* を見る必要がないからである。このあと③へ行く。

② (1)で求めた *bounded-context-pairs* が *Right* μ_{ij} のみではユニークに識別できないならば、*node* 分割を行なって *CFSM* グラフを作りかえねばならない、その手続きを以下に述べる。

Ti の *bounded-context-pairs* (σ_{ij}, μ_{ij}) に関して *CFSM* グラフ上を *nodeN* から σ_{ij} を逆にたどっていく、そして到着した *node* を M とすると M の子孫として $\{(\sigma_{ij}, \mu_{ij})\}$ の系列に対応する *node*, *arrow* を付け加える。^{*-2)}

次に、NからMへ逆にたどってゆき複数の子arrowを持つnodeにはじめて行き当たったとき、そのnodeの子孫から σ_{ij} 以下のarrow nodeを取り除く。

③ 他にSLR(k)が適用できないinadequate状態があれば、それを対象として(1)へ戻る。

(3) (2)において作りかえられた結果のグラフはnon-deterministicなグラフなので、これをdeterministicなグラフに変える。その方法は以下の通りである。

あるnode Nの子arrowに同一遷移記号を持つarrowが複数本あるとき、それらを1本のarrowにまとめる。そして、この操作をそのような状態のarrowが無くなるまで繰り返す。

*-1) (例4-5)がこの場合に該当する。

*-2) (例4-6)の(2)参照。

*-3) (図4-19)参照。

c. STAGEのL(m)R(k)処理例を次に示す。

(例4-5)

前記のように(表4-3)で示す構文からCF SMグラフを作ると(図4-16)になる。

(図4-16)においてはinadequate状態はnode 7及びnode 12である。このときnode 7の $\#_s$ arrow, node 12の $\#_s$ arrowに対して $F_T^2(A)$ を求めると、それぞれ $\{c \dashv, d \dashv\}$, $\{c \dashv, d \dashv\}$ となり、いずれもSLR(2)では処理できない。



- (1) まず node 7 からの各遷移 ($\text{arrow}\{c\}$, $\text{arrow}\{\#_s\}$) に対する $'BC'$ を求める。

$\text{arrow}\{c\}$ に対して $'BC'(c) = \{(e, d)\}$

$\text{arrow}\{\#_s\}$ に対して $'BC'(\#_s) = \{(e, c)\}$ となる。

よって node 7 からの各遷移に対する bounded-context-pair は各々ユニークになる。node 7 の inadequate 状態は $L(1)R(1)$ で分離できる。

- (2) (1)で求めた node 7 の bounded-context-pairs はそれぞれ Left は同じで Right がユニークである。それゆえ node 7 を分割し CF SM グラフを作りかえる必要はなく, node 7 からの遷移は Right を見るだけで確定する (node 7 から node 14 への遷移に対する look-ahead セットは $\{d\}$ となる)。同様に node 12 から node 14 への遷移に対する look ahead セットは $\{c\}$ となる。

上記の結果をまとめると (図 4-17) ができ上がる。

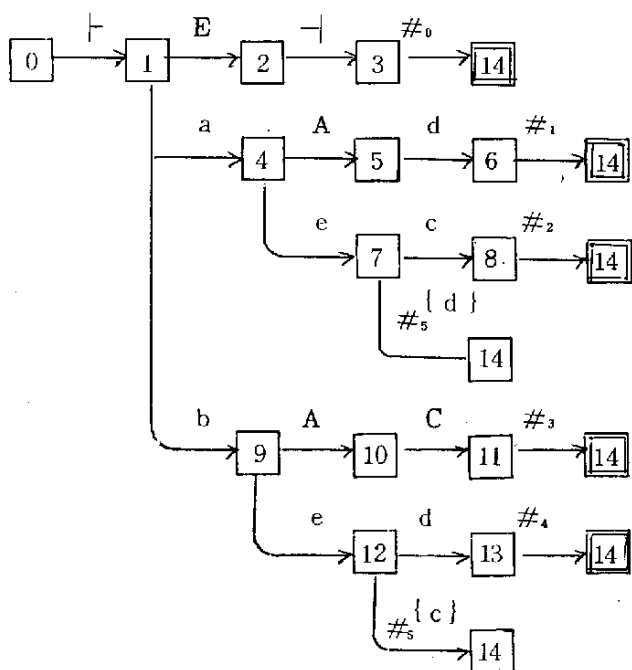


図 4-17 (表 4-3) に示す構文に対する L(1)R(1)処理後の CFMSM グラフである。

(例 4-6)

(表 4-4) に示す SYNTAX DEFINITION から CFMSM グラフを作ると (図 4-18) になる。

表 4-4

SYNTAX DEFINITION;

SYNTAX RULE;

<S> ::= |-<E>-| : EXEC 0;

<E> ::= a<A>d : EXEC 1 | ac : EXEC 2 |
b<A>c : EXEC 3 | bd : EXEC 4 |
a e f : EXEC 5;

<A> ::= e : EXEC 6;

 ::= e : EXEC 7;

SYNTAX RULE END;

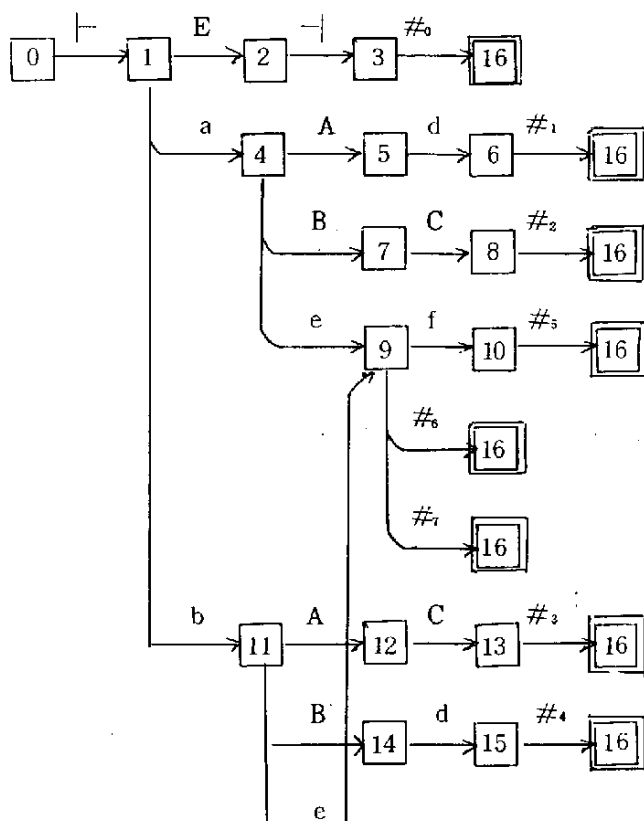


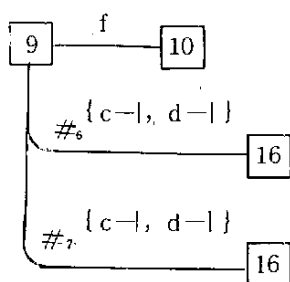
図4-18 (表4-4)に示す構文に対するCF SMグラフ

(図4-18)のCF SMグラフにおいてnode 9に inadequate 状態が生じている。

このとき node 9 の #6 arrow に対して $F_T^2(A)$ を求めると

$\{c \dashv, d \dashv\}$ であり, #7 arrow に対して $F_T^2(B)$ を求めると

$\{c \dashv, d \dashv\}$ となり, SLR(2) ではこの inadequate 状態は処理できない。



- (1) node 9 からの 3 つの遷移 ($\text{arrow}\{f\}$, $\text{arrow}\{\#_6\}$, $\text{arrow}\{\#_7\}$) を T' とし、それぞれに対して ${}^2BC^1(T')$ を求める。

$\text{arrow}\{f\}$ に対して ${}^2BC^1(f') = \{(ae, f)\}$

$\text{arrow}\{\#_6\}$ に対して ${}^2BC^1(\#_6) = \{(ae, d), (be, c)\}$

$\text{arrow}\{\#_7\}$ に対して ${}^2BC^1(\#_7) = \{(ae, c), (be, d)\}$

結果の bounded-context-pairs はユニークになる。それゆえ node 9 の inadequate 状態は $L(2)R(1)$ で分割できる。

- (2) (1) で求めた bounded-context-pairs は各々の Right を見るだけではユニークにならない。それゆえ node 9 を分割して CF SM グラフを作りかえる。

① (図 4-18) で node 9 から各 bounded-context-pair の Left すなわち $\{ae, be\}$ を逆にたどると node 1 に到達する。それゆえ, node 1 に ${}^2BC^1(f')$, ${}^2BC^1(\#_6)$, ${}^2BC^1(\#_7)$ に対応した arrow, node 群を子孫として追加する。^{*-4)}

node 4 の子孫, 及び node 10 の子孫から $\text{arrow}\{e\}$ をたどって到達できる node, arrow を取り除く, その結果 (図 4-19) となる。

- (3) (2) ででき上った non-deterministic な CF SM を deterministic な CF SM に変更する。

(図 4-19) のグラフは non-deterministic なものである。これを deterministic なグラフに変えると結果は (図 4-20) となる。

(図4-18)と(図4-20)を見比べると(図4-20)は(図4-19)のnode9をnode9'とnode9²の2つに分割した形となり、node9'及びnode9²からの遷移はlook-aheadのみに可能となる。

*-4) (図4-19)の点線で囲んだ部分。

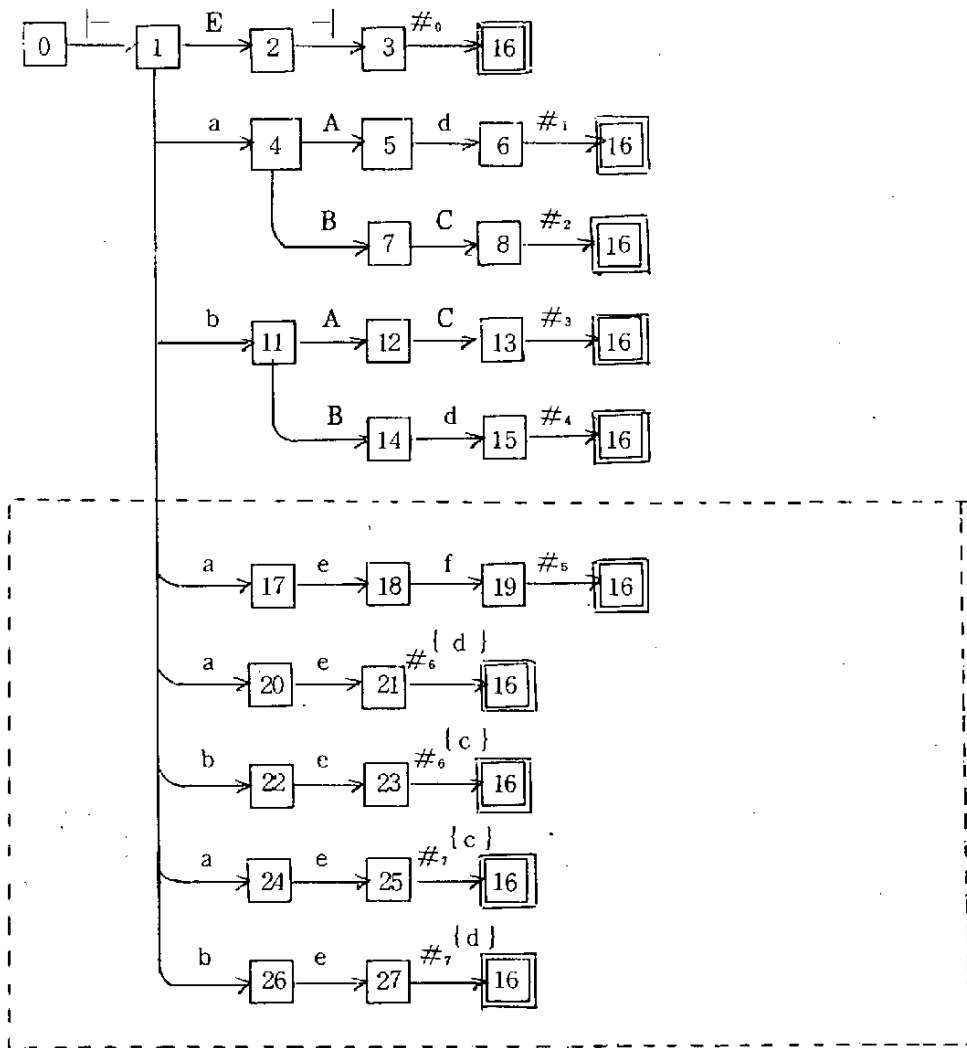


図4-19 (2)処理後の non-deterministicなCFSMグラフ

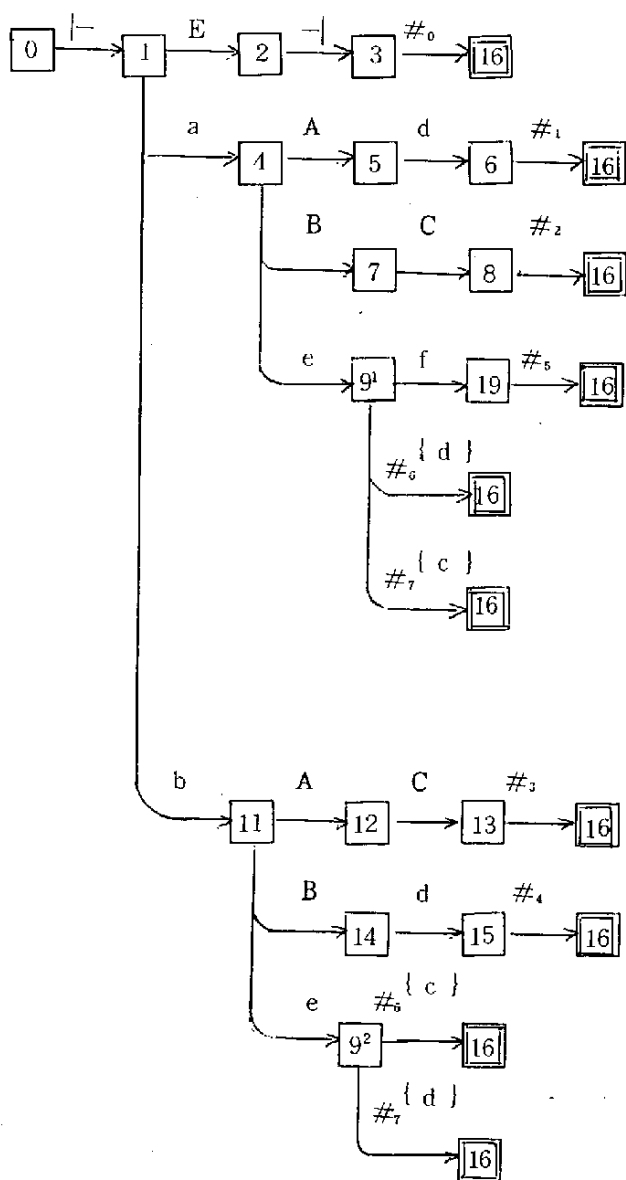


図 4-20 (3)処理後の CFMSM グラフ

4.7 遷移図作成処理

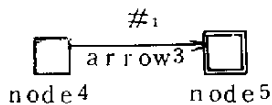
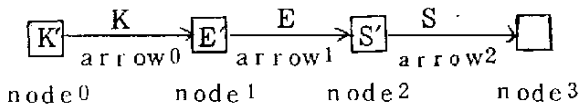
この処理は(図4-1)のstep 6に対応するものでCESMグラフより遷移図を作成する。

その場合必要があればlook-backセットを求める作業を行なう。遷移図作成の手続きを述べる前に記号の定義をしておく。

${}^k\text{LBN}(T)$: arrow{ T }のEND nodeに到達するまでにあらわれる長さ k の記号列の集合。

例 (図4-21)で ${}^1\text{LBN}(E) : \{K\}$

${}^1\text{LBN}(S) : \{KE\}$ である。



注) $\#_1$ は nonterminal 記号 E を生成する生成規則に対する $\#n$ 記号

図4-21

(1) CFSMグラフから遷移図を作成するアルゴリズム

T を nonterminal 記号とすると、記号 T を帰納記号とする $\#n$ arrow の ENDnode からの arrow で arrow{ T } を置きかえるという操作を全ての nonterminal 記号に対して行なう。

例 (図4-21) の場合 arrow 1 を node 5 から node 2 への arrow で置きかえる。

上記の操作を行なうにあたって考慮すべきことが2つある。

- ① 同一の nonterminal 記号を帰納記号とする $\#n$ arrow が複数本あるとき、ENDnode を1つにまとめる。
- ② 同一の nonterminal 記号をその遷移記号として持つ arrow が複数本あるとき、各セマンティックス処理 arrow の ENDnode から arrow

が複数本出るが、ドライブの際その経路選択を行なわねばならない。その為該当する nonterminal 記号に到着するまでの履歴が必要となり、これを見つけるのが Look-BACK 処理である。その場合必要なのが look-back セットである。

(2) look-back セットを求める方法

T を nonterminal 記号とすると、 $\text{arrow}\{T\}$ が複数本あるとき各 $\text{arrow}\{T\}$ に対して ${}^k\text{LBN}\{T\}$ を見つける、この操作は $k=1$ より始めて k を 1 ずつ増加しながら、各 ${}^k\text{LBN}(T)$ がユニークになるまで行なう。そして各 ${}^k\text{LBN}(T)$ がユニークになったとき、 $\#n \text{ arrow}$ の ENDnode からの各 arrow (各 $\text{arrow}\{T\}$ にかわるもの) の look-back セットは各 ${}^k\text{LBN}(T)$ となる。

例 4-7

(図 4-14) の CF SM グラフ上に上記の手続きを適用した結果できた遷移図は (図 4-22) で示される。

(図 4-22) の説明

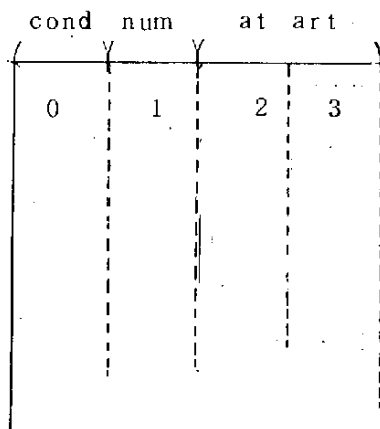
- ① node6, node7 における $\#n \text{ arrow}$ はともに nonterminal 記号 $\{E\}$ に対するものなので、その ENDnode は node17 の 1 つだけにしてある。
- ② (図 4-14) を見ると nonterminal 記号 E を、その遷移記号として持つ arrow は 2 本ある、そして各 arrow に対して ${}^1\text{LBN}(E)$ を求めると、各々 $\{ \mid \}$, $\{ (\}$ となる、よって (図 4-22) の node17 からの arrow は 2 本あって、各々 $\{ \mid \}$, $\{ (\}$ の look-back セットを持つ。

5. 遷移テーブル類

STAGEで作成される遷移テーブル類は(図4-22)のような遷移図を計算機で取り扱える形式(FASPソース)で表現したもので(図4-22)のnodeに対応するものとしてCG.NODE, arrowに対応するものとして, CG.ART, CG.SYM, CG.LABTとしてFASPソースの形で作成される。

(1) CG.NODEテーブル

CG.NODEテーブルは(図5-1)で示す様に3要素で1WORDを占め1つのnode情報を示す。



① condは第0バイト目を占め、このnodeの状態を示す。

0bit目ONの時: Start node

1 " :END node

2 " :テーブルドリブンの
 際ドライブの対象を
 このnodeまで引き
 戻す。

図5-1 CG.NODEテーブル

② at artは第2, 第3バイトを占め, CG.ARTへのポインタがBinayではいる。つまりこのnodeから出ているarrowの情報が入っているエリアの先頭番地を示す。

③ numは第1バイトを占め, CG.ARTのat artで示す所から何項目にわたって情報がはいっているかがBinayではいる。つまりこのnodeから出ているarrowの個数を示す。

(2) CG . ART テーブル

(3) CG . SYM テーブル

CG . ART及びCG . SYMは1対1の対応をしており，各1WORD
づつ合計2WORDで1つのarrow情報を示す。

CG . ARTテーブル

CG . SYMテーブル

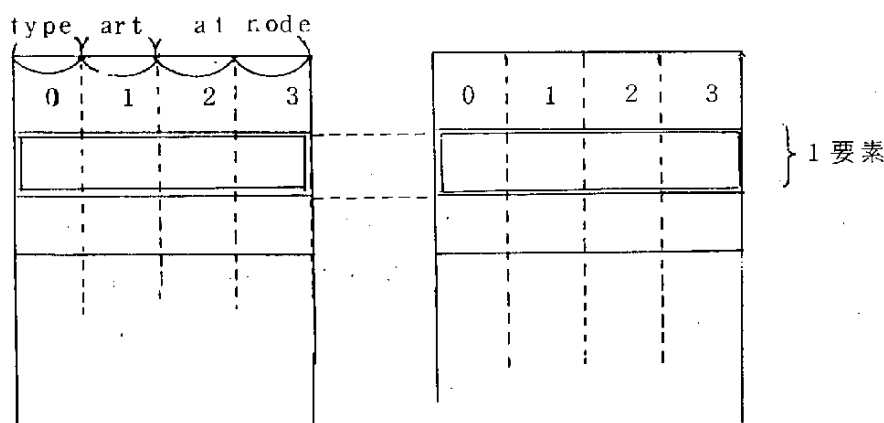


図 5-2 CG . ART及びCG . SYM
テーブル

表5-1 CG. ART, CG. SYMの内容説明表

0バイト	CG. ART		CG. SYM	
	1バイト	2, 3バイト	0, 1バイト	2, 3バイト
0bit目ON READ	0bit目ON	CG. NODEへの ポインタがBinary ではいる。	0バイト目に文字コードでterminal記号はいる。	
	1bit目ON		0, 1バイト目にBinaryでBASIC SYMBOL番号はいる	
	2bit目ON		0, 1バイト目にBinaryでnonterminal記号番号はいる	
1bit目ON LOOK-BACK	LOOK-BACKの 個数がBinary ではいる		CG. LABTのLOOK-BACK 情報組数がBinaryではいる。	このarrowに該当する CG. LABTへのポインター がBinaryではいる。
2bit目ON LOOK-AHEAD	LOOK-AHEAD の個数がBinary ではいる		CG. LABTのLOOK-AHEAD 情報組数Binaryではいる。	このarrowに該当する CG. LABTへのポインター がBinaryではいる。
3bit目ON EXEC			0, 1バイトにプッシュダウン すべきnonterminal 記号番 号がBinaryではいる。	2, 3バイトに実行すべきセマン ティックスルーテンググループ通番 がBinaryではいる。
4bit目ON POP				2, 3バイトにpop up すべき 数がBinaryではいる

(4) CG . LAB T テーブル

CG . LAB T テーブルは、LOOK-AHEAD, LOOK-BACK の場合の情報をセットしておくもので(図5-3)に示すように1項目2要素で1項目で1WORDを構成する。

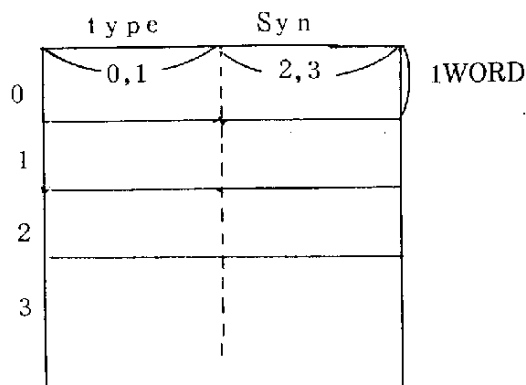


図5-3 CG . LAB T テーブル

- ① type はCG . LAB T を構成するWORDの0,1 バイトにあたり, Synにはいっている記号のタイプを示す。Synは2,3 バイト目に当り, typeで示された型に従って記号がはいる。

表5-2 CG . LAB T の内容に対する説明表

type	syn
0 bit目ON	terminal 記号が文字コード ではいる
1 bit目ON	BASIC SYMBOL番号が Binaryではいる
2 bit目ON	nonterminal 記号番号が Binaryではいる

例5-1

(表4-1)のSYNTAX DEFINITIONが与えられた場合, 遷

移テーブル類はどのようなになるかを示す。

(表4-1)から遷移図を作成すると(図4-22)になり、それを基としてCG・NODE, CG・ART, CG・SYM, CG・LATを作ると(図5-4)になる。以下に(図4-22)と(図5-4)の関係を少し述べる。(図4-22)のnode 2を見ると、子arrowはnode 3へのものと、node 4へのものの2本存在し、node 3へ行くarrowはREADしたものがterminal記号 $\nabla \mid \nabla$ であり、node 4へ行くarrowはREADしたものが $\nabla + \nabla$ のときである。(図5-4)のCG・NODEの第2項を見るとnumは2, at artは3である、つまりCG・ARTの第3項から2項目にわたって2本分のarrow情報がはいっている事を示す。

CG・ART第3項のtype(0バイト目)を見ると0bit目ONなのでこのarrowはREADであり、atr(1バイト目)は0bit目ONなのでCG・SYMの第3項にはterminal記号が入っていることを示す。実際CG・SYMの第3項0バイト目には $\nabla \mid \nabla$ が入っている。

CG・ARTの第3項のai node(2, 3バイト目)にはBinaryの3がはいっている、つまりこのarrowのEND nodeはnode番号3すなわちCG・NODEの第3項という事になる。CG・ARTの第4項はCG・NODEの第2項のもう1つの子arrow情報がはいっている。どうように(図4-22)のnode番号15の子arrowは全てLOOK-BACKである。(図5-4)のCG・NODEの第15項はnum:4, at art:21なのでこのnodeの子arrowは4本あってその情報はCG・ARTの第21項以降4項目に入っていることを示す。例としてCG・ARTの第21項を見るとtype:0bit ON, atr:1, at node:9であり、CG・SYMのupper:1, lower:3, つまりこのarrowはLOOK-BACKでその数及び組数はそれぞれ1で、その情報はCG・LABTの第3項に入っていることを示す。実際CG・LABTの第3項にはnon-terminal $\nabla \uparrow \nabla$ が入っている。

CG. NODE				CG. ART				CG. SYM			
cond num at art				type atr at node							
0	0	1	0	0	0	0	1	▽(			
1		2	1	1	0	1	10	▽(*2) ¹			
2		2	3	2	0	0	11	▽-			
3		1	5	3	0	0	3	▽T			
4		2	6	4	0	0	4	▽S			
5		1	8	5	3						1
6		1	9	6	0	1	10		1		
7		2	10	7	0	0	11	▽E			2
8		2	12	8	3		17	▽E			3
9		1	14	9	3		17	▽↑			
10		1	15	10	0	0	8		3		0
11		2	16	11	2	3	16	▽			1
12		2	18	12	0	1	10				4
13		1	20	13	0	0	11	▽T			
14	1			14	3		15	▽P			1
15		4	21	15	4	0	7		1		
16		1	25	16	0	1	10	▽			
17		2	26	17	0	0	11	▽)			
				18	0	0	13	▽+			
				19	0	0	4	▽P			6
				20	3		9				3
				21	1	1	9		1		4
				22	1	1	5		1		5
				23	1	1	6		1		6
				24	1	1	12		1		5
				25	3		15	▽T			7
				26	1	1	2		1		8
				27	1	1	12		1		

CG. LABT

0	0	▽-
1	0	▽+
2	0	▽)
3	0	▽T
4	0	▽P
5	0	▽↑
6	0	▽E
7	0	▽E
8	0	▽
9	0	▽

注)

- (* 1) CG. NODEの第0バイト, CG. ARTの第0バイト
CG. LABTの第0バイト及びCG. ARTの第0バイト
が0のときのCG. ARTの第1バイトは対応するbit
がONになっていることを示す。
- (* 2) ▼▼でかこまれているのは, 文字コード, 但し
nonterminal記号の場合, 実際には
nonterminal記号番号がBinaryではいる。
- (* 3) その他はすべてBinaryである。

図 5-4 (表 4-1)に示す文法に対する出力遷移テーブル類

6. 中間テーブル類

コンパイラ・ジェネレータは遷移テーブル類を作成するまでの中間テーブルとして次のような2グループのテーブルを持つ。

a. 第1グループ

(1) CG_NODEテーブル (TABLE)

これは1項目4要素で構成されnodeの情報をセットするのに用いられる。その内容は(表6-1.)に示される。

表6-1 CG-NODTの内容説明表

名 称	型	機 能	初 期 値
CG_NDY1	P(pointer)	node記号の入っているCG_NDNMへのポインター	null
CG_NDY2	B(Bit)	nodeのstate を示す (★-1)	B 000000000
CG_NDY3	P	このnodeをEND nodeとするarrow 情報が入っているCG_PARTへのポインター	null
CG_NDY4	P	子arrow情報の入っているCG_ARTへの ポインター	null

★-1) CG_NDY2

第1bit目ONの時 このnodeはStart node

2 " " End node

3 " " Basic Statement
対応node

(2) CG_ARTテーブル (TABLE)

これは1項目5要素からなりarrow情報セットの為に用いられる。

表6-2 CG-ARTテーブル内容説明表

CG-ARTテーブル CG-ARY1 (R)	CG-ARTテーブル CG-ARY2 (P)	CG-ARTテーブル CG-ARY3 (P)	CG-ARTテーブル CG-ARY4 (P)	CG-ARTテーブル CG-ARY5 (P)
1 bit目ON (READ)	シンボルテーブル CG-SYMへの ポインター	このarrowと同位の arrowが他にある 時それへのポインター	このarrowのEND -nodeへのポインタ -(CG-NODTの ポインター)	このarrowの start-nodeへ のポインター-(CG-NO DTへのポインター)
2 bit目ON (LOOK-BACK)	CG-LABT1 へのポインター			
3 bit目ON (LOOK-AHEAD)	CG-LABT1 へのポインター			
4 bit目ON (EXEC)	CG-SHAT へのポインター			
5 bit目ON (POP)				

(3) CG-SYMテーブル (TABLE)

1項目4要素からなりBASIC SYMBOL, terminal 記号,
nonterminal 記号の情報をセットする。

表6-3

CG-SYMY1 (B)		CG-SYMY2 (C)	CG-SYMY3 (P, N)	CG-SYMY4 (P)
1 bit目ON		terminal 記号 が入る	———	———
2 bit目ON		BASIC SYM BOL名が入る	Numeric型でBASIC SIMBOL番号が入る	———
3 bit 目ON	19 bit目ON (Start 記号)	nonterminal 記号名が入る	nonterminal 記号を生 成する生成規則グラフの SNODE へのポインター, CF SMグラフ作成後は nonterminal 記号番号	nonterminal 記号を 生成規則グラフが結合され た母型グラフ上の node へのポインター
	20 bit目ON (Basic Statement)			

(4) CG-SHATテーブル (TABLE)

1項目2要素からなり#n 記号に関する情報がセットされる。

表6-4

名 称	型	機 能
CG-SHAY1	P(pointer)	#nの帰納記号へのポインタ (CG-SYMへのポインタ)
CG-SHAY2	N(numeric)	CG-ARY1の4bit目がON(EXEC)の とき、セマンティックスルーチンググループ 通番 CG-ARY1の4bit目がON(POP) のときPOPすべき数

(5) CG-LABT1 テーブル (TABLE)

(6) CG-LABT2 テーブル (ARRAY)

CG-LABT1, CG-LABT2テーブルはLOOK-BACK, LOOK-AHEAD の時の情報をセットするもので, CG-LABT1 はインデックス, CG-LABT2はボディ部に当る。

CG-LABT1 (1項目3要素)

表6-5 CG-LABT1の内容説明表

CG-LABY11 (N)	CG-LABY12 (N)	CG-LABY13 (N)
LOOK-BACK, LOOK-AHEADの 個数がある	LOOK-BACK, LOOK-AHEADの 組数がある	CG-LABT2 へのポインタ がある

CG-LABT2 (1項目1要素)

表6-6 CG-LABT2の内容説明表

CG-LAB21(P)
シンボルテーブルCG-SYMへの ポインタがある

例 6 - 1

(図 6 - 1) で示すようなグラフがあったとき、CG-LABT1 及び CG-LABT2 はそれぞれ (図 6 - 2) で示すような形となる。

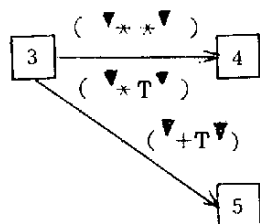


図 6 - 1

注) ① node 3 から node 4 への LOOK-BACK は (▽**▽) とする。

② node 3 から node 5 への LOOK-BACK は (▽*T▽), (▽+T▽) とする。

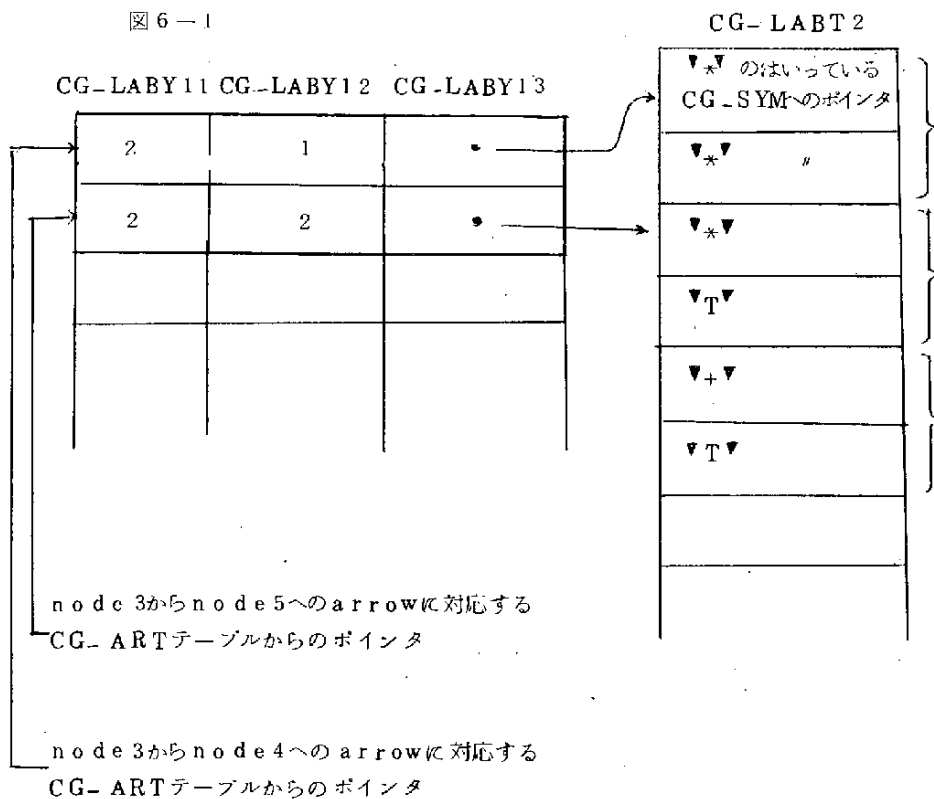


図 6 - 2

(7) CG-PARTテーブル (TABLE)

1項目2要素からなり nodeに入る arrow 情報がはいっている CG-ART テーブルへのポインタがはいっている。

表6-7 CG-PARTテーブルの内容説明表

名 称	型	機 能	初 期 値
CG-PART1	POINTER	流入 arrow 情報が入っている CG-ARTへのポインタ	null
CG-PART2	POINTER	流入 arrow が他にもある時, その Pointerが入っていると CG- PARTへのポインタ	null (他にない時も null)

(8) CG-NDNMテーブル (TABLE)

1項目2要素からなり node 記号に対する情報がセットされる。

表6-8

名 称	型	
CG-NDNY1	P	node記号に対応する nonterminal 記号の 情報が入っている CG-SYMへのポインタ
CG-NDNY2	P	nodeが以外に node記号を持つとき, その情 報がセットされている CG-NDNMへのポインタ

b. 第2グループ

第1グループの各種テーブル上の情報から第2グループのテーブルをセットする。これらは第1グループのテーブルをガーベージしたもので、遷移テーブル類の為のものである。

(i) OCG-NODTテーブル (TABLE)

OCG-NODTは node 情報セットの為に用い、1項目3要素からなる。

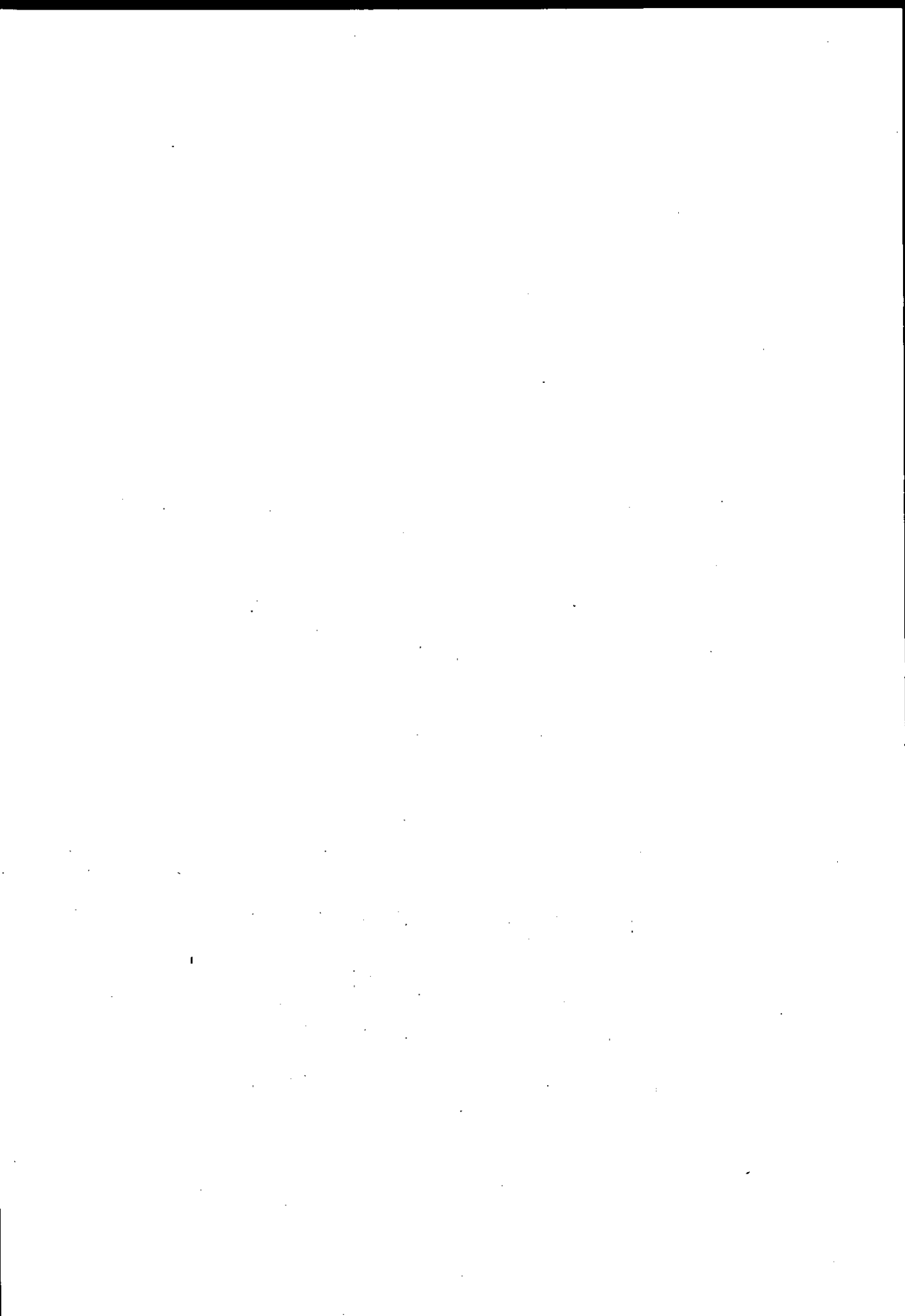
表 6-9 OCG_NODT テーブルの内容説明表

OCG_NDY1 (B)	OCG_NDY2 (P)	OCG_NDY3 (N)
Bit型でこのnodeの stateを示す。 1bit目ONStart node 2bit目ONEnd node 3bit目ONBasic Statement 対応node	Pointer型で OCG_ARTへのポインタ ーがはいる、つまりこの nodeから出ている arrow情報のはいついてい るエリアへのポインター。	Numeric 型でOCG_ NDY2 でしめされたエリ アから何項目にわたって arrow情報がはいついてい るかを示す、つまり arrowの数がはいる。

(2) OCG-ARTテーブル (TABLE)
 OCG-ARTテーブルは1項目5要素からなる。

表 6-10 OCG-ARTテーブル内容説明表

OCG_ARY1 (B)	OCG_ARY2 (B,N)	OCG_ARY3 (P)	OCG_ARY4 (P,C,N)	OCG_ARY4 (P,N)
1bit目ON (READ)	1bit目ON (B)	OCG-NODT へのポインター	terminal記号 (C)	
	2bit目ON (B)		BASIC SYMBOL 番号 (N)	
	3bit目ON (B)		nonterminal記号 番号 (N)	
2bit目ON (LOOK-BACK)	LOOK-BACKの 個数 (N)		LOOK-BACK情報 組数 (N)	CG-LABT2への ポインター (N)
3bit目ON (LOOK-AHEAD)	LOOK-AHEAD の個数 (N)		LOOK-AHEAD情報 組数 (N)	CG-LABT2への ポインター (N)
4bit目ON (EXEC)			プッシュダウンすべき nonterminal記号 番号が入っているエリアへ のポインター (P)	セマンティックスルーチン グループ通番 (N)
5bit目ON (POP)				POP UPすべき数 (N)



第 2 編

L S D 外 部 仕 様 書

目 次

第Ⅱ編 LSD 外部仕様書

1. プログラム要素	171
1.1 言語の基本的構造	171
1.2 文字	171
1.3 区切り記号	171
1.4 名前	172
1.5 コメント	173
1.6 プログラムの構造	173
2. データ要素	175
2.1 データの構成	175
2.2 データの型	177
2.3 変数の分類	180
2.4 変数の参照方法	182
3. 式	187
3.1 算術式	187
3.2 関係式	187
3.3 論理式	189
3.4 連結式	189
4. 宣言	191
4.1 DECLARE ステートメント	191
4.2 前後関係による宣言	192
4.3 省略時の解釈	193
4.4 属性	193
4.5 EQU ステートメント	198
5. 各ステートメントの説明	200

5. 1	PROCEDURE ステートメント	201
5. 2	END ステートメント	201
5. 3	ENTRY ステートメント	203
5. 4	代入ステートメント	203
5. 5	SETB ステートメント	208
5. 6	RESETB ステートメント	208
5. 7	SET ステートメント	209
5. 8	ERASE ステートメント	209
5. 9	SEARCH ステートメント	210
5.10	PUSH ステートメント	213
5.11	POP ステートメント	213
5.12	OPEN ステートメント	214
5.13	CLOSE ステートメント	214
5.14	WRITE ステートメント	214
5.15	READ ステートメント	215
5.16	SCAN ステートメント	216
5.17	DO ステートメント	216
5.18	GO TO ステートメント	217
5.19	IF ステートメント	219
5.20	CALL ステートメント	220
5.21	STOP ステートメント	220
5.22	RETURN ステートメント	221
5.23	ENTER ステートメント	221
5.24	EXIT ステートメント	221
5.25	NULL ステートメント	221
5.26	DEBUG ステートメント	221
5.27	INCREASE ステートメント	223

6. LSDとFASPのリンク	224
6.1 入口名のリンク	224
6.2 LSD プログラムにおける変数の 占有領域と方法	226

1. プログラム要素

1.1 言語の基本的構造

LSDのステートメントは、カードあるいはカード・イメージのファイルから入力されるものであるが、カードのカラムによる制限はなく、セミコロンで区切られるまでを1ステートメントとする。また、カード上の有効なカラムをユーザがコントロール・カードのパラメタで指定することができる。このパラメタを省略すると2～72カラムが有効、すなわち前のカードの72カラムが、次のカードの2カラムに続くものとする。

1.2 文 字

LSDでは、以下の55個の文字を用いることができる。

英 文 字 ... A|B|C|D|E|F|G|H|I|J|K|L|M|N|O|P|
Q|R|S|T|U|V|W|X|Y|Z (26文字)

数 字 ... 0|1|2|3|4|5|6|7|8|9 (10文字)

特殊文字 ... =|+|-|*|/|()|.|,|:|;|_|&|_|<|
>| (空白) (18文字)

下 線 ... _ (under - bar) (1文字)

これらの文字のほか、カタカナやこれ以外の特殊文字も文字ストリングやコメントの中に書くことができる。

1.3 区 切 り 記 号

特殊記号や特殊記号の組合せは区切り記号として用いられ、これらは演算子、カッコおよびその他の区切り記号に大別できる。

演算子

<算術演算子> ::= +|-|*|/

<比較演算子> ::= >|_|=|>|=|<|_|<|_|>|<=

<論理演算子> ::= $\neg$ | $\&$ | $\perp$

<連結演算子> ::= $\parallel$

カッコ

カッコはリストを囲むためや、いろいろなキーワードと結びついた情報を指定するのに用いる。

(左カッコ

) 右カッコ

/* コメントの始まりを示す

*/ コメントの終了を示す

▼ スtring定数を囲むために用いる

その他の区切り記号

, (コンマ) リストの区切りとして用いる

;(セミコロン) ステートメントの終わりを示す

:(コロン) ラベルのあとに続く

= 代入

空白 separator として用いる

. 名前の修飾(qualify)のために用いる

* (アスタリスク) テーブル etc. の要素の宣言のために用いる

1.4 名前

LSD において、値の参照、処理手続きの参照は名前によって行なう。

名前は英字で始まる 8 文字以内の英数字の列でなければならない。もし 8 文字を超える場合にはワーニング・エラーを表示し、はじめの 8 文字を有効とする(ただし 30 文字を超える名前があらわれた場合にはそのステートメントを無効とする)。但し、2～7 文字目までは下線(under-bar)があってもよい。

名前は変数名、ラベル名、手続き名、ファイル名などに用いる。LSD では

組み込み関数の名前は reserved word とするので名前として用いることができない。

組み込み関数としては ADDR, SUBSTR がある。

1.5 コメント

コメントは空白の出現し得るところならば、どこでも置くことができる。コメントは `/*`, `* /` で囲むことによって示し、コメントの中には使用可能なすべての文字を書くことができる。ただし、`*/` の組み合わせだけは書いてはならない。

1.6 プログラムの構造

LSD のプログラムは PROCEDURE ステートメントで始まり、END ステートメントで終わる 1 群のステートメントを procedure と呼び、ステートメントには 1 個のラベルをつけることができる。宣言文につけたラベルは無視される。

LSD では procedure は 1 レベルだけ入れ子にすることができる。すなわち 1 個の procedure の中に PROCEDURE ステートメントを書くことはできるが、その内側の procedure の中にさらに PROCEDURE ステートメントを書くことはできない。この外側の procedure のことを external procedure、内側の procedure のことを internal procedure と呼ぶ。

external procedure の中には、複数個の internal procedure を書くことができるが、その位置は external procedure の後の部分におかなければならない。すなわち、external procedure の中では internal procedure に含まれない部分が先に来て、そのあとに internal procedure が来るようにしなければならない。

名前の宣言で DECLARE ステートメントによるものは、その名前の参照以前に宣言されていなければならない。この DECLARE ステートメントは

internal procedureの中では書くことができないので注意すること。

さらに名前は、特にEXTERNAL と明示しないかぎり、その external procedure 内だけで有効である。

ただし external procedure の入口名、および external procedure 内で internal procedure に含まれない部分にある ENTRY ステートメントによる 2 次的な入口名はすべての external procedure 間で有効である。

2. データ要素

LSD の目的プログラムで実行時に操作できる情報をデータと呼ぶ。各データ項目は一定の型と表現を持っている。

2.1 データの構成

データはスカラ・データすなわち単一のデータとデータ項目の集合として構成される。

2.1.1 スカラ・データ

スカラ・データは定数あるいはスカラ変数のいずれかである。

2.1.2 集合データ

変数データ項目を配列、構造体、テーブル、スタックなどのようにまとめることができる。これらを総称して集合データと呼び、それを構成しているスカラ変数のことを集合データの要素と呼ぶことにする。

次に集合データの個々のものについて簡単に述べることにする。

(1) 配 列

同じ属性を持つスカラ変数の集合で1次元の順序づけられた要素の集合である。配列の大きさ(上限)は宣言文によって指定する。下限は1である。

配列を形成するスカラ変数のひとつひとつを配列要素と呼び、これを参照するときは配列名のあとに添字式を括弧()で囲んで書けばよい。配列要素の参照方法については後述する。

(2) 構 造 体

配列とは異なり、異なる属性のデータ項目をまとめるのには構造体を用いる。構造体は宣言文でその構造体の名前とそれに含まれる要素名とを指定することによって作られる。構造体の要素は単純変数あるいは配列である。

(3) テーブル・スタック

異なる属性のデータをひとまとめにし、さらにそのまとまりを複数個集めたものにテーブルとスタックがある。テーブルもスタックも同じ構成の構造体を複数個集めたものと考えてよいわけであるが、この構造体のことはテーブルあるいはスタックの項目と呼ぶことにする。

テーブルやスタックを宣言するときにはその項目の個数、すなわちテーブルまたはスタックのサイズとその1項目を構成する単純変数あるいは配列を構成順に列記する。テーブルの場合には、さらにこのテーブルの項目を参照するために用いるテーブル固有のポインタ変数を共に指定しなければならない。

テーブルとスタックの違いは、テーブル内の項目あるいはその項目の要素を参照したり、これに値を入れたりする時に、ポインタ変数によってその項目の位置を識別する（特にポインタ変数を指定しなければ固有のポインタ変数とみなされる）のに対してスタックの場合には、その項目あるいは項目の要素を参照したり、これに値を入れる時には常に現在使われているスタックの一番上の項目と考えられることにある。

例えば数値変数A 1と文字変数B 1から成るテーブルT（項目数10）と数値変数A 2と文字変数B 2から成るスタックS（項目数10）があるとする。

その構成は形式的には次の（図2-1）のようになる。

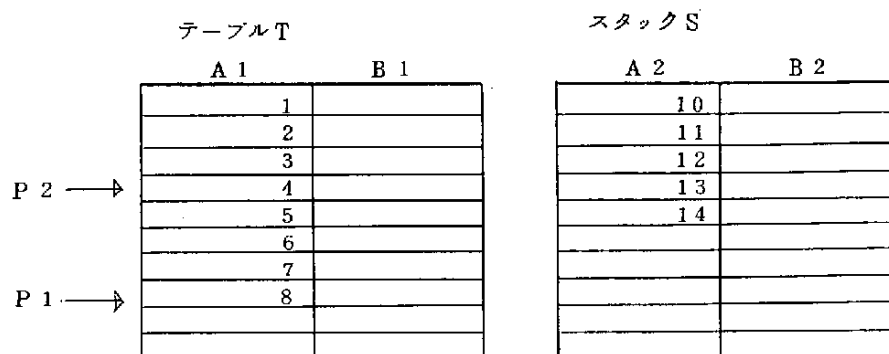


図2-1

Tを参照する時、もしTの固有のポインタ変数がP 1, P 2 が別のポインタ変数であれば

Tのみ書いたときには、P 1 で示されるTの項目すなわち、A 1 として8の入っている項目をとり出すことになる。

P 2 Tとポインタ変数を指定しているときには、このポインタ変数の指す項目、ここではA 1 として4の値をもつ項目をTとして参照することになる。

スタックSにおいては、もし第5番目の項目まで使われているとするとSあるいはA 2, B 2 と書けばこの第5番目の項目を指すことになる。

スタックの項目をスタック独自の方法でとり出したり、新しい項目を作ったりする時にはPUSH およびPOPステートメントを用いる。

さらにスタックの一番上から、いくつ目かの項目を参照する方法もあるが、詳しくは“変数の参照方法”として後述する。

テーブルの項目をとり扱かうのには、SET, ERASE, SEARCHの特別なステートメントがある。

2.2 データの型

LSDでとりあつかうデータの型は、ふつうのデータとプログラム制御用データの2つに大別される。

2.2.1 ふつうのデータ

ふつうのデータとは数値、ビットストリングおよび文字ストリングデータのことである。

(a) 数値データ

数値データは値として符号つき、あるいは符号なしの整数を持つものである。

〈整数〉 ::= 〈数字〉 | 〈整数〉〈数字〉

〈符号なしの整数〉 ::= 〈整数〉

〈符号付きの整数〉 ::= 〈整数〉 | +〈整数〉 | -〈整数〉

整数は $2^{35} - 1$ および -2^{35} の範囲内になければならない。

(1) 数値定数

数値定数は符号付き整数のことである。

(2) 数値変数

数値変数は値として符号付き整数を持つような変数で、宣言文では **NUMERIC** と宣言するか、あるいは型に関する属性を省略したものおよび **DECLARE** ステートメントにもよらず前後関係によっても解釈されない名前をもつものである。

(b) ビット・ストリング・データ

ビット・ストリング・データは 1 個以上のビットの列である。ビット・ストリング定数は、2 進表示のほか 8 進、16 進数として表わすこともできる。

(1) ビット・ストリング定数

① 2 進表現

2 進表現のビット・ストリング定数は 1 個以上の 2 進数 (0 or 1) の列を引用符で囲みその前に **B** をつけたもので、2 進数字は 36 個以内でなければならない。

例 2-1 **B**▼101▼ **B**▼10101▼

② 8 進表現

8 進表現のビット・ストリング定数は、3 ビットを 1 個の 8 進数で表わしたもので、1 個以上 12 個以下の 8 進数字を引用符で囲みその前に **O** をつけたものである。

例 2-2 **O**▼0473▼ ... **B**▼000100111011▼ の速記形
 O▼01101▼... **B**▼00000100100001▼ の "

③ 16 進表現

16 進表現のビット・ストリング定数は 4 ビットを 1 個の 16 進数で表わしたもので 1 個以上 9 個以下の 16 進数字を引用符で囲みその前に **X** をつけたものである。

なお、16進数では10進数10をあらわすのにA、11をB、12をC、13をD、14をE、15をFで示す。

例 2-3 X▼1234▼...B▼0001001000110100 ▼ の速記形
X▼049AC▼...B▼00000100100110101100▼の

(2) ビット変数

ビット変数は値としてビット・ストリング定数をとるような変数のことで宣言文でBITと宣言し、さらにビット数を指定する。ビット数は36以下の整数で、指定しなければ36とみなされる。

(c) 文字ストリング・データ

文字ストリング・データは値として1個以上の文字の列を持つものである。

(1) 文字定数

文字定数は1個以上256個以下の任意の文字の列を引用符で囲んだものである。文字定数の中で引用符を表現したい場合には▼▼と2つ重ねて書いて1個の▼を表わす。

例 2-4 ▼COMPILER-GENERATOR▼

▼CHARACTER▼▼X▼▼デス▼

(2) 文字変数

文字変数は値として文字ストリングを取るような変数であって宣言文でCHARACTERと宣言し、さらに文字数を指定する。文字数は1以上256以下の整数で、指定しなければ4とみなされる。

2.2.2 プログラム制御データ

プログラム制御データにはラベル・データ、ポインタ・データ、アドレスデータの3種がある。

(a) ラベル・データ

ラベル・データは値としてラベル定数すなわちステートメント・ラベルをとるものである。

(1) ラベル定数

ラベル定数とは、プログラム中にステートメント・ラベルとしてあらわれている名前のことで、この名前は出現によって定義される。

(2) ラベル変数

ラベル変数は値としてラベル定数（ステートメント・ラベル）をとるような変数のことで、宣言文で LABEL と宣言することによって定義できる。

(b) ポインタ・データ

ポインタ変数はテーブル内の項目を指定するのに用いる変数で宣言文で独立に POINTER と宣言するか、あるいはテーブルの宣言の中で CONTROLLED に続く句の中を書くことにより定義する。後者の場合には、この変数はそのテーブルの固有のポインタ変数となる。

(c) アドレス・データ

アドレス変数はデータのアドレスを値としてもつような変数で宣言文で ADDRESS と宣言することに定義する。

アドレス変数に何かのデータのアドレスをセットする時には組込関数 ADDR を用いる。

また、アドレス変数の内容である特定のアドレスに値を代入したり参照する場合には、仮想変数を用いてアドレス変数の修飾型を作る。

2.3 変数の分類

LSDの仕様書の中では変数に関していろいろな呼び方を用いているので、これについて説明を行なうことにする。

(1) 単純変数

ただひとつその変数名を書くことによって識別できる変数名でスカラ・データである。

単純変数は構造体、テーブルあるいはスタックの要素であってもよい。

(2) 添字付変数

配列要素のことである。

(3) 仮想変数

アドレス変数の値として持っているアドレスの中に値を入れる時に、その値の持つ属性を明示するために仮想変数を用いる。

この変数は型に関する属性だけを備えたもので、この変数に値を代入することはできないし、仮想変数のみを参照することもできない。

仮想変数は常にアドレス変数と共に出現してそのアドレス変数の内容の属性をあらわすのに使われる。

仮想変数は単純変数とはみなさない。

(4) ポインタ変数修飾型

テーブルの項目、およびその項目の要素を識別したり、参照したりする場合に、特定のポインタ変数によってその位置を指定することがある。

これをポインタ変数修飾型と呼ぶことにし、形式は

ポインタ変数 . テーブル あるいは

ポインタ変数 . テーブル項目の要素

とする。

注) この時のポインタ変数は単純変数、添字付変数である。

(5) アドレス変数修飾型

アドレス変数の内容を参照したり、これに値を代入する時には、前述の仮想変数を用いてその属性を明らかにする必要がある。

これをアドレス変数修飾型と呼び

アドレス変数 . 仮想変数

の形で指定する。

注) このときのアドレス変数は単純変数、添字付変数およびこれらをポインタ変数によって修飾した形であって、配列、ポインタ変数修飾型の配列であってはならない。

(6) ビット変数修飾型

ビット・ストリング・データの場合、そのスカラ変数の何番目のビットか

を指定して参照したり on/off にセットしたりすることができる。これをビット変数修飾型と呼ぶことにし

ビットストリング型 ビットを示す
スカラ変数 整数

の形とする。

この形は SETB, RESETB, IF ステートメントの中だけで用いられる。

(7) スカラ変数

単一のデータ項目のことで次のものを含む

- 単純変数
- 添字付変数 (配列要素)
- ポインタ変数修飾型でテーブル項目の要素が単純変数
- " " 添字付 "
- アドレス変数修飾型

(8) 擬変数

組み込み関数 SUBSTR は、代入ステートメントの左辺に出現し得るので擬変数 (pseudo variable) と呼ぶことにする。その一般型は次のとおりである。

SUBSTR (S , i [, k])

- S : ビットあるいは文字型のスカラ変数
- i : 抽出するビット・ストリングあるいは文字ストリングのはじめの位置を示す。符号のない整数あるいは数値型の単純変数
- k : 抽出するビット・ストリングあるいは文字ストリングの文字数を示す。符号のない整数あるいは数値型の単純変数

2.4 変数の参照方法

LSD で用いる変数の参照方法は、次のとおりである。

2.4.1 単純変数

① テーブルの要素でない場合

単にこの変数の名前を書くことにより参照する。

② テーブルの要素となっている場合

①と同じ形あるいはポインタ変数による修飾型を用いる。ただし、後者の場合はスカラ変数ではあるが、全体としては単純変数とはみなさない。

例 2-5 X
 P 1 . X

2.4.2 添字付変数

配列要素は、添字式をカッコでくくったものを配列名のあとに書くことにより参照する。

添字式の形は次のとおりである。

<添字式> = $i * j + k \mid i * j - k \mid$
 $i * j \mid j + k \mid j - k \mid + k \mid - k \mid k$

ここで i, j, k は数値型の単純変数あるいは整数であって単純変数は構造体、テーブル、スタックの一部であってもよいが、ポインタ変数によって修飾される形であってはならない。

また、アドレス変数修飾の形であってはいけない。

例 2-6 A を数値型の配列とする
 A (I 1 * I 2)
 A (1 - I 3)

添字付変数全体としては、ポインタ変数修飾の形をとることができる。

例 2-7 P 2 . A (I 1 + 3)

2.4.3 配 列

代入ステートメントや入出力ステートメントなどで、配列全体を1度にとり扱かう場合には配列名だけを書くことによって代入したり、参照したりすることができる。

さらに配列がテーブル要素である場合はポインタ変数による修飾型を作ることでもある。

A, B を配列とする時、次の代入ステートメントは有効である。

例 2-8 A = P . B ;

2.4.4 構 造 体

構造体を1度にとりあつかう時には、その構造体名だけで参照したり、値を入れたりすることができる。

例 2-9 PUSH S1 INTO S2 ;

ただし $\left\{ \begin{array}{l} S1 \cdots \text{構造体} \\ S2 \cdots \text{スタック} \end{array} \right.$

2.4.5 テーブル

テーブルを参照したりこれに値を入れたりする場合、2通りの意味に分かれる。

- ① テーブル内の1項目を指す場合で、通常はこの意味でテーブル名を用いる。

例えば、T1 というテーブルがあり、固有のポインタ変数がP1 であるとすれば、単に

T1

と書くと、これはP1 . T1 に等しく、テーブルT1 内のP1 で示される項目を指すことになる。

この形のものにはポインタ変数修飾型を作ることができる。

- ② テーブル全体を指す場合で、入出力ステートメントで用いられる場合である。

これらの意味の解釈は個々のステートメントにより異なるが、特に指定がなければ①の意味と解釈してよい。

2.4.6 スタック

スタックを参照したりこれに値を入れたりする場合にも、テーブルの場合と同じく2通りの場合が考えられる。

- ① スタック内の1項目を指す場合で、単にスタック名を書いてもよし、スタック名のあとに $(* [- \text{整数}])$ をつけてもよい。例えばスタック名をSとすると

S は スタックの一番上の項目

S (*) は 同じく 一番上の項目

S (* - 1) は スタックの上から2番目の項目

を指す。

- ② スタック名によって、スタック全体を意味する場合である。

これらの意味の解釈は個々のステートメントにより異なるが、特にことわらないかぎり①の意味とみなす。

注) スタックの項目の要素に対して $(* [- \text{整数}])$ を指定することはできない。

2.4.7 アドレス変数修飾型

アドレス変数の内容に値を代入したり参照したりする場合には、アドレス変数修飾型を用いる。

この形式は

アドレス変数・仮想変数

であって、ここでいうアドレス変数は

- ・ 単純変数
- ・ 添字付 "

- ポインタ変数（単）の修飾型
- ポインタ変数（添）の "

のいずれかである。

2.4.8 ポインタ変数修飾型

テーブルの項目およびその項目の要素を識別したり参照したりする場合には、ポインタ変数修飾型を用いる。

この形式は

ポインタ変数・テーブル あるいは

ポインタ変数・テーブル項目の要素

であって、ここでいうポインタ変数は

- 単純変数
- 添字付変数

のいずれかである。

3. 式

式はある値を計算するのに使われる算法である。

式は、演算の種類によって算術式、論理式、関係式および連結式となる。これらの中で用い得るオペランドの型は式の種類によって定まっている。

3.1 算 術 式

算術式は算術演算を行なう式であって、算術データから成るオペランドと算術演算子によって作られるものである。

LSDでは算術演算子は $+$, $-$, $*$, $/$ の4種、すなわち、四則演算を行なう演算子である。

次にLSDの算術式をBNFで記述してみる。

$$\begin{aligned} \langle \text{算術式} \rangle &::= \langle \text{因子} \rangle \mid \langle \text{算術式} \rangle + \langle \text{因子} \rangle \mid \langle \text{算術式} \rangle - \langle \text{因子} \rangle \mid \\ &\quad + \langle \text{算術式} \rangle \mid - \langle \text{算術式} \rangle \end{aligned}$$
$$\begin{aligned} \langle \text{因子} \rangle &::= \langle \text{1 次子} \rangle \mid \langle \text{因子} \rangle * \langle \text{1 次子} \rangle \mid \langle \text{因子} \rangle / \langle \text{1 次子} \rangle \\ \langle \text{1 次子} \rangle &::= \langle \text{符号のない整数} \rangle \mid \langle \text{算術型スカラー変数} \rangle \end{aligned}$$

演算の順位は $*$ 、 $/$ の方が、 $+$ 、 $-$ よりも高く、同順位の場合には左から右へと算定する。

ここで算術型スカラー変数とは、算術型の値をとるスカラー変数のことであ
って、すなわち単純変数あるいは添字付き変数のいずれでもかまわない。

3.2 關係式

関係式は2つのオペランドを比較演算子でつないだもので両オペランドの型は等しくなければならない。

いずれの場合にも関係式の結果は長さ1ビットのビット・ストリング・データとなり、true の場合 1 (on) false (off) となる。

$$\langle \text{比較演算子} \rangle ::= \langle \text{比較演算子 1} \rangle | \langle | \neg \langle | \leq | \rangle = | \rangle | \neg \rangle$$

〈比較演算子 1〉 ::= = | < | > | =

関係式をBNFで書くと次のようになる。

〈関係式〉 ::= 〈算術型関係式〉 | 〈ビット型関係式〉 | 〈文字型関係式〉
| 〈ポインタ型関係式〉

〈算術型関係式〉 ::= 〈算術型オペランド〉 〈比較演算子〉 〈算術型オペ
ランド〉

〈算術型オペランド〉 ::= 〈算術変数〉 | 〈符号付きの整数〉

〈ビット型関係式〉 ::= 〈ビット型オペランド〉 〈比較演算子 1〉 〈ビッ
ト型オペランド〉

〈ビット型オペランド〉 ::= 〈ビット変数〉 | 〈擬変数〉 | 〈ビット定数〉
| 〈8進定数〉 | 〈16進定数〉

〈文字型関係式〉 ::= 〈文字オペランド〉 〈比較演算子〉 〈文字オペラン
ド〉

〈文字オペランド〉 ::= 〈文字変数〉 | 〈擬変数〉 | 〈文字定数〉

〈ポインタ型関係式〉 ::= 〈ポインタ変数〉 〈比較演算子 1〉 〈ポインタ
変数〉 | 〈ポインタ変数〉 〈比較演算子 1〉

NULL

算術型の関係式は、算術変数あるいは符号付きの整数を比較演算子で結んだものである。

ビット型では、ビット変数やビット・ストリング型の定数、あるいは擬変数SUBSTRでビット型のものをオペランドとし演算子で結ぶ。ただしこの場合、大小関係の比較はゆるさず等しいか否かの関係だけを調べる。左右のオペランドのビット数が異なる場合には、最左端のビットをそろえ、不足分にゼロをつめて右に延長する。

文字型では、文字型の変数、文字オペランドおよび擬変数SUBSTRの文字型のものをオペランドとし演算子で結ぶ。照合順序は
9 > … > 0 > … > Z > … > A > 特殊文字とする。

左右オペランドの文字数が異なる場合は、最左端の文字をそろえ、不足分に空白をつめて右に延長する。

ポインタ型の関係式は、両オペランドがポインタ変数で等しいか否かを調べるものである。特別な形として

〈ポインタ変数〉〈比較演算子 1〉NULL

がある。NULL は何も入っていないことを示すポインタの値で、キーワードである。

3.3 論 理 式

論理式は論理和、論理積、論理否定を行なう演算で、オペランドとしてはビット・ストリング・データあるいは関係式を取るものである。その構文をBNFで示すと次のようになる。

〈論理式〉 ::= 〈論理因子〉 | ¬〈論理因子〉

〈論理因子〉 ::= 〈論理因子〉 | 〈論理 1 次子〉 | 〈論理因子〉 & 〈論理 1 次子〉 | 〈論理 1 次子〉

〈論理 1 次子〉 ::= 〈ビット定数〉 | 〈8 進定数〉 | 〈16 進定数〉 |
〈ビット変数〉〈ビット型擬変数〉 | 〈関係式〉

オペランドのビット長が異なる場合には、最左端のビットをそろえて不足分にゼロをつめて右に延長する。

演算の順位は、関係式の中の比較演算が一番高く、次に¬(NOT)、その次に | (OR) あるいは & (AND) であるとし、同一順位のものでは左から右へと算定することにする。

3.4 連 結 式

連結式は文字型のオペランドに対して連結演算(concatenation)を行なうものである。

〈連結式〉 ::= 〈文字オペランド〉 | 〈連結式〉 || 〈文字オペランド〉

＜文字オペランド＞ ::= ＜文字変数＞ | ＜擬変数＞ | ＜文字定数＞

連結式においては，文字数の合計が256文字を越えてはならない。

4. 宣 言

LSDのプログラムで用い得る名前は、いろいろな対象を表わすことができる。例えばファイルをあらわしていることもあるし、文字データのことも算術データのこともある。名前によって表現される対象を規定する性質を属性(attribute)という。

属性には、データの型に関するもの、有効範囲に関するもの、初期値の設定に関するものなどがある。

この名前と属性を結びつけるのが宣言であるが、この宣言にはDECLAREステートメントによって陽に示すもの、暗黙のうちにきまるもの、前後関係によってきまるものがある。

4.1 DECLARE ステートメント

DECLARE ステートメントは名前の属性を指定するために用いられる非実行ステートメントである。

一般型は次のとおりである。

```
      { DECLARE }
      { DCL      } DCL要素, DCL要素, ...;
```

DCL要素には option 1 と option 2 の2種類がある。

option 1

名前 [属性] [INITIAL属性]

option 2

名前 [(次元属性)] 属性 [属性],

*名前 [属性] [INITIAL属性]]

[, *名前 [属性] [INITIAL属性]]

(1) option 1 は単純変数、配列、ファイルの定義のときに用いる。

(2) option 2 は構造体、テーブル、スタックの場合に用いる。

このときには、まず構造体、テーブルあるいはスタックの名前を書き、その後、全体としての属性（スタック属性、構造体属性、テーブル属性）を指定し、コンマで区切って構成要素を列記する。

構成要素はアスタリスクのあとにつづけて書き、さらに要素の属性を指定する。構造体、テーブルの場合には要素の属性の後に、INITIAL属性を指定することができる。

スタックには、INITIAL属性は指定できない。

個々の属性については後述する。

- ① 属性はそれが結びつけられる名前の直後にブランクで区切って任意個つづけることができる。
- ② DECLARE ステートメントによる宣言は、その名前に対する代入や参照が行なわれる以前になければならない。
- ③ 一つの名前に対する宣言は同一 external procedure では一度しか書けない。
- ④ 構造体、テーブル、スタックの要素に指定することのできる属性は

次元属性

NUMERIC	属 性
BIT	"
CHARACTER	"
POINTER	"
ADDRESS	"
LABEL	"

のみである。

4.2 前後関係による宣言

- (1) ステートメントの前にコロンの区切ってつけられている名前は、ステートメントラベル定数である。

- (2) テーブルに関する宣言の中で CONTROLLED に続く括弧の中の名前は、ポインタ変数とみなす。
- (3) external procedure の名前、あるいは internal procedure に含まれない ENTRY ステートメントによって規定される 2 次的な入口名は EXTERNAL 属性が与えられる。
- (4) SYSPRT, SYSIN は宣言されなければ、ファイル名であるとみなす。

4.3 省略時の解釈

LSD では、DECLARE ステートメントにもよらず、前後関係によっても解釈できない名前が出現すると算術型の単純変数とみなす。

4.4 属性

属性には、いろいろな種類のものがある。例えば配列の大きさを指定するのは次元属性であり、データの型が算術型であるかビット型であるかを示すのは型に関する属性である。これらの属性には組み合わせて宣言することの可能なものと排他的なものがある。以下個々の属性について述べることにする。

(a) 次元属性

次元属性は配列あるいはテーブル、スタックの上限を定義するものである。

一般型は (符号のない整数)

であり、この属性は名前の直後に他のすべての属性に先立って書かなければならない。

この宣言のあとに TABLE あるいは STACK 属性が指定されなければ配列とみなす。

(b) データの型に関する属性

(1) 算術データ

変数は NUMERIC と指定されたとき、あるいは他のデータの属性が全

く指定されない場合に算術データであると宣言される。

(2) ビット・ストリング・データ

ビット・ストリング・データに対しては、データの長さを付けて **BIT** と宣言することによってその属性を規定する。

BIT [(長さ)]

- ① 長さは符号のない整数 (3 6 以下) である。
- ② 長さを指定しなければ 3 6 とみなす。

(3) 文字ストリング・データ

文字ストリング・データは **CHARACTER** と宣言し、そのあとに文字数を指定することによってその属性を規定する。

$\left\{ \begin{array}{l} \text{CHARACTER} \\ \text{CHAR} \end{array} \right\} [(長さ)]$

- ① **CHARACTER** は **CHAR** と略して書くことができる。
- ② 長さは符号のない整数 (2 5 6 以下) である。
- ③ 長さの指定を省略すると 4 とみなす。

(4) ポインタ属性

ポインタ変数であることを宣言するのには **POINTER** と宣言すればよい。

POINTER

- ① テーブル宣言の **CONTROLLED** のあとに出現した名前もポインタ変数とみなされ、この場合には **POINTER** という宣言は行なってはならない。

(5) アドレス属性

変数名は **ADDRESS** と宣言されることによってアドレス変数という属性を与えられる。

ADDRESS

(6) ラベル属性

変数名は LABEL と宣言されることによってラベル属性を与えられる。

LABEL

(c) データの構成に関する属性

(1) 構造体属性

変数が構造体であることを示すには、STRUCTURE と宣言する。

この属性とともに書くことのできる属性は EXTERNAL 属性と INITIAL 属性であり、構造体は DCL 要素の option 2 に従う。

(2) テーブル属性

変数がテーブルであることを示すには、次元属性のあとに TABLE と宣言し、さらにこのテーブルに固有なポインタ変数を CONTROLLED のあとに指定しなければならない。

TABLE CONTROLLED (変数名)

この場合も DCL 要素の option 2 に従う。

テーブル属性とともに指定できるのは次元属性、INITIAL 属性ある。EXTERNAL 属性は暗黙のうちに持つ。

(3) スタック属性

変数がスタックであることを示すには、次元属性と共に STACK と宣言する。

STACK

スタックの場合にも DCL 要素の option 2 に従う。

EXTERNAL 属性は暗黙のうちに持つ。

(d) 範囲の属性

範囲の属性は、宣言された名前の有効範囲を指定する属性である。

EXTERNAL

EXTERNAL と宣言した変数は、どんな procedure からでも参照できる。EXTERNAL 属性を宣言した場合には、他のすべての属性が一致しなけ

ればならない。

(e) VIRTUAL 属性

VIRTUAL 属性は、この属性を宣言した変数を仮想変数とする。

仮想変数は、なんら実体のない変数でアドレス変数修飾型の場合に、アドレス変数にそのアドレスが入っているデータの属性を明らかにするために用いる。

VIRTUAL 属性を宣言した場合の他の属性との関係

- ① データの型に関する属性を必ず宣言しなければならない。
- ② EXTERNAL 属性は宣言できない。

(f) INITIAL 属性

INITIAL 属性はデータに初期値（定数）を与えるのに用いる。

$$\left. \begin{array}{l} \text{INITIAL} \\ \text{INIT} \end{array} \right\} (\text{item} [, \text{item}] \dots)$$

① ここで定数というのは

算術定数、文字定数、ビット・ストリング定数、ラベル定数であって、ここで宣言している名前の型（データ属性）と一致していなければならない。

② INITIAL 属性はポインタ・データ、アドレス・データ、スタックに指定することはできない。

③ item は次の形である。

{ 定数 | くりかえし定数 * 定数 | none | くりかえし定数 * none }

ここで none というのは、特に指定をしないことで空白でもよいし、全く何もなくてもよい。

④ 変数名が配列の場合は、はじめから順に指定したい個数の item を書けばよい。

例 4 - 1

```
DCL A(10) NUMERIC  
INIT(1, 3*0, 5, 2*, 0);
```

この場合 A には次の値がセットされる。

```
A(1) ... 1  
A(2) ... 0  
A(3) ... 0  
A(4) ... 0  
A(5) ... 5  
A(6) ... 不定  
A(7) ... 不定  
A(8) ... 0  
A(9) ... 不定  
A(10) ... 不定
```

- ⑤ ラベル変数に INITIAL 属性で書くことのできるラベル定数は、その procedure 内で定義されているステートメントラベルでなければならない。
- ⑥ EXTERNAL 属性を持つ変数にも INITIAL 属性は指定できるが、二つ以上のプロシージャで同じ変数に異なる値を INITIAL 属性で入れると結果は保証しない。
- ⑦ テーブル構造体を宣言している場合には、INITIAL 属性は各要素のあとに書く。

例 4 - 2

```
DCL A(10) TABLE CONTROLLED(P),  
* B CHAR(4) INIT ('ABS', .....),  
* C BIT INIT (0'1', .....),  
* D INIT(1, .....);
```

(g) ファイル属性

変数がファイルであることを示すには、ファイル属性を指定する。ファイルは常にEXTERNAL 属性を持っており、ファイル属性と共に他の属性を指定することはできない。LSDで用いるファイルはカード・イメージの入力およびライン・プリンタへの出力用のファイルSYSIN, SYSPRT以外は、すべてDECLARE ステートメントによって宣言しなければならない。一般型は次のとおりである。

ファイル名 FILE (RCDSIZE = { 符号のない整数₁ [WORD]
" " " " [CHAR]
BLKSIZE = 符号のない整数₂ [RECORD]
[, TMOD = { $\frac{6}{8}$ }])

- ① RCDSIZE はこのファイルのレコードサイズで、WORDあるいはCHARをつけることによって何語分あるいは何字分であるかを指定する。
- ② BLKSIZE はこのファイルが何レコードを1ブロックとしているかを指定するもので、RECORDは省略してもよい。
- ③ TMODは転送モードのことで、指定がなければ8とみなす。
- ④ ファイルは自動的にEXTERNAL 属性を持っているので1つのプログラムを構成する2つ以上の external procedureで同一スタイルに関して異なった宣言を行なうと結果は保証しない。

4.5 EQU ステートメント

同じエリアに異なる名前をつけたいとき(データの型は異なってもよい)には、EQUステートメントを用いて変数を結びつけることができる。一般型は次のとおりである。

EQU (変数名1, 変数名2 [, 変数名3] ...)
[, (変数名1, 変数名2 [, 変数名3] ...)] ... ;

ここで変数名 1 と変数名 2……は次の組み合わせである。

	変 数 名 1	変 数 名 2 …	
①	単 純 変 数 名 配 列 名 構 造 体 名	{ 単 純 変 数 名 配 列 名 構 造 体 名	但し 変数名 1 の TOTAL WORD 数が 1 項目の占有WORD数が 変数名 2 …のそれと等しいか 大きくなければならない。
②	テーブル, スタック の 要 素 名	{ 単 純 変 数 名 配 列 名	
③	テ ー ブ ル 名 ス タ ッ ク 名	① と 同 様	

5. 各ステートメントの説明

LSD で使用可能なステートメントは次のとおりである。

1. PROCEDURE	ステートメント
2. END	"
3. ENTRY	"
4. Assignment	"
5. SETB	"
6. RESETB	"
7. SET	"
8. ERASE	"
9. SEARCH	"
10. PUSH	"
11. POP	"
12. OPEN	"
13. CLOSE	"
14. WRITE	"
15. READ	"
16. SCAN	"
17. DO	"
18. GO TO	"
19. IF	"
20. CALL	"
21. STOP	"
22. RETURN	"
23. ENTER	"

24. EXIT	ステートメント
25. NULL	"
26. Debug	"
27. INCREASE	"

以下に個々のステートメントの説明を述べることにする。

5.1 PROCEDURE ステートメント

PROCEDURE ステートメントはそれ以下のステートメントがLSDの Procedureであることを示すと同時にその入口名を定義する。

一般型は次のとおりである。

```
label : { PROCEDURE } [OPTIONS(MAIN)]
       { PROC }
       [ ( 仮パラメタ・リスト ) ] ;
```

- ① OPTIONS(MAIN)のある場合は、これがメイン・プログラムの external procedureであることを示し、そうでなければサブ・プログラムとみなす。internal procedureにはOPTIONS(MAIN)と指定することはできない。
- ② OPTIONS(MAIN)を指定した場合には仮パラメタ・リストは指定できない。
- ③ 仮パラメタ・リストは仮パラメタを2個以上あればコンマで区切って並べたもので、仮パラメタとなり得るものはスカラ・データ、配列、構造体である。

5.2 END ステートメント

END ステートメントは2種類あり1つは procedureを終了させるものであり、もう1つはDOループを終了させるものである。

一般は次のとおりである。

- END procedure 名 ;

- ```
〔 label:〕 END 〔ラベル名〕:
```

- ### 例 5 — 1

```
L 1 : D O ;
 .
L 2 : D O ;

 END L1 ;
```

- 例 5-2

(例5-1)をEND:の形で書くと次のようになる。

```
L 1 : D O ;
 |
L 2 : D O ;
 |
 END ;

 END ;
```

- ③ procedureを終了させる場合のEND ステートメントには必ず procedure  
名を書かねばならない。

### 5.3 ENTRY ステートメント

ENTRY ステートメントは `procedure` の 2 次的な入口名を指定するもので一般型は次のとおりである。

入口名 : ENTRY [ ( 仮パラメタ・リスト ) ] ;

- ① 仮パラメタ・リストに関しては `PROCEDURE` ステートメントを参照のこと。
- ② `MAIN-PROCEDURE` 内では、指定してはいけない。

### 5.4 代入ステートメント

#### 5.4.1 一般形

代入ステートメントは式を計算してスカラ変数、配列などに値を代入するのに用いる。一般型は次のように分類できる。

Option 1 (スカラ変数への代入)

$$\left\{ \begin{array}{l} \text{スカラ変数} \\ \text{擬変数} \end{array} \right\} = \text{式};$$

Option 2 (配列への代入)

$$\left\{ \begin{array}{l} \text{配列} \\ \text{擬変数の配列} \end{array} \right\} = \left\{ \begin{array}{l} \text{配列} \\ \text{スカラ・データ} \end{array} \right\};$$

Option 3 (構造体への代入)

$$\text{構造体} = \left\{ \begin{array}{l} \text{構造体} \\ \text{テーブル} \\ \text{スタック} \end{array} \right\};$$

Option 4 (テーブル、スタックへの代入)

$$\left\{ \begin{array}{l} \text{テーブル} \\ \text{スタック} \end{array} \right\} = \left\{ \begin{array}{l} \text{構造体} \\ \text{テーブル} \\ \text{スタック} \end{array} \right\};$$

Option 5 (ラベル変数への代入)

$$\text{ラベル変数} = \left\{ \begin{array}{l} \text{ラベル定数} \\ \text{ラベル変数} \end{array} \right\};$$

Option 6 (ラベル変数の配列への代入)

$$\text{ラベル変数の配列} = \left\{ \begin{array}{l} \text{ラベル定数} \\ \text{ラベル変数} \\ \text{ラベル変数の配列} \end{array} \right\};$$

Option 7 (ポインタ変数への代入)

$$\text{ポインタ変数} = \text{ポインタ変数};$$

Option 8 (ポインタ変数の配列への代入)

$$\text{ポインタ変数の配列} = \left\{ \begin{array}{l} \text{ポインタ変数} \\ \text{ポインタ変数の配列} \end{array} \right\};$$

① Option 1 のスカラ変数は算術、ビットあるいは文字型のものである。

左辺と右辺の型が異なる場合には、右辺の式の算定結果を変換ののち左辺に代入する。

② Option 2 における配列は、算術、ビットあるいは文字型のものであって、右辺の配列についても同様である。両辺の型が異なる場合は上記①と同様の変換を行なう。

Option 2, Option 6, Option 8 で両辺が配列で両配列の大きさの異なる場合は次のようになる。

(左辺の配列の大きさ)  $\geq$  (右辺の配列の大きさ) のとき

右辺に対応する左辺の要素すべてに代入を行ない、

左辺の余ったものには代入を行なわない。

(左辺の配列の大きさ)  $<$  (右辺の配列の大きさ) のとき

右辺に対応する要素の存在するところまで左辺全体に代入を行なう。

さらに、Option 2, 6, 8 で右辺がスカラ・データ (Option 2 では、算術、ビットあるいは文字型でなければならない、Option 6 で

はラベル・データ、Option 8ではポインタ数である)の場合には、  
左辺の配列の要素すべてに右辺の同じ値を代入する。

③ Option 3では両辺の構造体の構成要素の順序、型、大きさはすべて等しくなければならない。

④ Option 4では、テーブルやスタックの1項目と右辺の構造体の構成は、全く等しいものでなければならない。

テーブルの場合には、このテーブルと共に宣言されているポインタ変数の示している項目に対して代入を行なう。代入ステートメントの実行後もこのポインタ変数は変らない、スタックの場合にも一番上の項目に代入が行なわれる。

⑤ Option 3, 4以外の場合に、両辺に出現する変数がスタックやテーブルの要素であることも可能であるがこの場合は、テーブルならばこのテーブルと共に宣言したポインタ変数の指している項目に属するもの、スタックならば一番上に入っている項目の中のものが用いられる。

#### 5.4.2 代入の際の変換規則

(1) 算術変数=算術式；

変換は行なわず、右辺の式を算定し左辺に代入する。

(2) 算術変数=論理式；

右辺の式の算定結果のビット・ストリングを右シフトして左辺に代入する。右辺のビット数が36ビットであれば変換は行なわない。

例5-1 DCL B1 BIT(10)；

B1 = B<sup>▼</sup>11100001100<sup>▼</sup>；

XX = B1；

この結果XXにはB<sup>▼</sup>000.....011100001100<sup>▼</sup>が入る。

26ビット

(3) 算術変数=連結式；

右辺の文字ストリングの値としては文字表現の〔+|-〕数値定数の形しかゆるさない。ただし、前後の空白はあってもかまわない。この場合には、右辺の値を内部表現に変換ののち左辺に代入する。

例 5-2 DCL C1 CHAR(5);

C1=▼+10▼;

XX=C1;

この結果 XX には内部表現の10すなわち

B ▼ 0...01010 ▼ が入る  
32ビット

(4) ビット変数=算術式;

右辺の算術式の内部表現を下から有効ビット数分だけ取り出す。内部的には左シフトを行なうことになる。

例 5-3 DCL B1 BIT(10);

XX=12;

B1=XX;

この結果XX には12、B1 にはB ▼0000001100▼ が入る

(5) ビット変数=論理式;

変換は行なわず、右辺の値を算定し、左辺に代入する。左辺のビット数の方が短い場合には右辺の頭からそのビット数分が取られ、右辺の方が足りない場合には、後に必要な個数だけ0をつけ加える。

(6) ビット変数=連結式;

右辺の文字ストリングが1, 0 から成る集合であればそのままこれを B ▼▼ の中の値と考えて変換し左辺に代入する。文字ストリングの文字数が左辺のビット数を越える場合には、頭からそのビット数分だけを代入し残りは捨てる。また、左辺のビット数が足りない場合は文字ストリングの後に必要な個数だけ0を補って考える。

空白は0とみなす。



例5-4 DCL B1 BIT(10), B2 BIT(3);

DCL C1 CHAR(8);

C1=▼11001010▼;

B1=C1;

B2=C1;

この結果B1にはB▼1100101000▼

B2にはB▼110▼

が代入される。

(7) 文字変数=算術式;

右辺の数値変数の値を文字コードに変換して必要ならば- (マイナス) を付け加えて左辺に代入する。変換の結果文字ストリング (左辺) の方が長ければ頭に空白をつめる。短かすぎる場合は結果は保証されない。

例5-5 DCL C1 CHAR(10), C2 CHAR(3);

XX=1521;

C1=XX;

C2=XX;

この結果C1には▼□□□□□1521▼, C2は不定となる。

(8) 文字変数=論理式;

右辺の算定結果のビット・ストリングの値を文字コードに変換して左辺に代入する。文字数がそろわない場合は(9)に準じる。

例5-6 DCL B1 BIT(5);

DCL C1 CHAR(8);

B1= B▼10111▼;

C1= B1;

この結果C1には文字コードの▼10111□□□▼が入る。

(9) 文字変数=連結式;

右辺のストリングの文字数が左辺と一致していればそのまま代入する。

左辺の方が短い場合には、右辺の頭から有効文字数を取って残りを捨て  
長い場合には後に空白をつめる。

例5-7 DCL C1 CHAR(10), C2 CHAR(3), C3 CHAR(5);

C1 = 'ABCDEFGH IJ';

C2 = 'XYZ';

C3 = C1;

C3 = C2;

はじめの代入でC3には 'ABCDE' が入り、あとの方では 'X  
YZ ' が入る。

## 5.5 SETB ステートメント

SETB ステートメントは、指定された変数の指定されたビットをONにする  
ステートメントである。

一般型は次のとおりである。

[label:] SETB ビット変数、符号のない整数

- ① ここでビット変数というのは、ビット属性を持ったスカラー変数のこと  
である。
- ② 符号のない整数とは、第何ビット目をONにするかを示すもので第1ビ  
ット目を1とし順に右の方へと番号をつけていく。
- ③ ビット変数のビット数はONにしようとするビットが含まれているよう  
に宣言されていなければならない。

## 5.6 RESETB ステートメント

RESETB ステートメントは、SETB ステートメントと対をなすもので指  
定された変数の指定されたビットをOFFにするステートメントである。

一般型

[label:] RESETB ビット変数、符号のない整数

## 5.7 SET ステートメント

SETステートメントはテーブルの一番はじめの空き領域に1項目をセットするためのステートメントで一般型は次のとおりである。

[ label : ] SET 構造体 INTO テーブル名

[ [ WHEN OVERFLOW [ THEN ] ] 単純型GOTO ステートメント ] ;

- ① SETの次に指定するものは構造体であるがこれはテーブルの1項目の構成と一致していなくてはならない。
- ② このステートメントの実行の結果、テーブルの中の空いている項目のものの中に構造体がセットされ、このテーブルと共に宣言したポインタ変数にはこの項目を示す値が入れられる。
- ③ WHEN OVERFLOW…以下の節がある場合には、テーブルがオーバーフローしているか否かをチェックしてオーバーしている時には、指定したラベルにJumpすることになる。
- ④ WHEN以下を指定しなければテーブル・オーバーフローは起らないという前提があるものとしてチェックは行なわない。

## 5.8 ERASE ステートメント

ERASEステートメントはテーブルの項目を削除するためのステートメントで一般型は次のとおりである。

[ label : ] ERASE テーブル [ FROM ポインタ変数 ] ;

- ① FROMポインタ変数を指定すると、そのテーブルの中でそのポインタ変数が示している位置からさいごまでの項目を削除する。
- ② ①の場合そのポインタ変数の値は変わらない。
- ③ FROM以下を省略するとテーブル内のすべての項目を削除する
- ④ いずれの場合にもこのステートメントを実行するとそのテーブルと共に宣言されたポインタの値は不定となる。

例5-8

いま次のようなDCLステートメントを持つテーブルがあるとする。

DCL T1(10) TABLE CONTROLLED(P1),

\* A1 CHAR(8),

\* A2 POINTER;

そしてさらにP2もポインタ変数で、このテーブルの中が次のようになっているとする。

|     | T1 | A1 | A2 |
|-----|----|----|----|
|     | A  |    | 5  |
|     | B  |    | 4  |
|     | C  |    | 6  |
|     | D  |    | 7  |
| P2→ | E  |    | 9  |
|     | F  |    | 4  |
|     | G  |    | 3  |
|     | H  |    | 2  |
|     |    |    |    |
|     |    |    |    |

(図5-1)

P2が第5番目の項目を指しているとするとき

ERASE T1 FROM P2;

を実行するとE~Hで始まる項目が削除されP1は不定となる。

## 5.9 SEARCH ステートメント

SEARCH ステートメントはテーブルの中から指定した変数あるいは定数と一致する項目をさがしてそのポインタをセットするステートメントで一般型は次のとおりである。

[ label: ] SEARCH テーブル要素 KEYED {  $\left. \begin{array}{l} \text{スカラ変数} \\ \text{定数} \end{array} \right\}$   
 [ FROM ポインタ変数<sub>1</sub> ] [ TO ポインタ変数<sub>2</sub> ]  
 {  $\left. \begin{array}{l} \text{BACKWARD} \\ \text{FORWARD} \end{array} \right\}$  [ INTO ポインタ変数<sub>3</sub> ]

{ [ IF EXIST [ THEN ] ] } 単純型 GOTO ステートメント ;  
[ IF NOT [ THEN ]

- ① テーブル要素は、この SEARCH ステートメントによって探し出そうとするデータの入っている可能性のあるテーブル要素のことでスカラ変数である。(配列名を指定することはできない)
- ② テーブル要素とスカラ変数あるいは定数 (KEYED の次に指定するもの) との型は一致していなくてはならない。
- ③ KEYED に導かれるスカラ変数あるいは定数の値と一致するものがテーブル要素の中にあれば、ポインタ変数<sub>3</sub> の指定されているときはこのポインタ変数に、もし指定がなければテーブルと共に宣言されたポインタ変数に、この項目を示すポインタの値をセットして IF 以下を実行する。
- ④ テーブル要素の中にさがしているものがない場合には、ポインタ変数の値を NULL にセットする。
- ⑤ FROM ポインタ変数<sub>1</sub> および TO ポインタ変数<sub>2</sub> は、この SEARCH ステートメントの対象がテーブル全体ではなく、これらのポインタ変数によって示される範囲に限定することを指定するものである。FROM を省略するとテーブル内の第 1 番目の項目、TO を省略するとテーブル内の一番さいごの項目が指定されているものとみなす。

FROM, TO に導かれているが、これらはテーブル内の位置を指しているのものであって、SEARCH を始める位置を示すのではない。このため、ポインタ変数<sub>1</sub> ≦ ポインタ変数<sub>2</sub> でなければならない。

- ⑥ BACKWARD あるいは FORWARD はテーブル内を探す search の方向を指定するもので BACKWARD の場合にはテーブルをさいごの方からはじめの方へと search を行ない、FORWARD は、はじめの方からさいごの方に向かって行なう、省略時には FORWARD とみなされる。
- ⑦ IF EXIST を指定すると、一致したものがあれば GOTO ステートメントを実行し、IF NOT であれば該当するものがないとき、GOTO ス

テートメントの実行を行なう。

- ⑧ IF EXISTは省略可能である。

例5-9

次のようなテーブルがあつてその内容として図のような値が入っているものとする。

DCL TAB1 (10) TABLE CONTROLLED (P),

\* A CHAR

\* N NUMERIC ;

|       |     |   |
|-------|-----|---|
|       | A 1 | 1 |
| P 1 → | A 2 | 2 |
|       | B 1 | 3 |
|       | B 2 | 4 |
| P 2 → | B 3 | 5 |
|       | A 1 | 6 |
| P 3 → | A 3 | 7 |
|       | B 4 | 8 |
|       |     |   |
|       |     |   |

( 図 5 - 2 )

今ポインタ変数 P 1 , P 2 , P 3 が図のようにテーブルの項目を指しているとする, この時NAMEという文字変数をキーとしてsearchを行なう。

SEARCH A KEYED NAME FROM P1

INTO P IF EXIST THEN GO TO L1;

NAMEに▼A 1▼が入っているとするとテーブルの第6番目に入っているA 1が探し出され, Pにこの項目を指す。ポインタの値がセットされ, L1へコントロールを移す。また,

SEARCH A KEYED▼B1▼ FROM P2 TO P3 INTO Q  
IF NOT GO TO L2;

ならばテーブル内の第5の項目から始めて第7まで探し, 存在しないので

L 2 に行く。このときポインタ変数 Q には NULL がセットされる。

```
SEARCH N KEYED 1 FROM P1 TO P2 BACKWARD
GO TO L3;
```

この場合にはテーブルの第 5 番目の項目から先頭に向かって P 1 の指している 2 番目のものまでを探して該当するものがないのでこのステートメントの次のステートメントにコントロールを移す、このテーブルと共に指定されたポインタ変数 P には NULL がセットされる。

## 5. 10 PUSH ステートメント

PUSH ステートメントはスタックを Push down して指定した 1 項目を入れるステートメントで一般型は次のとおりである。

```
[label:] PUSH 構造体 INTO スタック
```

```
[(WHEN OVERFLOW [THEN]) 単純型 GOTO ステートメント] ;
```

- ① ここで用いる構造体はスタックの項目と同じ構成でなければならない。
- ② WHEN OVERFLOW…以下の節があれば、スタックがオーバーフローしているか否かをチェックしてオーバーしている場合には、GOTO ステートメントを実行することになる。
- ③ WHEN OVERFLOW…以下を指定しなければオーバーフローのチェックは行なわない。

## 5. 11 POP ステートメント

POP ステートメントはスタックから 1 項目を pop up して指定したエリアに移すもので一般型は次のとおりである。

```
[label:] POP 構造体 FROM スタック名
```

```
[(WHEN NOTHING [THEN]) 単純型 GOTO ステートメント] ;
```

- ① 構造体はその構成においてスタックの 1 項目と等しくなければならない。
- ② WHEN 以下の節を指定した場合にはスタックが空になってとり出しても

ののない場合が起こると、GOTO ステートメントを実行する。この場合指定した構造体への値のセットは行なわない。

- ③ WHEN 以下を指定しない場合、ユーザーがスタックの空になる状態はあり得ないと考えているとみなしてチェックは行なわない。

## 5.12 OPEN ステートメント

OPENステートメントは、ファイル名とデータ・セットとを結合させ、ファイルに対する属性を指定するものである。

一般型は次のとおりである。

[ label: ] OPEN ファイル名 ( INPUT | OUTPUT );

- ① INPUT, OUTPUT の指定を省略するとINPUT とみなされる。  
② OPENステートメントが省略されると、最初のREAD/ WRITEステートメントでファイルの open を行なう。

## 5.13 CLOSE ステートメント

CLOSE ステートメントは open されているファイルを、結びつけられていたデータ・セットから切り離すもので一般型は次のとおりである。

[ label: ] CLOSE ファイル名;

- ① SYSIN, SYSPRT のファイルを用いるときはCLOSE ステートメントはなくともよい。  
② 上記以外のファイルを用いた場合必ずCLOSEしなければならない。

## 5.14 WRITE ステートメント

WRITEステートメントはレコードを内部ストレージ内の変数から出力用エリアへmove し、ブロッキングを行なって指定されたファイルへ書き出すものであって一般型は次のとおりである。

[ label: ] WRITE ファイル名 FROM 変数



[ IMMEDIATE ] ;

- ① WRITE ステートメントでは、データの転送の際に変換は行なわない。  
したがってユーザーは1レコード分のエリアを用意してそのエリアの変数を介してレコードの転送を行なう。
- ② WRITE ステートメントは SYPRT に対して行なうこともできる。  
この場合出力がライン・プリンタに出るわけであるがそのためにはユーザーが文字型の値にデータを変換して出力エリアの形式をととのえなくてはならない。
- ③ レコードサイズの小さいもの、レコード数の少ないものを write する場合には IMMEDIATE と指定すること。

## 5.15 READ ステートメント

READ ステートメントはレコードを入力用エリアから内部ストレージの指定エリアに転送する。(ファイルから入力用エリアへの転送を行ない、そのうちデブロッキングを行なう)。  
一般型は次のとおりである。

[ label : ] READ ファイル名 INTO 変数

[ AT END 単純型 GOTO ステートメント ] ;

- ① AT END 以下を指定すると End of File になったとき制御の行く先を示すことになる。
- ② SYSIN ファイルからの read の場合には、ブロックサイズ15レコード、レコードサイズ80バイトのファイルとみなす。
- ③ READ ステートメントでは、データの転送の際に変換は行なわない。  
従って変換はユーザーが自分で行なわなければならない。特に SYSIN からの read では、データは文字コードでエリアに入るものとする。

## 5.16 SCAN ステートメント

SCANステートメントは実行時にSYSIN ファイルからデータを読み込み、その中から指定された文字数のデータを指定エリアにセットするもので一般型は次のとおりである。

[ label : ] SCAN { { 符号のない整数 } } INTO 変数 ;  
                                  { 単純変数 }

- ① 符号のない整数は読み込む文字数を示すもので省略時には1と仮定する。
- ② 変数は文字変数でこのステートメントで指定している文字数を満たすものでなくてはならない。
- ③ 読み込まれるデータはカードの形式に基づいているものとし、80カラムのカードのどの部分を有効なものとするか、および1枚のカード毎にレコードがきれていると考えるか、あるいは 特定のカラムが次のカードの何カラムかに続いているものとするかは、ユーザーがラン・タイムのコントロール・カードのパラメタで指定する。

|                                             |                       |                       |
|---------------------------------------------|-----------------------|-----------------------|
| BOUND = ( n <sub>1</sub> , n <sub>2</sub> ) | n <sub>1</sub> …開始カラム | n <sub>2</sub> …終了カラム |
| RECORD/STREAM                               | RECORD …カード毎に区切る      | STREAM …続けて考える。       |

このパラメタを省略すると  $\text{BOUND} = (2, 72)$ ,  $\text{RECORD}$  とみなされる。

- ④ End Mark は1文字分全ビット ON の 0 7 7 7 である。
- ⑤ SCAN ステートメントをくりかえし実行している間、SYSIN ファイルに対して READ ステートメントを行なうことを禁止する。
- ⑥ SCAN の次に単純変数を指定してもかまわない。この場合も読み込みの文字数を示しているものであるから数値変数で1以上の値をとるものでなければならない。

## 5.17 DO ステートメント

DO ステートメントは DO ループの開始を規定するステートメントで一般型

は次のとおりである。

Option 1 DO 制御変数 = {  $\begin{smallmatrix} \text{スカラー変数} \\ \text{定数} \end{smallmatrix}$  } TO {  $\begin{smallmatrix} \text{スカラー変数} \\ \text{定数} \end{smallmatrix}$  }

{ BY {  $\begin{smallmatrix} \text{スカラー変数} \\ \text{定数} \end{smallmatrix}$  } } [ WHILE (論理式) ] ;

Option 2 DO WHILE (論理式) ;

- ① Option 1 の場合、制御変数とは算術型の単純変数であり、変数および定数も算術型のスカラーデータである。
- ② 制御変数 = の次を書いてある値を初期値、TO のあとのものを終値、BY のあとのものを増分と呼ぶことにすると、初期値  $\leq$  終値であつてかつ、WHILE 以下があれば、この論理式が真でないと DO ループを1度も実行しない。DO ループを1回実行するごとに増分を制御変数値に加へ変数の値が終値よりも大きくなるか、あるいは論理式の値が偽になるまでくりかえしこのループを実行する。
- ③ 増分は特に指定しなければ1である。
- ④ Option 2 では、論理式の値が真のとき DO ループを実行し、偽になると実行をやめて DO ループの次のステートメントに制御を移す。

## 5.18 GO TO ステートメント

GO TO ステートメントは制御を指定されたステートメントへ移させるものであり、一般型は次のように単純型のもの (Option 1) と計算型 (Option 2) の2種類がある。

Option 1 (単純型 GO TO ステートメント)

{ label : } {  $\begin{smallmatrix} \text{GO TO} \\ \text{GOTO} \end{smallmatrix}$  } {  $\begin{smallmatrix} \text{ラベル変数} \\ \text{ラベル定数} \end{smallmatrix}$  } ;

Option 2

{label:} { GO TO } 単純変数,  
                  { GOTO }

(ラベル定数[, ラベル定数]…);

- ① ここでラベル変数というのは、ラベル属性をもつスカラ変数のことである。
- ② Option 2の単純変数というのは、算術型の添字のない変数のことである。
- ③ Option 1においては、指定されたラベル変数の内容、あるいはラベル定数へ制御を移す。
- ④ Option 2においては、指定された算術型スカラ変数の値に従って、ラベル定数に制御を移す。

例 5-10

```

DCL LBL LABEL ;
 ⋮
 II = 2
LBLO: GO TO II, (LBL1, LBL2, LBL3) ;
 ⋮
LBL2: II = 3 ;
 ⋮
 GO TO LBLO ;
 ⋮
LBL3: LBL=L1 ;
 ⋮
 GO TO LBL ;
 ⋮
L1: ⋮

```

この例では最初にLBLOに入るとIIの値2に従ってLBL2に制御が移り、次にLBLOに来るとIIが3なのでLBL3に行く、ここでLBLというラベル変数にL1という値をセットしGO TO LBLの実行ののちL1にJumpする。

- ⑤ Option 2では、label定数の並びの個数は、算術型のスカラ変数のとり得る最大値と等しくしなければならない。もし、算術変数の値が大きいかあるいは0以下の場合にはこのGO TOステートメントの直後のステートメントに制御を移す。

## 5.19 IF ステートメント

IF ステートメントはプログラムの流れを式の値によって変えるもので一般型は次のとおりである。

### Option 1

IF 論理式 [ THEN ] IF, DO 以外のステートメント ;

### Option 2

IF 文字ストリング IS {  
NUMERIC  
ALPHA  
ALPHANUMERIC  
ENDMK

[ THEN ] IF, DO 以外のステートメント ;

### Option 3

IF スカラ・データ = スカラ・データ [ , スカラ・データ ] ...

[ THEN ] IF, DO 以外のステートメント ;

### Option 4

IF ビット変数・符号のない整数 IS { ON  
OFF }

[ THEN ] IF, DO 以外のステートメント ;

- ① Option 1 では論理式が true の値を取れば IF, DO 以外のステートメントを実行する。

- ② Option 2 では、文字ストリング変数は値として文字を持つような 1 文字の文字変数（スカラー変数）でなければならない。内容の文字コードが数値であるか、英数字か、英字かあるいはエンド・マーク O ▼ 777 ▼ であるかを調べ、条件を満足すると後半のステートメントを実行する。
- ③ Option 3 では、両辺のデータの型は関係式に準ずる。ただしスカラー変数でなければならない。
- 左辺が右辺のリストのどれかに該当していたら後半を実行する。
- ④ Option 4 ではビットの ON/OFF によって後半を実行するか否かをきめる。

## 5. 20 CALL ステートメント

CALL ステートメントは procedure を呼び出して procedure の指定された入口へ制御を移すのに用いる。

一般型は次のとおりである。

[ label : ] CALL 入口名 ( 実パラメタ・リスト ) ;

- ① 入口名として書くことのできるものは、同じ external procedure 内の procedure 名、あるいは ENTRY ステートメントによって定義される 2 次的な入口名、あるいは他の external procedure 名および入口名である。
- ② 実パラメタ・リストは実パラメタをカンマで区切って並べたもので、実パラメタとして与えられるものは、スカラデータ、配列、構造体、である
- ③ recursive call は与えないのでユーザーが注意すること。

## 5. 21 STOP ステートメント

STOP ステートメントはユーザープログラムの実行を終了させるステートメントで一般型は次のとおりである。

[ label : ] STOP ;

## 5.22 RETURN ステートメント

RETURN ステートメントはその RETURN ステートメントを含む procedure を終了させ、制御を呼び出し側の procedure に返すもので一般型は次のとおりである。

```
[label:] RETURN ;
```

## 5.23 ENTER ステートメント

ENTER ステートメントはこれ以降 EXIT ステートメントまでがユーザーのアセンブラによる own-coding であることを示し、一般型は次のとおりである。

```
[label:] ENTER ;
```

ユーザーの own-coding に関しては、LSD では check を行なわない。

## 5.24 EXIT ステートメント

ENTER ステートメントに対応し、ユーザーの own-coding を終了させこれ以後のステートメントが再び LSD 言語で書かれていることを示す。

一般型     EXIT ;

① EXIT ステートメントには label はつけられない。

## 5.25 NULL ステートメント

LSD では何も行なわない NULL ステートメントの使用をゆるしている。

一般型：

```
[label:] ;
```

## 5.26 Debug ステートメント

LSD ではプログラムのデバッグを容易にするためのステートメントを設けている。これには指定した範囲内のラベルの履歴をたどるものと、指定した変数

の内容を出力させるものとの2種類があり、次のステートメントに分類できる。

- (1) TRACE START ステートメント
- (2) TRACE END ステートメント
- (3) SNAP ステートメント

これらのステートメントに関しては、次に述べるが、TRACEの場合はSTARTからENDまでの間でコントロールの移ったラベル名を出力するステートメントでこの有効範囲は、スタティックに決まる。すなわちプログラムの物理的な順序によって決まり、論理的な流れにはよらない。またSNAPステートメントは、変数の内容を出力するものである。

いずれの場合にもこれらのステートメントに関連する出力はSYSPRTファイルに対して行なわれる。

- (1) TRACE START ステートメント

TRACE START ステートメントは、このステートメントからTRACE END ステートメントあるいはこのTRACE STARTステートメントを含むprocedureのEND ステートメントまでに出現する人口名、あるいはステートメント・ラベルを出力するものであって一般型は次のようになっている。

TRACE START ;

- ① TRACE START ステートメントにはラベルをつけることはできない。
- (2) TRACE ENDステートメント

TRACE ENDステートメントは上記のTRACE STARTステートメントの効果を終らせるもので次の形である。

TRACE END ;

- ① TRACE END ステートメントにはラベルをつけることができない。
- (3) SNAPステートメント

SNAPステートメントはこの位置での変数の内容を示すもので一般型は次のとおりである。

SNAP 変数のリスト ;



- ① 変数として指定できるものはスカラ変数, 配列, 構造体, テーブル, スタックであって出力形式は

|               |       |
|---------------|-------|
| ビット, ポインタ型ならば | 8 進定数 |
| 文字型           | 文字モード |
| 数値型           | 10 進数 |

の形である。変数のリストは変数をカンマで区切って並べたものである。

- ② テーブルを指定した場合にはそのポインタ変数の示す項目, スタックならば一番上のものを出力する。

- ③ SNAP ステートメントにはラベルをつけることはできない。

## 5.27 INCREASE ステートメント

INCREASE ステートメントは, 指定されたポインタ変数の値を BY 以後に示される項目数だけ increase する。

```
[label:] INCREASE ポインタ変数 [OF テーブル名]
 [BY { 定数 }
 { 変数 }] ;
```

- ① ポインタ変数が固有ポインタである場合は, OF テーブル名は省略可能である。
- ② 定数は数値定数, 変数は数値型のスカラ変数である。
- ③ BY 以後省略の場合は, 1 とみなす。

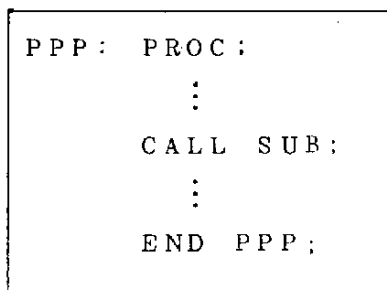
## 6. LSDとFASPのリンク

### 6.1 入口名のリンク

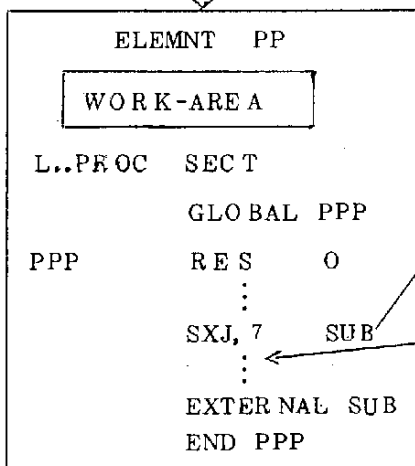
LSD プログラムにおいて external procedure の入口名, および external procedure 内で internal procedure に含まれない部分にある ENTRY ステートメントによる 2 次的な入口名は, GLOBAL 宣言がなされているので FASP プログラム内から LSD プログラムを参照する場合, その入口名を EXTERNAL として宣言しなければならない。又 FASP プログラム内の入口名を GLOBAL 宣言しておく, LSD プログラムから FASP プログラムを参照できる。

例 6-1

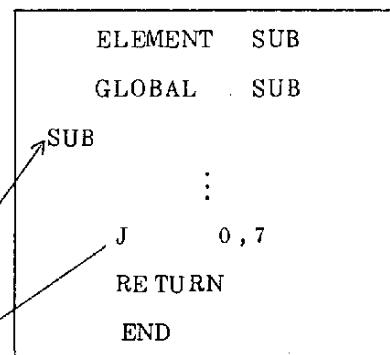
LSD プログラム



OBJECT



FASP プログラム



例 6 - 2

LSD プログラム

```

SUB: PROC;
 DCL...
 :
 RETURN;
END SUB;

```

OBJECT

```

ELEMENT SUB
 WORK-AREA
L..PROC SECT
GLOBAL SUB
L.IN0000 :
 J L.IN0001
 :
SUB インデックス
 退 避
 J L.IN0000
L.IN0001
 インデックス
 復 帰
 J 0,7
END

```

FASP プログラム

```

ELEMENT PPP
PPP
:
SXJ, 7 SUB
:
EXTERNAL SUB
END

```

## 6.2 LSDプログラムにおける変数の占有領域と参照の方法

### 6.2.1 単純変数

#### (1) 単純変数の占有領域

|      |                     |     |
|------|---------------------|-----|
| 数値型  | }                   | 1 語 |
| ビット型 |                     |     |
| ポインタ |                     |     |
| アドレス |                     |     |
| ラベル  |                     |     |
| 文字型  | (文字数 - 1) / 4 + 1 語 |     |

上のように領域がとられ、局所セクションの型でまとめられる。

#### 例 6-3

LSDプログラムのDCL ステートメント

```
DCL A, B BIT(3) INIT(0'7'), C POINTER,
 D ADDRESS, E LABEL, F CHAR(4) INIT ('AAAA');
```

OBJECT:

| L.. WORK | SECT |               |
|----------|------|---------------|
| A        | RES  | 1,            |
| B        | OCT  | 7000000000000 |
| C        | RES  | 1,            |
| D        | RES  | 1,            |
| E        | RES  | 1,            |
| F        | OCT  | 301301301301  |

#### (2) 単純変数を参照する方法

- ① LSD プログラム内でEXTERNAL 宣言を出したものを, FASP

プログラム内で参照するときには、そのプログラム内では、COMMON  
SECTIONを用いる。

例 6 - 4

LSDプログラムのDCLステートメント

```
DCL A EXTERNAL ;
```

OBJECT

```
A COMMON
A RES 1, Δ
```

FASPプログラム

```
A COMMON
A RES 1, Δ
```

(2) LSDプログラムのCALLステートメントの引数として引渡す。

例 6 - 5

LSDプログラム

```
DCL A ;
:
:
CALL SUB(A);
:
```

OBJECT

```
L..WORK SECT
A RES 1, Δ
:
L..PROC SECT
:
SXJ,7 SUB
NOP A,,1
:
EXTERNAL SUB
```

FASPプログラム

```
ELEMENT SUB
GLOBAL SUB
A OCT O
:
SUB

インデックス
退 避

LR =0/70777777
LA 0,7
KSTA *+1
LA *-*
STA A
```

### 6.2.2 配 列

#### (1) 配列の占有領域

1 要素の占有語数に配列の大きさをかけたものである。

例 6 - 6

LSDプログラム

```
DCL A(10), B(10);
```

OBJECT

```
L..WORK SECT
A RES 10, 2
B RES 10, 2
```

名前は先頭のエリアにつけられる。

#### (2) 参照方法

6.2.1 (2)と同じである。

### 6.2.3 構 造 体

#### (1) 構造体の占有語数

6.2.1 (1)と 6.2.2 (1)の各要素の語数を加えたものである。

例 6 - 7

LSDプログラム

```
DCL STR STRUCTURE,
 *A,
 *B CHAR(5) INIT(0 7 7 7 7),
 *C(2) BIT INIT(0 7 7 7 7 , 0 6 6 6 6);
```

OBJECT



|         |      |              |
|---------|------|--------------|
| L..WORK | SECT |              |
| STR     | RES  | 0            |
|         | EQU  | A = STR + 0  |
|         | EQU  | B = A + 1    |
|         | EQU  | C = B + 2    |
|         | RES  | 1, ▽, ▽      |
|         | OCT  | 301134134134 |
|         | OCT  | 301100100100 |
|         | OCT  | 700000000000 |
|         | OCT  | 600000000000 |

## (2) 参照方法

6.2.1 (2)と同じである。

## 6.2.4 テーブル

### (1) テーブルの構成

同じ構成の構造体を複数個集めたものである。

例 6-8

|     |    |                  |     |   |                          |
|-----|----|------------------|-----|---|--------------------------|
| TAB | 0  |                  | 11  | } | 内部ポインタ                   |
|     | 1  |                  | 14  |   | テーブルサイズ                  |
|     | 2  |                  | 1   | } | 1項目 { 数値型<br>文字型<br>ビット型 |
|     | 3  | 'B'              | 'B' |   |                          |
|     | 4  | 7000000000000000 |     |   |                          |
|     | 5  |                  | 2   |   |                          |
|     | 6  | 'B'              | 'B' | } | 1項目                      |
|     | 7  | 6000000000000000 |     |   |                          |
|     | 8  |                  | 3   |   |                          |
|     | 12 | '△'              | ~   |   |                          |
|     | 13 | '△'              | ~   |   |                          |

(例6-8)をLSDで記述すると

```
DCL TAB(4) TABLE CONTROLLED(P)
 *A INIT(1,2,3),
 *B CHAR(4) INIT(▼AAAA▼,▼BBBB▼,▼CCCC▼),
 *C BIT INIT(0▼7▼,0▼6▼,0▼5▼);
```

となる。

## (2) テーブルの占有領域

6.2.3(1)と同じく計算した1項目分に、SIZEをかけ2語を加えたものである。

### 例6-9

LSDプログラム

```
DCL TAB(2) TABLE CONTROLLED(P)
 *A INIT(1,2),
 *B CHAR(5) INIT(▼AAAAA▼,▼BBBBB▼);
```

#### OBJECT

```
TAB COMMON
P OCT 0
TAB OCT 2...(*1)
 OCT 10...(*2)
 EQU A=TAB+0
 EQU B=A+1
 ORIGIN TAB
 OCT 10
 ORIGIN *+1
 OCT 1
 OCT 301301301301
 OCT 301100100100
 OCT 2
 OCT 302302302302
 OCT 302100100100
```

#### COMMON領域

|       |             |
|-------|-------------|
| P     | 0           |
| TAB=A | 810         |
| B     | 810         |
|       | 1           |
|       | ▼A▼A▼A▼A▼A▼ |
|       | ▼A▼         |
|       | 2           |
|       | ▼B▼B▼B▼B▼B▼ |
|       | ▼B▼         |



( \* 1 ) 内部ポインタ

( \* 2 ) テーブルサイズ

(3) 参照方法

( 例 6 - 9 ) のように LSD プログラムにおいては、テーブルは  
COMMON SECTION としてエリアが確保されるため、参照する側で  
も COMMON SECTION としてエリアを確保しなければならない。

例 6 - 10

FASP プログラム

|       |        |                |
|-------|--------|----------------|
| TAB   | COMMON |                |
| P 1   | RES    | 1, ▼▼ (ポインタ領域) |
| TAB 1 | RES    | 8              |

(4) FACP プログラム内でテーブル・ポインタへ値をセットし、そのポイ  
ンタを LSD プログラムで用いる場合。

テーブル・サーチ等で一致する要素が見つかった時、その要素が示す TOP  
からの相対値をポインタへセットするのではなく、その要素が属する項目の  
TOP からの相対値をポインタへセットする。

例 6 - 11

|     |                |
|-----|----------------|
| TOP |                |
| 1   |                |
| 2   | 1              |
| 3   | ▼A▼A▼A▼A▼A▼A   |
| 4   | 70000000000000 |
|     |                |
|     |                |

文字列 ( ▼ A A A A ▼ ) が指定され  
た時相対値は 3 であるが、テーブル  
ポインタへは、文字列 ( ▼ A A A A ▼ )  
が属する項目の相対値 2 がセットさ  
れる。

## 6.2.5 スタック

### (1) スタックの構成

6.2.4 (1)と同じである。

### (2) スタックの占有領域

6.2.4 (2)と同じである。

例6-12

LSDプログラム

```
DCL STK(20) STACK,
 *A,
 *B CHAR(7),
 *C(2) BIT ;
```

OBJECT

```
STK COMMON
STK OCT 2...(*1)
OCT 146...(*2)
EQU A=STK+0
EQU B=A+1
EQU C=B+2
RES 100, ▼▼
```

COMMON 領域

|           |         |   |
|-----------|---------|---|
| STK<br>=A | 2       | 0 |
| B         | 102     | 1 |
|           | ▼▼ ~ ▼▼ | 2 |
| C         | ▼▼ ~ ▼▼ | 3 |
|           | ▼▼ ~ ▼▼ | 4 |
|           | ▼▼ ~ ▼▼ | 5 |

(\*1) 内部ポインタ

(\*2) テーブルサイズ

### (3) 参照方法

6.2.4 (3)と同じである。

スタックの項目あるいは項目の要素を参照したり、これに値を入れる

|         |     |
|---------|-----|
| ▼▼ ~ ▼▼ | 100 |
| ▼▼ ~ ▼▼ | 101 |

時には常に現在使われているスタックの一番上の項目を対象とするものである。



## 参 照 文 献



## 参 照 文 献 一 覽

1. J.W. Alsop  
A canonic translator  
MAC-TR-46 1967
2. Brooker and Morris  
A general translation program for phrase structure languages  
J.ACM 9 Jan. 1962 PP.1~10
3. ————  
Trees and Routines  
Comput .J. 5 1962 PP.33~47
4. ————  
The Compiler - compiler  
Annual Review in Automatic Programming Vol. 3 1963 PP.229~275
5. Brooker , Morris , Rohl  
Compiler - compiler facilities in Atlas Autocode  
Comput . J . 9 1967 PP.350~352
6. ————  
Experience with the Compiler - compiler  
Comput . J . 9 1967 PP.345~349
7. C.S.Carr , et.al.  
The TREE META Compiler - compiler system : A META  
Compiler system for the Univac 1108 and the Genral Electric 645  
TR 4-12 (RAD C - TR - 69-83)
8. A.Colmerauer  
Total Precedence Relation  
J.ACM Vol .17 No.1 1970 PP.14~30
9. F.L.Deremer  
Practical Translators for LR(k) Languages  
MAC - TR - 65

10. \_\_\_\_\_  
Simple LR (k) Grammars  
C.ACM Vol.14 No.7 1971 PP.453~458
11. P.J Dixon , J. Sable  
DM-1 a generalized data management system  
STCC 1967 PP.185~198
12. J.J. Donovan , H.F. Ledgard  
A formal system for the specification of the  
syntax and translation of computer languages  
FJCC 1967 PP.553~569
13. R.K.Dove  
Design highlights of ABAC - a compiler-compiler  
FJCC 1968 PP.1321~1328
14. J.A. Feldman  
A Formal Semantic for Computer Languages and its Application  
in a Compiler-compiler  
CACM Vol.19 No.1 1966.PP.3~9
15. J.A. Feldman , D. Gries  
Translator writing systems  
C.ACM Vol. 11 No.2 1968
16. E.N. Ferentzy , J. R. Gabura  
A syntax directed processor writing system  
FJCC 1968 PP.637~647
17. R.W. Floyd  
Syntactic analysis & operator precedence  
J.ACM Vol.10 Jan. 1963 PP.316~333
18. A descriptive language for symbol manipulation  
J.ACM Vol.8 Oct. 1961 PP.579~584
19. 二村良彦他  
PL/IW による TSS用コンパイラ・コンパイラの作成  
第12回プログラミング・シンポジウム報告集 1971 PP.43~53
20. 藤原宏, 渡辺勝正 -236-

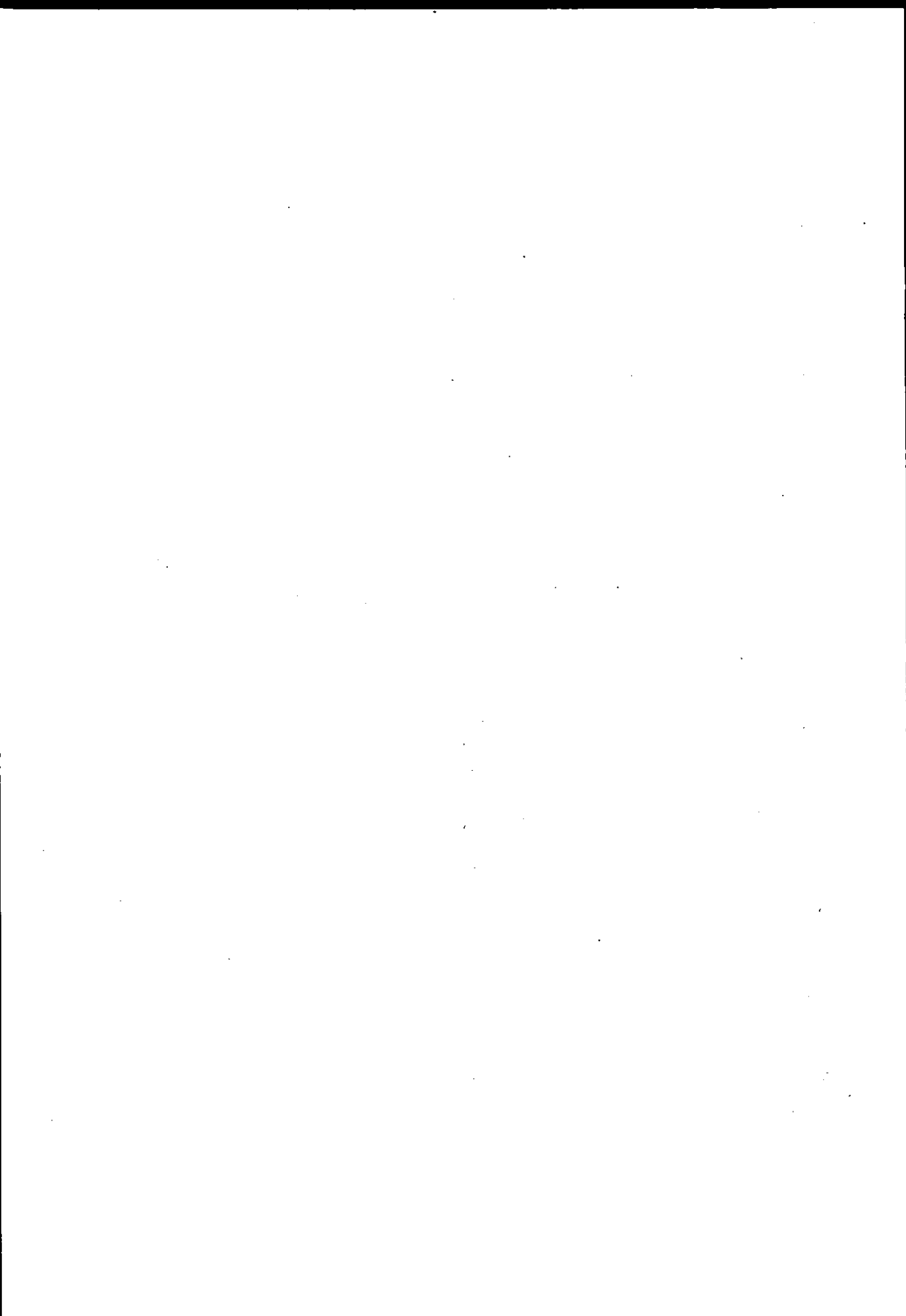


- Compiler 記述言語 : COL  
情報処理 Vol.9 No.4 1968 PP.187~196
21. \_\_\_\_\_  
COL で書かれた Compiler について  
情報処理 Vol.10 No.6 1969 PP.375~383
22. J. D. Ichbiah, S.P. Morse  
A Technique for Generating Almost Optimal Floyd - Evans Production  
for Precedence Grammars  
C.ACM Vol.13 No.8 1970 .PP.501~508
23. 井 上 謙 蔵  
順位言語の右順位解析  
情報処理 Vol.11 No.4 1970 PP.231~233
24. \_\_\_\_\_  
右順位文法  
情報処理 Vol.11 No.8 1970 PP.449~456
25. \_\_\_\_\_  
コンパイラ・コンパイラ  
産業図書
26. 井 上 他  
Compiler-compiler における意味記述言語について  
コンパイラ・コンパイラ シンポジウム報告集1971 PP.158~163
27. D. E. Knuth  
On the translations of languages from left to right  
Inf. Contr. Vol.8 Oct.1965 PP.607~639
28. 小 久 保 靖 世 , 佐 谷 鉄 夫  
コンパイラ記述言語 BPL  
情報処理 Vol.11 No.6 1970 PP.342~349
29. W.M. McKeeman, et.al.  
A Compiler Generator  
Prentice -Hall Inc.
30. \_\_\_\_\_  
The XPL Compiler Generator system  
FJCC 1968 PP.617~635

31. W.M. McKeeman  
An Approach to Computer language design  
Stanford Univ. Technical Report No. CS 48 1966
32. L. F. Mon dshein  
VITAL Compiler - compiler System reference manual  
Lincoln Laboratory Technical Note 1967 - 12
33. 日本電子工業振興協会  
超高性能高子計算機のソフトウェアに関する調査報告書 昭和42年3月
34. J. T. ONeil Jr.  
META-PI An On - line interactive compiler - compiler  
FJCC 1968 PP. 201~218
35. P. J. Pankhurst  
GULP - A compiler - compiler for verbal and graphic languages  
Proc. ACM 1968 PP. 405~421
36. M. Resnich , J. Sable  
INSCAN: a system directed language processor  
Proc. ACM 1968 PP. 423~432
37. D. T. Ross  
The AED approach to generalized computer - aided design  
Proc. ACM 1967 PP. 367~385

# 付 録

Translator Writing Systems



# Translator Writing Systems

Jerome Feldman, David Gries 共著

(カリフォルニア州, スタンフォード大学)

## —概念と原理の研究—

この報告書は、プログラミング言語のトランスレータの作成 (Translator Writing) を自動化することに関する最近の研究成果を批判的に検討しようとするものである。

第2章ではトランスレータ作成のためのシンタックスおよびその応用に関する形式についての研究を述べ、第3章ではTranslator Writingにおける post-syntax (すなわちセマンティックス) の自動化へのさまざまな試みについて述べることにする。さらにこれらに関連するトピックスについて第4章でとりあつかうことにする。



# 目 次

## I. INTRODUCTION

## II. SYNTAX

- A. Terminology
- B. Automatically Constructed Recognizers
  - 1. Operator precedence (Floyd)
  - 2. Precedence (Wirth and Weber)
  - 3. Extended precedence (McKeeman)
  - 4. Transition matrices (Samelson and Bauer, Gries)
  - 5. Production language (Floyd, Evans, Earley)
- C. Formal Studies of Syntax
  - 1. Bounded context grammars (Eickel, Floyd, Irons, Wirth and Weber)
  - 2. Deterministic pushdown automata (Ginsburg and Greibach)
  - 3. LR(k) grammars (Knuth)
  - 4. Recursive functions of regular expressions (Conway, Tixier)
  - 5. Summary

## III. SEMANTICS

- A. Syntax-Directed Symbol Processors
  - 1. TMG (McClure)
  - 2. The META Systems (Schorre, Schneider and Johnson)
  - 3. COGENT (Reynolds)
  - 4. ETC (Garwick, Gilbert, and Pratt)
- B. Compiler-Compilers
  - 1. FSL and its descendants (Feldman)
  - 2. TGS (Plaskow and Sherman, Cheatham)
  - 3. CC (Brooker, Morris, et al.)
- C. Meta-Assemblers and Extendible Compilers
  - 1. General discussion and METAPLAN (Ferguson)
  - 2. PLASMA (Graham and Ingerman)
  - 3. XPOP (Halpern)
  - 4. Extendible compilers - basic concepts
  - 5. Definitional extensions (Cheatham)
  - 6. ALGOL C (Galler and Perlis)

## IV. RELATED TOPICS AND CONCLUSIONS

- A. Other Uses of Syntax-Directed Techniques
- B. Formal Studies of Semantics
- C. Summary and Research Problems

## V. BIBLIOGRAPHY





## I Introduction

「…なぜならば、そのこと総ては、私や他の人々が今まで簡単には理解することのできなかつた長い詩の中に含まれているからだ」とジョーの「人間と超人」の中で悪魔が語っている。

コンパイラ・ライティングということは、プログラミング部門で長い間注目を集めてきた分野であり、また長く言い伝えられてきた夢物語であった〔Knu 62, Ros 64b〕。

最近では、コンパイラ作成者の仕事のいろいろな部分を自動化する方法へと研究者の注意が向きつつあるのである。そこで、この報告書ではこれらの研究の成果を批判的に概説してみようと思う。この同じ研究に関して前に発行され、広く配布されている報告書である Stanford Computer Science Report CS69 (1967年6月発行)、およびこれに対して我々のもとに寄せられた沢山の行き届いたコメントが、この報告書を正確かつ概念上明快たらしめるのに計り知れない程の助けとなったことをここで述べておきたい。

個々のシステムについて述べる前に translator writing に関する一般的な問題について、二、三述べておきたい。我々が特にコンパイラに着眼している理由はコンパイラがアセンブラやインタープリタに含まれる基本的な問題を含んでいるためである。compiler writing の研究には多大な努力がはらわれているのに、この問題に関して発表された書類は比較的少ないのである。そのため Translator Writing System (以下TWSと略す)の設計者たちは、それまでに注意を払って記述されたことのない方法をひとつひとつ形式化して行かざるを得ないのである。さらに困難なことは、トランスレータの性能に関しては効率というような合言葉があるだけで一般に認められた評価基準のない事である。コンパイラの効率はトランスレータやオブジェクト・プログラムの実行に要する時間とスペースによって決まってくる。使いやすさ、エラーの検出と回復の能力、編集機能、そしてリコンパイル (recompile) の速さなどが効率に重要な影響を与えるのである。これらは互いに矛盾する点も含んでいるので、コンパイラの効率についての絶対的な尺度を期待することはできない。ここで考えているTWSの設計者たちは多様性に富んだ独自の評価規準を採用することになっている。

compiler writing というものは、多くの視点から成る大規模なプログラミング作業なのでコンパイラ作成者を助けるために多種多様な方法が提唱されてきているのも不思議はない。実際大きなプログラムを作る時に役立つシステムの特徴 (たとえば、トレースやエディット) は、すべて compiler writing の道具となるものである。このことは、我々が compiler writing に関していろいろなシステムの特徴を調査する上で関係を持ってくるのである。術語として一般に認められているものがないので、ここで Translator Writing System (TWS) という語

は、ここで考えているところのプログラムおよび提唱されているプログラムを表わすものと定義する。TWSで用いられているトランスレータという言葉は、インタープリタ、コンパイラ、インクリメンタル・コンパイラ、アセンブラのいずれであってもよい。

この報告書は TWSの研究に関する歴史や紹介にはふれていない。この節の終りに記した参考文献の目録は 概説書としてどのような文献があるかを示すためのものである。また ここでは個々のシステムやいろいろな方法の比較を行なっているが、この報告書は Translator Writing System の使用者のための案内書を目的とするものではない。この報告書の目的は将来の調査に備えて統一された科学的基礎を作るべく、現在行なわれている研究に注意深い考察を加える事にある。このような事を目的とするため 我々はある特定のプログラムの実行に役立つシステムの特徴よりは、むしろ種々のTWSの設計の知的内容に重きを置くものである。

TWSの商業用コンパイラへの応用は、最近ようやく普及してきた。過去三年間の研究の進展は 決して遅々としたものであるとは言えないが はかばかしいものでもなかったのも、そのため 総てのTWSの開発を軽視する向きもあるようである。またある特定のTWSが 特定のコンパイラの特性に依存するということもありうるわけであるが、これはTWSを使用する上において大きな障害にはならないと思われる。ここで注目したいのは、TWSの成功例がすべて企業に追従しない人々によってなされた研究にあらわれているという点である。さらに 重要なことは、TWSの研究に固有な柔軟な言語 (flexible language) という概念が、企業が強調する規格化への動きとは全く反対の方向に進んでいくものであるという事である。従って TWSが商業用のコンパイラの作成者に採用されはじめたとはいえ、これがユーザーの手に委ねられるようになるまでには幾多の小さな改変を経なければならないであろう。

不幸にしてTWSに関する文献を読もうとするものは、かなりの注意を払って これに当たらねばならない。書物に形式的に表わされたシステムは、それを実際に使う段になるとかなり不備な点が目につくからである。また シンボリックな記法を用いることによって、TWS関係の書類の有用性を高めようとする数学主義 (mathematism) の流れもある。しかし 各研究者の間の連絡はほとんどなく、この分野に関する限り ひとつの発見が重複して行なわれることが多く、他の研究者の行なった研究を参照する事が少ないように思われる。そして このような状況と 評価するものの寛大さがあいまって、他の点に関しては申し分のないTWSの研究の歴史に汚点を残しているのである。

このTWSに関する調査は、シンタックスとセマンティックスのふたつの項目に大別することができる。自動構文法に関する研究は、TWSに関する研究の中で最も古く また最もよく理解されてきた部門である。この構文法は、さらに II章Bで述べるシステムに使われている限定

された一般論と、II章Cに述べるような理論的には有力であるが まだ実現の段階には到っていない方法とに分けられる。

一方 セマンティックスを三つの部分に分けているが、この分け方には まだ議論の余地があるように思われる。先ず、III章Aの *syntax-directed symbol processor* は一般に最も良く研究されてきた問題の特殊ケースとして、ここに取り上げたトランスレータの研究の一翼をになっていると言えよう。またIII章Bのコンパイラ・コンパイラはプログラムの *post-syntactic* な処理の機械化に多大な助力を与えんとするものである。最後のIII章Cでは、従来からある *macro-assembler* をさらに推し進めて、TWSに拡張しようとする動きに関連した二つの企画を扱っている。

IV章で扱っている関連するトピックスは、各章をふりかえって批評を加えようとするもので詳しい事は述べていない。すなわち、*syntax-directed technique* の他の利用法、および関係のある数値研究について書き添えておいたのは、TWSの研究とのつながりを明確にするためである。最後に 将来のTWSの開発に関係の深い、現在の実り多き研究活動にもふれておいた。なお 参考文献はアルファベット順に並べ、各章で参照した文献については その章の終りに記したことを述べておく。

## [ 参 考 文 献 ]

### REFERENCES FOR I

The *Communications of the ACM* and to a lesser extent the *Computer Journal* of the British Computer Society are the major journals for publications on translator writing. See especially *Comm. ACM* 4 (Jan. 61), 7 (Feb. 64), and 9 (Mar. 66).

Other general references: Che 64a, Flo 64b, Hals 62, Knu 62, Ran 64, Ros 64b, Weg 62, Wil 64b.

Formal descriptions of various programming languages: Bac 59, Ber 62, Brook 61, BroS 63, EvA 64, Gor 61, IBM 66, Naur 60, 63b, Rab 62, Samm 61, Shaw 63, Tay 61, Wir 66b, 66c.

## II SYNTAX

### A. Terminology

TWSに関する文献の内容を読む場合にわずらわしいことのひとつに 一定した言葉の定義 (notation) に欠けている点がある。この報告書を より読みやすくするために いろいろな著述者たちの使用している記号や、時にはシンタックスも自由に変更するなどのことを行なった。シンタックスについて論ずる場合には、我々はGinsburgによって用いられた記法を採ることにした。[Gin 66a, pp. 8, 9] しかし これとは矛盾しないようにしながら、構文的な超言語である Backus - Naur Form (BNF) の記法 特に記号  $::=$  の使用などの記法は、随時併用して読みやすくなるように努めることにした。

この報告書では、多くの言葉を形式的な意味と非形式的な意味の両方に用いている。たとえばシンタックスに関するこのセクションでは 形式的な意味で用い、第三章や第四章ではそれ以外の意味に用いるというようにしている。

▼シンタックス▼や▼セマンティックス▼のような用語の形式的な定義に関しては、一般的には意見がまとまっていない。これらの定義については、セクションIVでより詳しく話を進めるつもりである。非形式的な意味では、シンタックスとは構造を記述する機構を通常含んでいる ある言語のよく整ったステートメント (well-formed statement) を詳しく吟味することであり セマンティックスとはこの statement をいかにして 実際のあるいは理論上のコンピュータによって実行するかについて詳しく調べていくことであろう。

この報告書全体を通じて次のように定義する。すなわち 集合  $a^*$  はアルファベット  $a$  の symbol から成るすべての有限記号列の集合であり、言語  $\mathcal{L}$  はこの  $a^*$  の部分集合である。どのような記号列が言語  $\mathcal{L}$  の中にあるか ( $\mathcal{L}$  のシンタックス) は syntax metalanguage で記述する。syntax metalanguage は手続きであり、そして  $\mathcal{L}$  の記号列を作り出すためのアルゴリズム (synthetic syntax), あるいは  $a^*$  の任意の要素が  $\mathcal{L}$  に含まれることを認識するアルゴリズム (analytic syntax) を記述するものである。syntax metalanguage 中の state - ment はふつう production と呼ばれる。

有効な (nontrivial な) analytic syntax を利用するプロセスはすべて syntax-directed であると言われる。

アルファベット  $a$  に属する symbol は端記号 (terminal symbol) と呼ばれ、第II章ではこれらを  $T, T_1, T_2$  などせいうように表わすことにする。syntax metalanguage はアルファベット  $a$  に属さない非端記号 (nonterminal symbol)  $\eta$  の集合を含みうる。非端記号は

その言語を定義するために用いられる。この非端記号は ALGOL report では常に括弧<と>に囲まれて登場し、同じ report の formal syntax rule 中に出ているものである。

この章ではシンタックスについて もっと形式的にとり扱いたいので、上記の括弧を省略し nonterminal をあらわすのには ラテン大文字の U, V, および Z を用いることにする。さらにシンタックスに関するこれらの章では、以下に述べるような、やや広義の術語を必要としている。

語彙  $\mathcal{U}$  は、terminal symbol の集合  $\mathcal{a}$  と nonterminal symbol の集合  $\mathcal{V}$  の和集合として定義する。記号  $S, S_1, S_2$  等々は  $\mathcal{U}$  の構成要素を示すのに用い、一方 記号の列 (要素を1つももたない空な列  $\Lambda$  を含む) は ラテン小文字の  $u, v, w, \dots$  を用いて表わすものとする。

$Z = xy$  がひとつの記号列である時  $x$  は  $Z$  の head であり、 $y$  は  $Z$  の tail であるという。

$U_i \rightarrow u_i$  の形の production の有限集合が次のような性質を持つものとして定義できるとき、その文法  $G$  を句構造文法 (phrase structure grammar) と呼び、これによってある言語  $\mathcal{L}_G$  を規定することにする。その性質とは、

- (1) 各  $U_i$  が空でない任意の記号列であって、その記号列に含まれる記号が語彙  $\mathcal{U}$  に属している。
- (2) 各  $U_i$  は任意の1個の nonterminal symbol であって  $U_i \in \mathcal{V}$  である。
- (3) distinguished symbol  $Z$  と呼ばれる  $U_i$  が唯一個あり、これはいかなる  $u_i$  の中にも現われない。

$U_i, u_i$  はそれぞれ production  $U_i \rightarrow u_i$  における left part, right part と呼ばれる。

図1(a)は  $Z = \langle \text{program} \rangle$  というひとつの文法の例である。図1(b)は 同じ句構造文法に対して Backus Naur Form (BNF) の表記法を用いたものである。ここでは  $\nabla \rightarrow \nabla$  のかわりに  $\nabla ::= \nabla$  を使用し、metasyymbol  $\nabla \mid \nabla$  は 同じ left part に いくつかの異なる right part が対応する時 これらを分けるのに用いられるのである。

$\langle \text{program} \rangle \rightarrow \perp E \perp$

$E \rightarrow T$

$E \rightarrow E + T$

$T \rightarrow P$

$T \rightarrow T * P$

$P \rightarrow (E)$

$P \rightarrow \perp$

$\langle \text{program} \rangle ::= \perp E \perp$

$E ::= T \mid E + T$

$T ::= P \mid T * P$

$P ::= \perp \mid (E)$

(a) formal な記法

(b) BNF の記法

Nonterminal symbol :  $\langle \text{program} \rangle$  E T P

Terminal symbol : I ( ) + \*  $\perp$

図 1 句構造文法の例

いくつかの研究をふりかえてみると 常にそうであるとは言えないが, right part が空であるような production が認められていることがある。

句構造文法がどのようにして ある言語を定義するかを示すためには, より詳細な定義を与える必要がある。  $w = x U y$  で かつ  $v = x u y$  であるような記号列  $x, y$  がある時 ( $x, y$  は空であってもよい),  $U \rightarrow u$  という production の適用によって  $v$  は  $w$  の direct derivation であるという ( $w \Rightarrow v$  と書く)。

$\nabla \Rightarrow \nabla$  の transitive closure は  $\nabla \xRightarrow{*} \nabla$  という記号で表わされ,  $w = w_0$ ,  $w_0 \Rightarrow w_1$ , ...,  $w_{i-1} \Rightarrow w_i$ ,  $w_i \Rightarrow v$  という記号列  $w_0, w_1, \dots, w_i$  ( $i \geq 0$ ) が存在する時,  $w \xRightarrow{*} v$  となる。この時  $v$  を  $w$  の derivative と呼び,  $w = w_0$ ,  $w_0 \Rightarrow w_1$ , ...,  $w_{i-1} = w_i$ ,  $w_i = v$  という連続を  $w$  から  $v$  の derivation (導出) という。

distinguished symbol  $Z$  からの derivative を Sentential form という。言語  $\mathcal{L}_G$  は文 (sentence) の集合として定義される。例えば terminal symbol だけから成る sentential form の集合は,

$$\mathcal{L}_G := \{ x \mid Z \xRightarrow{*} x \text{ and } x \in a^* \}$$

と表わされる。

図 1 の文法  $G$  において  $\mathcal{L}_G$  は算術式の集合

(演算子として  $+$  と  $*$ , 括弧として  $($  と  $)$ , operand として  $I$  を使用) である。この算術式の始めと終りは記号  $\perp$  によって明示してある。

(図 1 の文法によれば)  $\perp P + T * P \perp$  という sentential form は次に示すような小なくとも 2 つの derivation を持っている。

$$\begin{aligned} \langle \text{program} \rangle &\Rightarrow \perp E \perp \Rightarrow \perp E + T \perp \Rightarrow \perp T + T \perp \\ &\Rightarrow \perp P + T \perp \Rightarrow \perp P + T * P \perp \end{aligned} \quad (2.1)$$

$$\begin{aligned} \langle \text{program} \rangle &\Rightarrow \perp E \perp \Rightarrow \perp E + T \perp \Rightarrow \perp E + T * P \perp \\ &\Rightarrow \perp T + T * P \perp \Rightarrow \perp P + T * P \perp \end{aligned} \quad (2.2)$$

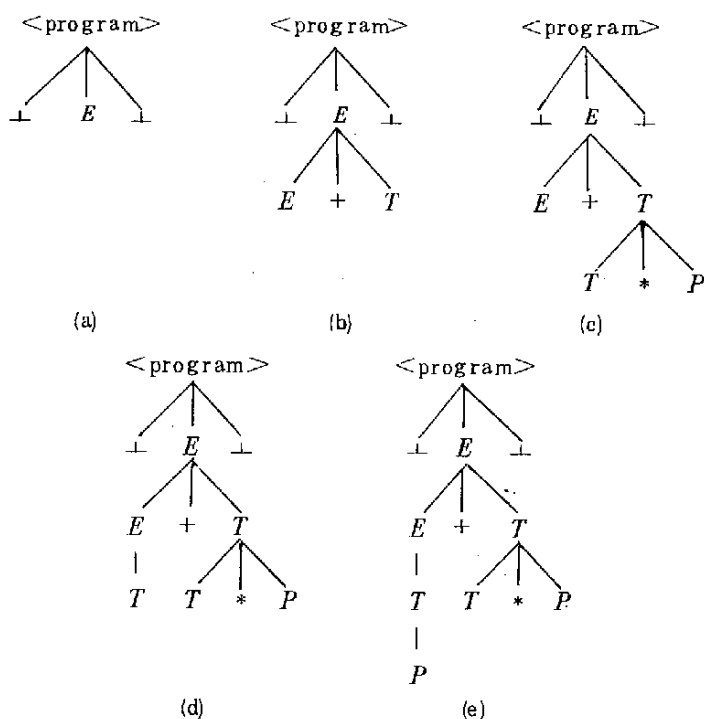


図 2 Syntax tree

derivation は syntax tree によって図示することができる。例えば derivation (2.2) を書きあらわしてみると次のようになる。

記号  $\langle \text{program} \rangle$  から始まって枝 (branch) は図 2 (a) のように描かれる。枝とは node  $\langle \text{program} \rangle$  から 下向きにのびている線の集りで、これらの線の終りの node (label) を含めたもののことである。連結されたこれらの node は、derivation において  $\langle \text{program} \rangle$  と置き換えられた記号列  $\perp E \perp$  を形成する。このようにして続けていく。すなわち production  $U ::= x$  (各 direct derivation) を適用することにより、置き換えられた node  $U$  より 枝を描く。この枝はその node が置き換えるべき記号列  $x$  を形成しているものである。このことは図 2 が derivation (2.2) に対する一連の syntax tree を図示するものであり、特に図 2 (e) が derivation を完了した形を表わすものであることを示している。ここで注意すべきことは、連結されている時に tree の end node (下方へ枝の出ないもの) が 最終的な sentential form を作るということである。

$\perp P + T * P \perp$  はふたつの derivation を持ち、ふたつとも同じ syntax tree を持っている。このとき、derivation は production が direct derivation として適用される順序

についてのみ 異なっているのである。文法  $G$  は、 $\mathcal{L}G$  の文の中で  $G$  に関して2つ以上の syntax tree を持つようなものがひとつでも存在する場合に  $\nabla$  あいまいである  $\nabla$  (ambiguous) と言われる。

任意の総合的な句構造文法が与えられる時、derivation によって その言語の文を自由に作り出すこともでき、また distinguished symbol  $Z$  から syntax tree を作ることもできる。ところがコンパイラは、これとは反対の問題をかかえている。それは ある文  $x$  と文法  $G$  とを与えられて  $x$  の derivation を作り上げ、対応する syntax tree を見い出すということである。このことは 文を解析する (parsing) とか、認識する (recognizing — この言語は recognizer に関連して出てくる) とか、分析する (analyzing) と呼ばれる。

parsing の方法には top-down と bottom-up と呼ばれる二つの異なる方法があり、このふたつは時として混同されることがある。その理由は、ひとつには 人によって syntax tree の書き方が異なるからである。例えば図2(b)の syntax tree (これが我々の表記法であるが) は、図3のようにも書けるのである。そして いまひとつの理由は 前記のふたつの方法が、recognizer が複雑になるにつれて、実際にひとつのものになってきたからである。このことについては、この節のこのあとの部分と 第II章C5においてふれることにする。これらのふたつの方法のどちらも共に left-right と呼ばれる。というのは 文の中で記号を処理する順序が、可能なかぎり左から右へと進行するからである。

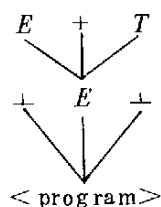


図3 Upside-down syntax tree

純粋な top-down recognizer は、完全に goal-oriented である (goal に依存している)。main goal は もちろん distinguished symbol  $Z$  であり、prediction は認識される記号列が実際にある文になるように作られている。したがって 先ず第一に、記号列をある production  $Z \rightarrow S_1 S_2 \dots S_n$  の right part の  $S_1 S_2 \dots S_n$  に還元 (reduce) することができるか否かを確かめなければならない。

このことは次のようにして行なう。すなわち  $S_1$  が terminal symbol であれば、その production の適用が有効なものであるためには その記号列は terminal  $S_1$  から始まらな



くてはならない。 $S_1$  が nonterminal である場合は、subgoal を設け、これをテストして、その記号列の head が  $S_1$  に還元できるかどうか調べる。そしてこれが正しいことがわかると、 $S_2$  をテストし、同じ方法で  $S_3$  以降を調べていく。ある  $S_i$  に適合するものが見つからない場合には、別の production  $Z \rightarrow S_1' S_2' \dots S_m'$  を適用して試してみる。

同じようにして subgoal  $U$  もテストする。すなわち production  $U \rightarrow S_1 S_2 \dots S_n$  が適用できるか否かをテストするのである。このように、新しい subgoal が適当でない場合は、失敗であったことをもうひとつ高いレベルに報告し、この高い方のレベルで別のものを試みるようにする。

left - right top - down recognizer においては時々左回帰 (left recursion) が生じて困ることがある。すなわち  $U_1 ::= U_1 x$  という形の production は無限のループをひき起すことがありうるのである。この理由は、 $U_1$  が新しい subgoal になった場合、次のステップで新しい subgoal として再び  $U_1$  を作り出すことである。しかしこの左回帰の問題は、時には文法を変えることにより、また時には (次のように) recognizer を修正したりすることによって解決されることもある。同一の left part に対応するいくつかの異なる right part のテストの順序がここでは重要な役割りを果たことになる。左回帰を含む right part がひとつしかない場合には、テストの際にこれを一番最後にすればよい。しかしこうすると他の順序の問題、例えば right part の短いものを一番さいごにテストするというような問題に矛盾することもありうるのである。

top - down recognizer という名前は syntax tree が構成される方法に由来している。すなわち解析のあらゆる点においても (おそらくは正しくないものもあるであろうが)、一番上の node から出てその記号列に至るというように tree を作り上げていくのである (図 4 (a))。このような recognizer は各段階で作られる関係を予知 (predict) しようとするので、しばしば predictive と呼ばれる [Kun 62 を参照]。

top - down recognizer は、回帰的なサブルーチンや 1 個の stack を用いて動く単一のルーチン、その他多くの異なる方法でプログラムすることができる。ここで重要な特徴は、これらが goal - oriented であるという事である。

反対に 純粋な bottom - up recognizer は (もちろん、implicit な goal  $Z$  は除くが) 基本的にはいかなる長期的見通しにたつた goal も持っていない。記号列は production の right part である部分記号列 (substring) としてさがされるのである。

次にこれらの記号列は、それに対応する左側のものと置き換えられるのである。図 4 (b) はこのことを示すものである。ここで後の節で必要な術語を導入するため、少々詳しく述べよう

と思う。

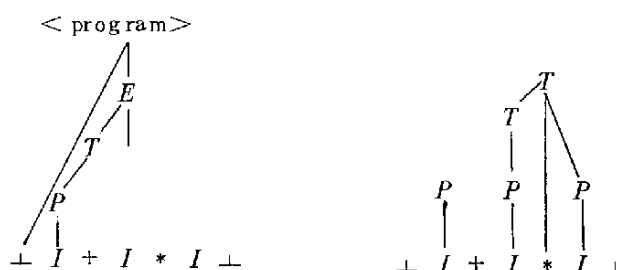


図4 (a) top-down 解析の一部 (b) bottom-up 解析の一部

syntax tree は 我々がある文の解析を始めようとするときには まだないのであるが、bottom-up 解析の背後には それが存在するという概念が在ることは明らかである。したがって、sentential form  $S$  とその syntax tree が存在すると仮定する。 $S$  の phrase とは、その syntax tree のある subtree の end node の集合であると定義する。また ある syntax tree に関係のある  $S$  の handle (把手) とは、それ自身の他には phrase を含まないような最左端の phrase であるとする。例えば、図2(c)の syntax tree には  $P$ ,  $T * P$ ,  $P + T * P$ ,  $\perp P + T * P \perp$  の4つの phrase がある。この中で  $P$  と  $T * P$  のふたつは他のいかなる phrase も含んでいない。handle とは、 $P$  のような最左端の phrase を指すのである。

次のアルゴリズムは、left-right bottom-up 解析の背景にある一般的な規則を指すものである。sentential form  $S = S_0$  (およびそれに対する syntax tree) から始めて、 $S_n = Z$  が作り出されるまで  $i = 0, 1, \dots, n$  に対して次の各ステップをくりかえす。

1. ( $S_i$  についての syntax tree を見ることによって)  $S_i$  の handle を見つける。
2. 対応する枝に名前をつけ、sentential form  $S_{i+1}$  を生み出すような node の label で  $S_i$  の handle を置きかえる。
3. handle をそれから取り去ることによって tree を刈り込む。

すると  $Z = S_n \Rightarrow S_{n-1} \Rightarrow \dots \Rightarrow S_1 \Rightarrow S_0$  という一連のつながりは、 $S_0$  の derivation である。例えば sentential form  $S_0 = \perp P + T * P \perp$  と図2(c)のその syntax tree が与えられる時、handle  $P$  を  $T$  に置き換え、 $S_1 = \perp T + T * P \perp$  とその syntax tree 図2(d)ができる。図2(d), 2(c), 2(b), 2(a)における syntax tree に対する handle は、それぞれ  $T$ ,  $T * P$ ,  $E + T$ ,  $\perp E \perp$  である。図2(c)の syntax tree は2(d)の handle を刈りこむことによって生じ、2(b)は2(c)の handle を刈りこむことにより、そして2(a)は同じ

ようにして2(b)から生じる。この連続は、derivation (2.2)を生み出すのである。

純粹の top-down recognizer と同じように 純粹の bottom-up recognizer は、通常正しくない結果を返す reduction (あるいは結合) をもたらすことがある。これは次のふたつの異なる方法のいずれかで操作しうる。第1の方法は別の可能性のあるものを試しうる点まで逆もどりすることである。(backup, backtracking) これは その記号列の一部を前の形に戻したり、作られていた結合のいくつかを消去することを含むものである。第2の方法は すべての可能な解析列(parse)を並行して作るものである。それらのいくつかが"dead end"(もう reduction や結合が不可能になること)に陥いると、それらを捨て去るのである。

COGENT (Ⅲ章 A.3) は top-down の方法の中でこのやり方を用いている。(〔Kun 62〕を参照のこと)

まちがった reduction の可能性を減ずるために、より複雑な recognizer が開発されてきた。例えば ある修正を施した recognizer では、新しい subgoal から出発する前に ある subgoal のいくつかの derivative が、実際に今問題となっている部分記号列の最初の記号から出発しうるかどうか(look ahead)、あるいは調べようとしている subgoal がずっと先まで形成された tree の一部にできているものかどうか(memory)をみるために、前以って作られたテーブルに目を通すことをする。

修正した top-down recognizer の例は (Ir 61, May 61, War 64)にある。第Ⅲ章Aの syntax-directed symbol processor の多くは、修正した top-down recognizer を用いている。

同様に 修正した bottom-up recognizer は、決定の際に助けとなる可能性のある handle のまわりの文脈を見る。実際 これらの recognizer は、文法にある種の限定をほどこしているので逆戻りや並行に進む方法(parallelism)が不用であるが、そのかわり 大変に複雑なものである。

このような修正は、これらふたつの概念の連合(混同)に役立つものである。ある recognizer については、bottom-up とか top-down とかの区別をすることが大変にむずかしいことがある。例えば Earley のアルゴリズム(第Ⅱ章 B.5)によって作られた production language は、両方の性質を持っている。ある recognizer が出会うはずの explicit な goal や subgoal を持っている時には、これを(修正した) top-down であると言いたい。その理由は goal-oriented であるからである。この問題について詳しくは第Ⅲ章 C.5を参照のこと。

残された術語の多くはコンピュータ・サイエンスの一般的な知識を有する人々には耳新しいものではないだろう。定義の必要な二、三の data structure の用語を用いるので 次に示しておく。

リスト構造 ( list structure ) システムという語は、ポインター ( links ) とダイナミック・アロケーションを使用する いくつかのプログラミング・システムを述べるのに用いる。任意の二つの node の間に 1 つしか道 ( path ) のないリスト構造が tree である。一般的な結合が explicit に許されていて、各々の要素がいくつかの ( 大ていは 2 つの ) link を含む storage のあつまりとなっているようなリスト構造は plex とよばれる。我々は逐次的な情報をとりあつかうための一般的規則として LIFO ( last - in - first - out ) と FIFO ( first - in - first - out ) という言葉を用いる。

#### [ 参 考 文 献 ]

#### REFERENCES FOR II. A

Che 64c, Chom 63, Flo 64b, Gin 66a, Ir 61, Kun 62, Naur 60, War 64.

## B. Automatically Constructed Recognizers

この節では 文法に定められた言語の文を解析 (paring), あるいは 認識 (recognizing) するためのいくつかの実用的な方法を述べ 評価する。"実用的な"方法とは、コンパイラを作成するのに今まで用いられてきた、あるいは現在用いられているもののことである。このような recognizer は、基本的なルーチンの集合によって使われるテーブルの形 あるいは ある言語によるプログラムの形であろう。ここで述べられている recognizer は、任意の適当な文法からそれを作り出すために constructor と呼ばれる関連するアルゴリズムを持っている。結局、これらはすべて 1 個の LIFO スタックを持って働き、逆戻りの機能を持たない left - right recognizer である。

automatic generation という特徴はコンパイラ作成者にとって大変に重要なものである。多くの constructor では、実際に recognizer を作成する前に その文法にあいまいさがあるかを調べる (これは、はっきりした利点である)。コンパイラの大部分の automatic construction はまた 大した仕事量でなくなり、code の optimization のようなことを考えるのに十分な時間ができるのである。さらに automatic construction はその recognizer が形式的なシンタックスを追いかけることを保証するものである。

不幸にして これらの recognizer とそれらの constructor とは、すべての問題を解決できるというわけにはいかない。第 1 に 現在存在しているシンタックスの形式的概念は、多くのプログラミング言語のシンタックスを完全に記述する習慣がない。第 2 に セマンティックスというもののプログラミング言語において果す役割がより大きく より複雑である。— セマンティックスに正しく適合するようにその文法や、生成された recognizer を変更しなければならないことがよくあるのである。第 3 に ある方法が理論的に正しいと思われても、その使用に必要な制限がそのプログラミング言語の慣習的文法に本質的な変更を要するかもしれないのである。

我々は、ついでに数種の recognizer の "効率" について Griffiths と Petrick が比較を行ったことを記しておく。[ Grif 65 ] この比較は 理論的な面白さはあるが、各 recognizer に対応する Turing machine の効率を主としているので、コンパイラの作成においては実際の価値を持っていない。我々は 次のような点について recognizer を比較したことには興味を抱いた。

- (a) recognizer の使用するスペースはどの位か。
- (b) その recognizer の速さはどの位か。
- (c) その constructor に受け入れられるために普通使っている文法をどれ程変更しなければならないか。

(d) 一旦 recognizer が作られたあとで、その文法によって表現されていなかったセマンティックスやシンタックスの特性をどのようにしてそれに追加できるか。

読者は 我々の比較が形式的に適用されるとき、これらの recognizer をもとにしていることおよびこれらが一般的な観察であるということに注意されたい。例えば bit - pushing や工夫をほどこした速いリストのサーチの方法で、特定 implementation の効率をいぢるしく上げることができるのである。上記の特徴はまた 問題の言語を作り出すコンパイラのタイプなどにも依存している。特に問題(d)に対する答えは ソース・プログラムの内部的な形 またはマシーン・コード自体が生成されるかどうかによっているし、そのコンパイラで有効なセマンティックな処理の能力にも依存している。

逆戻りをゆるす bottom-up recognizer を作ることは可能であるにもかかわらず、ほとんど行なわれていない。その文法のあいまいでないことを保証し、任意の sentential form について一意的な handle を効果的にさがし出し、reduction を行なうことができることを保証するために、通常は 制限を設ける。

一方 今日使われている top-down recognizer の多くは、逆戻りを許しているか あるいは 可能な解析をすべて並行して行なっているかである。唯一つの制限は左回帰をゆるさないことである。(第II章Aを参照) 現在存在する top-down recognizer はそのために より広いクラスの文法を受けとることができるが、余り効率の良くない傾向がある。文法が不完全な場合には、逆戻りを許すことは大変に効率のよくない recognizer を作ることになる。

純粋な top-down recognizer については 第II章Aで簡単に述べたので、ここでは論じないことにする。修正した top-down recognizer を用いるコンパイラについて詳しくは [War 61] を参照されたい。さらに [Che 64c] はコンパILINGにおける top-down recognizer の使用に関する指導的な論文であり、[Flo 64b] もこの方法についての十分な内容を含んでいるものである。

ここに記述する recognizer のあるものは 多くのコンパイラの中で多くの人々によって用いられているものであるが、ここではこれらのすべてを参照できるように列挙することはできない。各 recognizer に対して、その recognizer と constructor が述べられている論文名を挙げておくことにする。ここで述べられているシステムの形式的性質と同様に機械的に組み立てられうる理論的に面白い recognizer については第II章C節で簡単に述べることにする。

第II章Aの図1の文法を例として この節全部を通して用いることにするので、読者は第II章A節の定義を簡単に見直しておかれる方がよいかもしれない。

### B. 1. Operator Precedence (Floyd [Flo 63])

この文法は、最初 演算子文法に限定する。すなわち 任意の記号列  $x$  および  $y$ , nonterminal  $U_1, U_2$  に対して、 $U \rightarrow x U_1 U_2 y$  の形となる production があってはならない。このことは、2つの隣りあった nonterminal を含むような sentential form はないということの意味している。これは決して重大な制限ではない。なぜなら 多くのプログラミング言語の文法はすでにこの形になっているし、この形でないプログラミング言語の文法の大部分のものをその言語の構造を本質的に混乱させることなく 演算子文法に変えることができるからである。

ある演算子文法があるとして、 $S$  がその sentential form であるとする。 $S$  の prime phrase とは、自分自身以外の phrase を含まず 少なくとも1つの terminal  $T$  を含むような phrase であると定義する。(これを第II章Aの phrase と handle の定義と比較してみると)。例えば 図5(a)において phrase は  $P$ ,  $T * P$ ,  $P + T * P$ ,  $\perp P + T * P \perp$  であり, prime phrase は  $T * P$  である。同様に図5(b)において prime phrase は  $P + T$  であり, 図5(c)では  $\perp E \perp$  である。ここで作られる recognizer は、各ステップにおいて handle ではなく、最左端の prime phrase を還元 (reduce) する。しかしこれも基本的には左から右へと進むので bottom-up, left-right recognizer と呼ぶことにする。

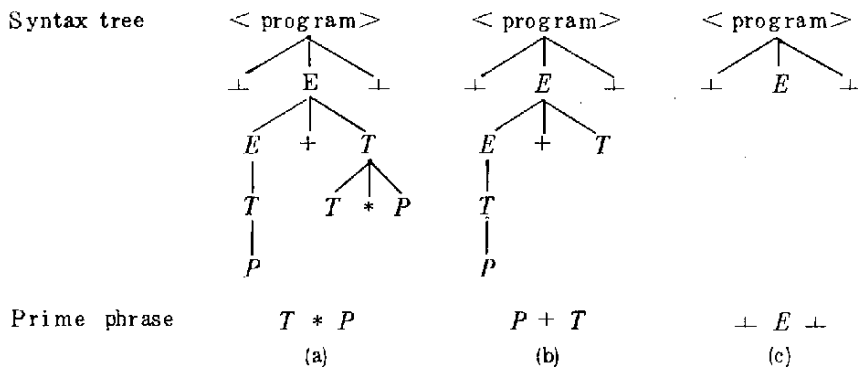


図5 Operator Precedence を用いた解析

同様に  $x$  が少なくともひとつの terminal を含む かつ  $U \rightarrow x$  という production が存在するか、あるいは  $x' \xRightarrow{*} x$  であって derivation  $x' \xRightarrow{*} x$  において適用される唯一ひとつの production の形が  $U_i \rightarrow U_j$  である時  $U \rightarrow x'$  という production が存在するかのいずれかである場合、またこの場合にかぎり  $x$  は少なくともひとつの sentential form  $S$  から成る prime phrase である。

文  $T_1 \dots T_m$  の解析の間中、ひとつの LIFO スタックが部分的に還元した記号列  $S_0, S_1 \dots S_i T_j T_{j+1} \dots T_m$  の記号  $S_0, S_1 \dots S_i$  を含んでいる。各ステップにおいて、(1)  $S_i$  がそのスタック中の最左端の prime phrase の tail symbol であるかどうか、あるいは(2)  $S_i$  が tail ではなく  $T_j$  がスタック中に push されなければならないかどうかを記号  $S_{i-1}, S_i, T_j$  を見て答えることができる必要がある。

このことを行なうために、演算子文法の terminal symbol  $T_1$  と  $T_2$  の間に次の3つの関係を定義する。

1.  $U \rightarrow x T_1 T_2 y$  あるいは、 $U \rightarrow x T_1 U_1 T_2 y$  の production が存在する時  $T_1 \dot{=} T_2$  である。ここで  $U_1$  は nonterminal である。
2. production  $U \rightarrow x U_1 T_2 y$  が かつ ある  $z$  と  $U_2$  に対して derivation  $U_1 \Rightarrow z T_1$  あるいは  $U_1 \xRightarrow{*} z T_1 U_2$  が存在する時  $T_1 > T_2$  である。
3. production  $U \rightarrow x T_1 U_1 y$  が存在し かつ ある  $z$  と  $U_2$  に対して derivation  $U_1 \Rightarrow T_2 z$  または  $U_1 \xRightarrow{*} U_2 T_2 z$  が存在する時  $T_1 < T_2$  である。

terminal symbol  $T_1, T_2$  の任意の順序対 (ordered pair) の間にただか1個の関係 (relation) が保たれている時、その文法を演算子順位 (operator precedence) 文法とよび、その言語を演算子順位言語という。

演算子順位言語においては これらの一意的な関係は、還元され得る部分記号列 (すなわち prime phrase) を見出すために用いられる。 $T_0 x T$  が sentential form  $S = x_1 T_0 x T x_2$  の部分記号列であり、この部分記号列  $x$  内の terminal symbol は 順番に  $T_1, T_2, \dots, T_n$  ( $n \geq 1$ ) であるとする。また  $T_0, T_1, \dots, T_n$  および  $T$  の間に次の関係があると仮定する：

$$T_0 \leq T_1 \dot{=} T_2 \dot{=} \dots \dot{=} T_n > T$$

( $x$  の nonterminals は、ここでは何の役割も果していないことに注意すること)。

この時  $x$  は prime phrase である。さらに ある  $U$  に対する  $x$  の reduction は、常に sentential form  $x_1 T_0 U T x_2$  を生み出すように実行されるのである。

文 (またはプログラム) の解析は、大変素直なものである。(図6を参照) terminal スタックの一番上にある記号  $T$  と 次に入ってくる記号  $T$  との間に  $T_n > T$  の関係が存在するまで、記号をそのスタックに push down する。その記号列が確かにその言語の文であれば、スタックの一番上の要素は 前述のように記号列  $T_0 x$  を持っている。この関係を用いて、 $T_0$  と  $x$  のはじまりを見つけるためにスタックを逆向きにサーチする。そこで  $x$  は prime phrase でありスタック内に  $T_0 U$  を作り出しながら、ある  $U$  に還元される。 $T_0$  と  $T$  とを比較することによ



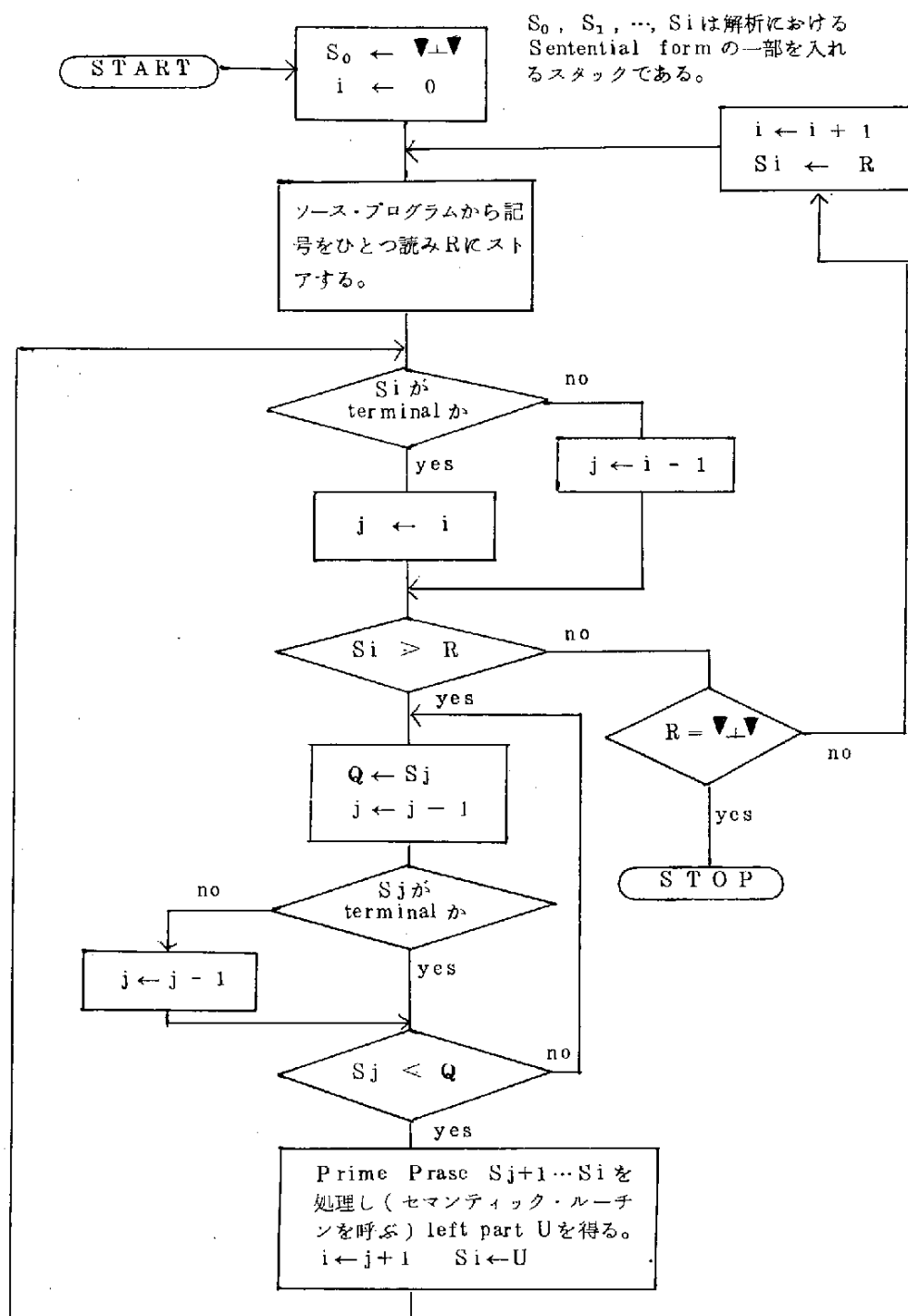


図 6 演算子順位を用いた recognizer

てこの過程をくりかえす。

例を挙げれば, sentential form  $\perp P + T * P \perp$  は, 図5に示されるように (図1の文法を用いて) 解析される。ここでは 各 tree は, その直前の syntax tree の prime phrase を刈りこむことによって導出される。

|         |                   | 表 I     |        |        |
|---------|-------------------|---------|--------|--------|
| $T_2$   | $( I * + \perp )$ | $T$     | $f(T)$ | $g(T)$ |
| $T_1$   |                   |         |        |        |
| )       | > > > >           | )       | 5      | 1      |
| I       | > > > >           | I       | 5      | 6      |
| *       | < < > > > >       | *       | 5      | 4      |
| +       | < < < > > >       | +       | 3      | 2      |
| (       | < < < < =         | (       | 1      | 6      |
| $\perp$ | < < < < =         | $\perp$ | 1      | 1      |

いま 1 をこの文法の terminal symbol の個数とすると,  $>$ ,  $=$ ,  $<$  からなる関係は  $1 \times 1$  の行列の中に置くことができる。( [ F10, 63 ] ) によれば ALGOL 式の言語に対するこの行列は,  $35 \times 35$  ぐらいである。) 行は スタックの一番上にある terminal symbol に対応し, 列は入ってくる記号に対応するものと定義する。すると 行列の要素の間の関係をテストすれば, 前述の比較を行なうことになる。

もし,  $T_1 \leq T_2$  が  $f(T_1) < g(T_2)$ ,  $T_1 = T_2$  が  $f(T_1) = g(T_2)$ ,  $T_1 \geq T_2$  が  $f(T_1) \geq g(T_2)$  を含むような二つの integer precedence function  $f(T)$  と  $g(T)$  が見つければ, この関係に必要なスペースは長さ 1 のふたつのベクトルにまで減らすことができる。

Floyd は 順位関係 (precedence relation) の行列を見出すためのアルゴリズム およびそれらが存在するならば, そしてそれらが存在するときに限り,  $f$  と  $g$  を見出すアルゴリズムの概略を述べている。図1の言語に対する順位行列と関数とは表1のように作られる。

関数  $f$  と  $g$  を用いる時には, よいエラーの復旧法を図示することは かなりむずかしい。というのは, エラーは可能性のある prime phrase がひとつではないという場合にのみ 検出されるからである。したがって この関数はこれに先立つ pass で完全なシンタックス・チェックを行なったとき はじめて用いることができる。(あるコンパイラでは, 実際にプログラムを2回解析 (parse) する。第1回の解析では 完全なシンタックスのチェックを行ない, また変数やブロックなどについての global な情報を集めることもできるようにする。第2回の解析では

効率のよい演算子順位の方法 — すなわち 関数 — を用い、第1の解析で集めた情報を code generation のために用いる。しかし コードがより簡単に generate でき、第2の解析を不要とするために syntax-checker にソース・プログラムの変化した形(逆ポーランドや tripl など)を作り出させる方法を採用傾向がある。)

この方法についてのひとつの障害は、この言語が まだ あいまいな文を含みうることである。解析樹(parse tree)の構造(structure)は、その文法が演算子順位であるかぎりあいではないが、node の名前はあいまいでないというわけにはいかない。production の right part が一意的であるという制限がないので、ひとつの prime phrase  $x$  に対して それを還元できる二つ以上の nonterminal が存在しうるのである。この障害は nonterminal が通常 セマンティック・ルーチンによってとりあつかわれ、シンタックスでは それ程 扱われないということによって大部分解決できる。すなわち シンタックスは構造を定義するものであり、ある node が例えば  $\langle \text{integer expression} \rangle$  と名付けられているか  $\langle \text{real expression} \rangle$  といわれているかは セマンティックスの話である。

セマンティック・ルーチンは prime phrase が認識され、それが還元されうる場合に呼び出される。異なった prime phrase の各々を処理するために、別々のルーチンが書かれるのである。このことはもちろん 実行されるべきセマンティックな処理に依存して、時には文法の変更を必要とすることもある。例えば production  $\langle \text{cond} \rangle \rightarrow \text{if} \langle \text{be} \rangle \text{ then} \langle \text{expr} \rangle \text{ else} \langle \text{expr} \rangle$  は、test や jump がセマンティック・ルーチンによって 正しい位置に挿入できるように次のごとく変更する：

$\langle \text{ifcl} \rangle \rightarrow \text{if} \langle \text{be} \rangle$

$\langle \text{if-then} \rangle \rightarrow \langle \text{ifcl} \rangle \text{ then} \langle \text{expr} \rangle$

$\langle \text{cond} \rangle \rightarrow \langle \text{if-then} \rangle \text{ else} \langle \text{expr} \rangle$

しかし 変形された文法は、多分 その言語のもともとの参照される文法とは本質的に異なっていない。(例えば [Flo 63] の Floyd の言語を見てみること。)

我々の知るかぎりでは、どんなコンパイラも機械的に組み立てられたこの種の recognizer を持っていない。けれども 順位(precedence)の方法自体は、ALGOL, MAD, FORTRAN のコンパイラの中で 度々用いられている。この方法はわかりやすく、柔軟で大変に効率のよいものである。

## B. 2. Precedence (Wirth と Weber [Wir 66c])

Wirth と Weber とは、Floyd の順位の概念に変更を加えた。その文法は 必ずしも演算

子文法に制限することなく、その関係  $\odot$ ,  $\ominus$ ,  $\otimes$  が 任意の記号  $S_1$ ,  $S_2$  のすべての対の間に成り立つ。

1. production  $U \rightarrow x S_1 S_2 y$  がある時,  $S_1 \ominus S_2$  である。
2. production  $U \rightarrow x U_1 S_2 y$  (または  $U \rightarrow x U_1 U_2 y$ ) が存在し, ある  $z, w$  に対する derivation  $U_1 \xRightarrow{*} z S_1$  (そして  $U_2 \xRightarrow{*} S_2 w$ ) がある時,  $S_1 \otimes S_2$  である。
3. production  $U \rightarrow x S_1 U_1 y$  とある  $z$  に対する derivation  $U_1 \xRightarrow{*} S_2 z$  の時,  $S_1 \odot S_2$  である。

記号  $S_1$ ,  $S_2$  の任意の対の間にただかひとつの関係が成り立つとき, そしていかなる production も同一の right part を持たない場合に, この文法を順位文法 (precedence grammar) とよび その言語を順位言語 (precedence language) という。順位言語の文は, 一意的な syntax tree を持つ。解析の途中で スタックの一番上の記号  $S_i$  と入ってくる記号  $T$  の間に  $\odot$  か  $\ominus$  かの関係が成り立つ時には,  $T$  をスタックに push down する。  $S_i \otimes T$  の場合には,

$$S_{j-1} \odot S_j \ominus \dots \ominus S_{i-1} \ominus S_i$$

の状態をさがして, 下向きにスタックをサーチする。この時 handle は  $S_j \dots S_i$  であり, 一意的な production  $U ::= S_j \dots S_i$  の left part  $U$  によっておきかえられるのである。

このやり方と Floyd のものとの主な違いは, この方法では その関係 (relation) は terminal symbol 間に成り立つだけでなく 任意の2つの記号の間に成り立ちうるという点である。それゆえに, prime phrase ではなく handle が還元されるのである。順位に関する行列と関数  $f, g$  を作り出すアルゴリズムは Floyd の場合とよく似ている。これについては [Wir 66c] を参照されたい。

表 II

| $S_2 \backslash S_1$ | $E' E T T P ( I * + ) +$ | $S$     | $f(S)$ | $g(S)$ |
|----------------------|--------------------------|---------|--------|--------|
| $E'$                 |                          | $E'$    | 1      | 1      |
| $E$                  |                          | $E$     | 2      | 2      |
| $T'$                 |                          | $T'$    | 3      | 2      |
| $T$                  |                          | $T$     | 3      | 3      |
| $P$                  |                          | $P$     | 4      | 3      |
| )                    |                          | )       | 4      | 1      |
| $I$                  |                          | $I$     | 4      | 4      |
| *                    |                          | *       | 3      | 3      |
| +                    |                          | +       | 2      | 2      |
| (                    |                          | (       | 1      | 4      |
| $\perp$              |                          | $\perp$ | 1      | 1      |

図1の文法については、 $+ \ominus T$ ,  $+ \otimes T$ ,  $\perp \ominus E$ ,  $\perp \otimes E$ ,  $( \ominus E$ ,  $( \otimes E$  の関係が成り立っている。このような矛盾は、文法を次のように変えることによって処置することができる。

$$\begin{aligned}
 \langle \text{program} \rangle &\rightarrow \perp E \perp \\
 E' &\rightarrow E \\
 E &\rightarrow T' \\
 E &\rightarrow E + T' \\
 T' &\rightarrow T \\
 T &\rightarrow P \\
 T &\rightarrow T * P \\
 P &\rightarrow ( E' ) \\
 P &\rightarrow I
 \end{aligned}$$

この文法に対する順位行列と関数を表IIに挙げておく。実際に 任意の句構造文法をその句構造 ( phrase structure ) を曲解することなく、任意のふたつの記号の間に たかだかひとつの順位関係が存在するように修正することができる。production の右側が一意的でない時には、あいまいさが目だつのである [ Michael Fisher, ハーバード大学 ]。右側が複数個あるという問題は、その文法があいまいでない時でさえ このやや実際的でない問題をひきおこすのである。

Floyd の recognizer の場合と同じく、この場合も順位行列と関数  $f, g$  のいずれをも用いることができる。ここで使う行列は 任意の2つの記号間に関係が成立するので、Floyd の場合よりも ずっと多いのである ( ALGOL に対して  $70 \times 70$  )。ここでも handle が見つかった場合だけ、セマンティック・ルーチンが呼び出される。

理論的には、この方法は非常に安全で 効率のよいものである。その関係が任意の2つの記号間に成り立ちうるということからFloydの場合よりも信頼性が高いように思われる：すなわち順位文法においては、各文に対して一意的な canonical parse ( 正準解析列 ) が存在することがわかっているからである。しかしながら 実際問題としては、一意的な right part に関する制限がまだ残っている。二つ以上のやり方で還元されうる1個の handle に対する各セマンティック・プロセッサは、文脈とグローバルな情報から  $x$  を置きかえる正しい nonterminal を決定しなければならない。このことは次のような production が必要なことを示している。

$$\begin{aligned}
 \langle \text{array identifier} \rangle &\rightarrow \langle \text{identifier} \rangle \\
 \langle \text{procedure identifier} \rangle &\rightarrow \langle \text{identifier} \rangle
 \end{aligned}$$

さらに ある文法が順位文法である以前に、それが平均的なプログラミング言語であるとして巧みにとり扱わねばならない。その理由は、順位関係を決定するのに 十分な文脈を用いられないことである：すなわち 二つの記号の間に二つ以上の関係が存在することが多々あるからである。したがって（上の例のように）中間的な production を挿入するとか、時には その文脈に依存する異なった記号（例えばカンマ）を用いたりすることが必要である。後の場合には、prescanner をその文脈に注意するように変更し、どのような内部記号（internalsymbol）を各カンマに対して用いたらよいかを決めなければならない。最終的な文法は、その言語に対する参照（reference）として プログラマに示すことができなかった。

### B. 3. Extended Precedence (McKeeman [McKee 66])

McKeeman は順位行列を2つのテーブル — すなわち ひとつは handle の tail をさがすためのものと、もうひとつは handle の head を見つけるためのもの — に分け、次に 順位の矛盾がほとんど起らないようにするために もっとよく文脈を調べる recognizer を作ることによってWirthの概念を拡張した。したがって この constructor は、より広いクラスの文法を受けとることができるのである。

- (a) スタックの一番上から2個の記号  $S_{i-1}$ ,  $S_i$  および次に入ってくる記号を  $T$  とし、 $T$  をスタック内に置くべきかどうかを決定するため、そして  $S_i$  が handle の tail であって reduction を始めるべきかどうかをきめるために用いる。
- (b) 同様に handle の最初の記号を見つけるために スタック内に戻るときにも、2つではなく 3つの記号が用いられる。

この方法は Eickel et al. が提唱した方法 [Ei 63] と比較すべきである。第II章C.1を参照のこと。実際には これは異なる triple の個数が多すぎ (PL1 のひとつの dialect (方言) に対して 10000 以上)、また多くの場合 すべきことを一意的に決定するのに 2つの記号で十分である。

McKeeman の recognizer は、可能なかぎり Wirth の2個のアーギュメントの順位を用い、必要があれば triple に切り替える折衷的なものである。スタックが handle を含むかどうかを調べるために右側を見る場合には、エントリ  $\odot$  ( $\otimes$  または  $\ominus$ ),  $\oslash$ , および  $\otimes$  ( $\oslash$  および  $\ominus$  または  $\otimes$ ) をもつ行列 MATRIX1 が用いられる。スタックの一番上の記号  $S_i$  と入ってくる記号  $T$  との間に  $\otimes$  が成り立つとき、3つのアーギュメントを持つ次のような関数の値を見出すために triple のリストをサーチする：

MATRIX 1

| $S_2$   | $S_1$ | $E$ | $T$ | $P$ | $($ | $I$ | $*$ | $+$ | $\perp$   |
|---------|-------|-----|-----|-----|-----|-----|-----|-----|-----------|
| $E$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗     |
| $T$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $P$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $)$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $I$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $*$     |       |     |     |     |     |     |     |     |           |
| $+$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $($     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗ ⊗ |
| $\perp$ |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗ ⊗ |

② という矛盾が生じていないので  
関数 P 1 は不要である。

MATRIX 2

| $S_2$   | $S_1$ | $E$ | $T$ | $P$ | $($ | $I$ | $*$ | $+$ | $\perp$   |
|---------|-------|-----|-----|-----|-----|-----|-----|-----|-----------|
| $E$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗     |
| $T$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $P$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $)$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $I$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $*$     |       |     |     |     |     |     |     |     |           |
| $+$     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗   |
| $($     |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗ ⊗ |
| $\perp$ |       |     |     |     |     |     |     |     | ⊗ ⊗ ⊗ ⊗ ⊗ |

関数 P 2

(いくつかの sentential form の  
有効な部分記号列をも形成している  
必要な triple のみ列挙する)

$P2(\perp, E, +) = \text{TRUE}$   
 $P2(\perp, E, \perp) = \text{FALSE}$   
 $P2((, E, +) = \text{TRUE}$   
 $P2((, E, )) = \text{FALSE}$   
 $P2(+, T, *) = \text{TRUE}$   
 $P2(+, T, +) = \text{FALSE}$   
 $P2(+, T, )) = \text{FALSE}$   
 $P2(+, T, \perp) = \text{FALSE}$

$$P1(S_{i-1}, S_i, T) := \begin{cases} \text{true} & \text{文脈 } S_{i-1} S_i T \text{ において } S_i \otimes T \text{ の場合。} \\ & (S_i \text{ は handle の tail である}) \\ \text{false} & \text{文脈 } S_{i-1} S_i T \text{ において } T_i \otimes S \text{ が成り} \\ & \text{立つとき。} \end{cases}$$

もちろん この関数は 任意の triple に対して値を持たねばならず, constructor はこのこ  
とをチェックするのである。エントリ ⊗, ⊗ および ⊗ (⊗ および ⊗ または ⊗) を持つ 同  
じような行列 MATRIX 2 と関数 P 2 はスタックのその handle の最初の記号を探す時に用いる。

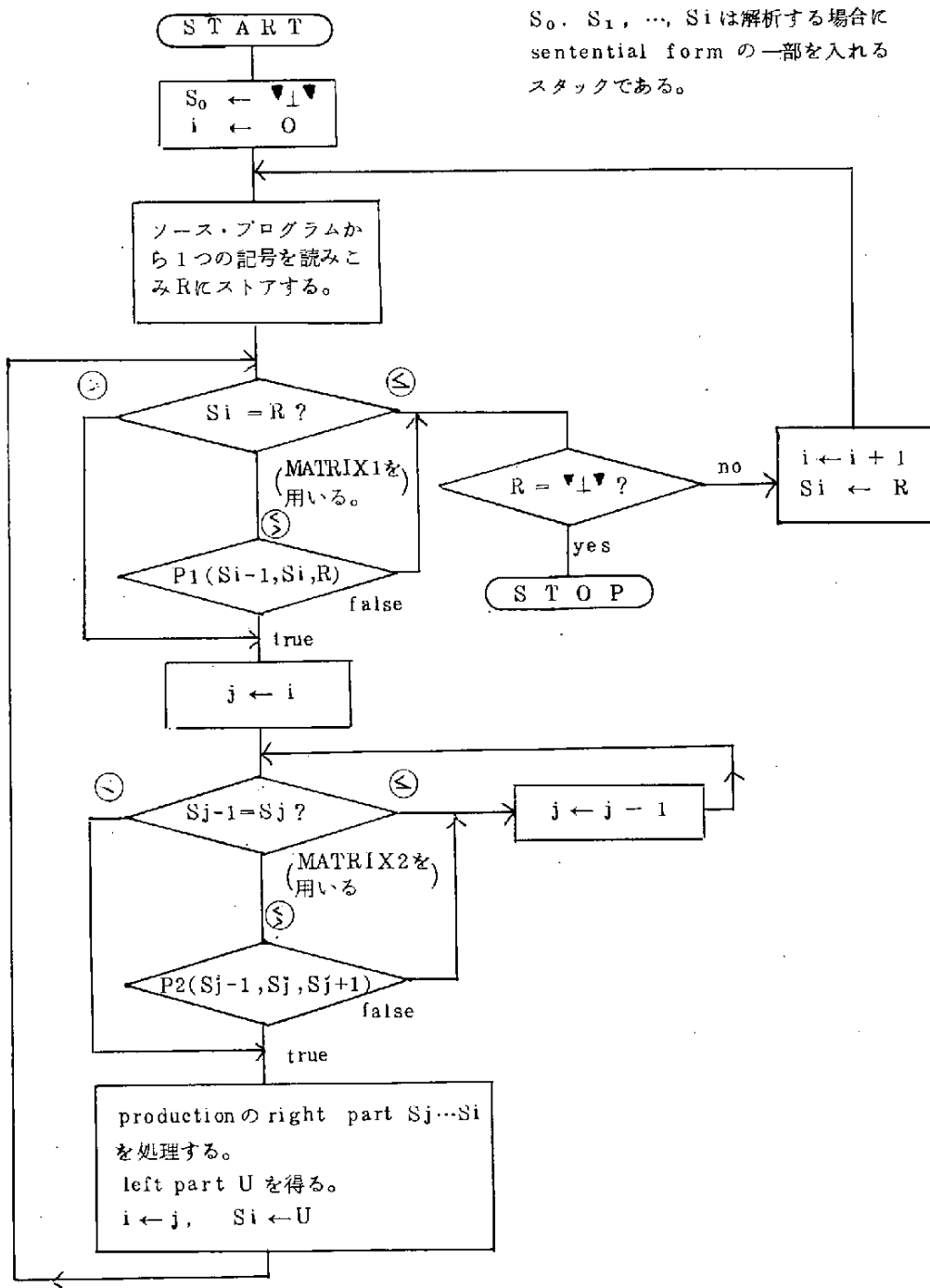


図7 Wirth の順位と McKeeman の triple を用いた recognizer



$$P2(S_{j-1}, S_j, S_{j+1}) = \begin{cases} \text{true} & \text{文脈 } S_{j-1} S_j S_{j+1} \text{ において } S_{j-1} \textcircled{<} S_j \\ & \text{の場合。}(S_j \text{ は handle の head である}) \\ \text{false} & \text{文脈 } S_{j-1} S_j S_{j+1} \text{ において } S_{j-1} \textcircled{\geq} S_j \\ & \text{が成り立つとき。} \end{cases}$$

図1の文法に対する行列と関数は表IIIのように作ることができる。その recognizer は図7に示すごとくである。

tripleの使用は、順位文法で遭遇する不都合な点を避けるのに役立つのである。しかし、ここでも、同じ理由で文法を変更する必要があるかもしれないので、handleが見つかった時のみセマンティック・ルーチンが呼び出されるであろう。

McKeemanはこの方法を用いてIBM 360のPL/1のdialectに対するコンパイラを作成している。行列MATRIX1は約 $90 \times 45$ 、MATRIX2は $90 \times 90$ （各行列の要素は2ビットの長さである）で、約450個のtripleが必要である。今考えられている別のアプローチは、handleのheadを見つけるために用いる $90 \times 90$ の行列を捨ててしまうことである。この時には、スタックの中にhandleが存在する時適用するのに正しいproductionを決めるためにすべての可能なright partをスタックの内容と比較することになるであろう。

#### B. 4. Transition Matrices（遷移行列）

（Samelson と Bauer [Sam, 60], Gries [Grie 67a]）

文を解析する（parsing）ためのこの方法は、Samelson と Bauer によってはじめて紹介されたものである。彼らの叙述では、2つのスタック — operatorのスタックとoperandのスタック — を用いている。文を左から右へと処理し／入ってくるidentifierはoperandスタックにpush downする一方、入ってくるoperator（+、/, beginなど）はoperatorスタックの一番上のものと比較する。reductionが実行できない場合には、入ってきたoperatorをoperatorスタックにpush downする。reductionが実行できる場合には、operatorスタックの一番上の要素とoperandスタック内のオペランドとを用いてサブルーチンを実行し、使ってしまったこれらの要素を消去して、結果のオペランドをoperandスタックにpush downする。演算子順位の方法に似ていることに注目されたい。すなわち2つのterminal（operator）が実行すべき過程を決めるのに用いられ、nonterminal（operand）はここでは何も役割を果たさないのである。特別なoperandスタックというものは、terminalとnonterminalを別々に参照するのを容易にするためにのみ用いられるもので、これは理論的なものではなく、

実際的な考え方である。唯ひとつの本当の相異は、演算子順位の方法では順位による行列を用いるのに対して、遷移行列の方法では サブルーチン名の行列を用いることである。operator スタックの一番上の要素と入ってくる記号によって、実行すべきサブルーチンの名前である行列要素を決定する。次に このサブルーチンは、必要な reduction を行なったり operator あるいは operand スタックに記号を push down したりする。このようにして production は、行なうべき reduction を決定するために各ステップにおいてサーチされる必要がなくなるのである。

遷移行列は、数多くのコンパイラの中で分析的な構文言語 (analytic syntax language) として用いられてきている。Gries は大きなクラスの演算子文法についての遷移行列の recognizer を組み立てる constructor を作成した。組み立てた recognizer が手書きの recognizer に似るように 演算子文法に制限が加えられた。

この constructor は、以下のようなやり方で prime phrase (handle ではない) の始めを見つけるためにテストしなければならないスタック上のいくつかの要素を還元することを始める。いま、

$$\langle \text{cond} \rangle \rightarrow \underline{\text{if}} \langle \text{be} \rangle \underline{\text{then}} \langle \text{expr} \rangle \underline{\text{else}} \langle \text{expr} \rangle \quad (4.1)$$

がこの文法のひとつの production であるとする。これは 次の4つの production に変えることになる。

$$\begin{aligned} \text{"if"} &\rightarrow \underline{\text{if}} \\ \text{"ibt"} &\rightarrow \text{"if"} \langle \text{be} \rangle \underline{\text{then}} \\ \text{"ibtee"} &\rightarrow \text{"ibt"} \langle \text{expr} \rangle \underline{\text{else}} \\ \langle \text{cond} \rangle &\rightarrow \text{"ibtee"} \langle \text{expr} \rangle \end{aligned}$$

これらの中間的な reduction は3つの記号 if <be> then から成る表現 "ibt" をスタックに push down するのを許すためにのみ用いる。同様に if <be> then <expr> else は "ibtee" で表わす。これらの新しい 引用符で囲んだ nonterminal は、 $V, V_1, V_2, \dots$  などで表わすことにする。

解析の任意の段階において、そのスタックは 引用符で囲まれた記号である  $V_j$  ( $j=1, \dots, i$ ) —あるもとの production の right part の head についての表現—および多分最終的な nonterminal operand の  $U$  から成るはずである。したがって 一連の

$$V_1 V_2 \dots V_i (U) T_j T_{j+1} \dots T_m$$

は、sentential form ではなく ひとつの表現 (representation) である。(  $U$  のまわりの括弧は、 $U$  がそこにあってもなくてもよいことを示している。) 一般に 解析の各段階にお

いて存在する可能性のある action は次の 3 つである：

1.  $V_1 V_2 \dots V_i V_{i+1} T_{j+1} \dots T_m$  という reduction を生ずる場合には、 $(U) T_j$  は 他 の right part の head を形成している。
2.  $V_1 V_2 \dots V_{i-1} U T_j \dots T_m$  あるいは  $V_1 V_2 \dots V_{i-1} U_1 T_{j+1} \dots T_m$  のいずれかを作り出す時、 $V_i (U)$  および おそらく  $T_i$  は還元されうる right part を形成している。
3.  $V_1 V_2 \dots V_{i-1} V_i' T_{j+1} \dots T_m$  という reduction を生ずる時、 $V_i (U) T_j$  はある right part の head を形成している。

constructor はそれぞれの  $V$ ,  $T$  の対および  $V$ ,  $U$ ,  $T$  から成る各 triple を上記の 3 つの可能性があるかどうかを見るためにチェックする。各対や triple に対して、このうちのたかだかひとつが存在し 実行されるべき reduction が一意的である場合には、遷移行列とサブルーチンが次のように組み立てられる。各行は引用符で囲まれた記号  $V$  のそれぞれに対して割り当て、各列は可能な terminal symbol  $T$  の各々に割り当てる。行列要素の各々において もし  $V$  が  $V_i$ ,  $T$  が入ってくる記号であるとするならば、実行するサブルーチンの個数が  $M_{V, T}$  にストアされる。対応するサブルーチンは、nonterminal  $U$  の順位を check し 適当な reduction を実行するために組み立てられる。我々は唯ひとつのスタック  $V_1, V_2, \dots, V_i$  を用い、 $U$  を location  $U_1$  におくことにする。図 8 は基本的な recognizer である。図 1 の文法から作られる行列とサブルーチンは表 4 のごとくである。

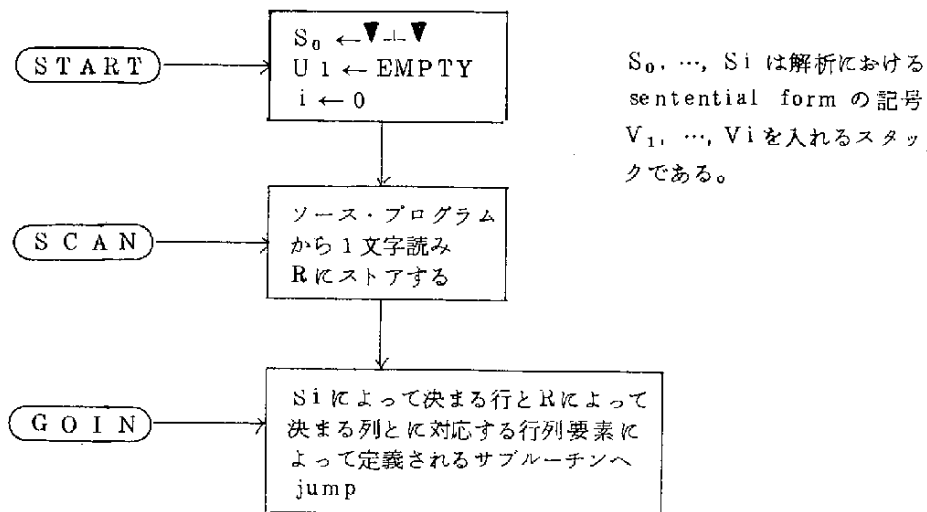


図 8 遷移行列の recognizer

表 IV

| Vi \ Tj     |         |   |   |   |   |   |
|-------------|---------|---|---|---|---|---|
|             | $\perp$ | + | * | ( | ) | I |
| " $\perp$ " | 1       | 4 | 5 | 6 |   | 8 |
| "E+"        | 2       | 2 | 5 | 6 | 2 | 8 |
| "T*"        | 3       | 3 | 3 | 6 | 3 | 8 |
| "("         |         | 4 | 5 | 6 | 7 | 8 |

- 1: if  $U1 = E$  or  $U1 = T$  or  $U1 = P$   
then *SUCCESS EXIT* else *ERROR*;
- 2: if  $U1 = T$  or  $U1 = P$   
then begin  $i \leftarrow i - 1$ ;  $U1 \leftarrow E$ ; go to *GO IN* end else *ERROR*;
- 3: if  $U1 = P$   
then begin  $i \leftarrow i - 1$ ;  $U1 \leftarrow T$ ; goto *GO IN* end else *ERROR*;
- 4: if  $U1 = E$  or  $U1 = T$  or  $U1 = P$   
then begin  $i \leftarrow i + 1$ ;  $S_i \leftarrow "E+"$ ;  $U1 \leftarrow \text{empty}$ ; go to *SCAN* end else *ERROR*;
- 5: if  $U1 = P$  or  $U1 = T$   
then begin  $i \leftarrow i + 1$ ;  $S \leftarrow "T*"$ ;  $U1 \leftarrow \text{empty}$ ; go to *SCAN* end  
else *ERROR*;
- 6: if  $U1 = \text{empty}$   
then begin  $i \leftarrow i + 1$ ;  $S_i \leftarrow "("$ ;  $U1 \leftarrow \text{empty}$ ; go to *SCAN* end else *ERROR*;
- 7: if  $U1 = E$  or  $U1 = T$  or  $U1 = P$   
then begin  $i \leftarrow i - 1$ ;  $U1 \leftarrow P$ ; go to *SCAN* end else *ERROR*;
- 8: if  $U1 = \text{empty}$   
then begin  $U1 \leftarrow P$ ; go to *SCAN* end else *ERROR*;

ALGOLに対する行列は約  $50 \times 40$  であり、約 500 のサブルーチンを持つ。  $U1 = \text{empty}$  であるか否かのチェックは、行列のいくつかの行を二重にすることによって削除される。

( [Grie 67a] を参照 )。一旦 recognizer が作り出されるから いくつかの変更が通常必要であるが、セマンティックスは解析の任意の段階で挿入されるはずであって ( 表 IV の 1 ~ 8 の任意のサブルーチン内 )、right part が認識された時に限りはしない。この文法は多くの変更を要したりはしないが演算子文法でなければならない。この constructor 自体は、コンパイラでまだ用いられていないが、しかし ここで generate される recognizer に大変よく似た recognizer が同じ方法を手書きで行なうことによって作られている。( [Grie 65] )

この方法は多分一番速いものである。一般に スピードが大切な場合にはいつも、テーブルを

切り替えること (Switching table) が用いられる。なされるべき reduction を決め 実行すべきセマンティック・ルーチンを決定する度ごとに production をサーチする必要はないことに注意すること。

この弱点は、用いられるスペースおよびこの方法をインプリメントするのに必要なサブルーチンの個数の多いことである。

## B. 5. Production Language

(Floyd [Flo 61], Evans [Ev 64], Earley [Ear 65])

production language は、Floyd によって紹介され Evans によって修正されたコンパイラ作成のための分析的な (analytic) syntactic metalanguage である。これは、プロダクションの集合から成り、(語 production (構文規則) とは、違った使い方であることに注意すること)。例えば次のようなものである:

$$LO: S_3 S_2 S_1 \mid \rightarrow S_2' S_1' \mid *G \mid$$

これに対するもっと普通の名前は reduction である。というのは、それは記号列を還元あるいは解析する方法を示すものであるからである。

|             |                 |               |            |              |
|-------------|-----------------|---------------|------------|--------------|
| PROGRAMO    | $\perp$         |               |            | *E0          |
|             | $\sigma$        |               |            | ERROR EXIT   |
| E0: T0: P0: | (               |               |            | *E0          |
|             | I               | $\rightarrow$ | P          | *P1          |
|             | $\sigma$        |               |            | ERROR EXIT   |
| E1:         | $\perp E \perp$ |               |            | SUCCESS EXIT |
|             | (E)             | $\rightarrow$ | P          | *P1          |
|             | E+              |               |            | *T0          |
|             | $\sigma$        |               |            | ERROR EXIT   |
| T1:         | T*              |               |            | *P0          |
|             | E+T $\sigma$    | $\rightarrow$ | E $\sigma$ | E1           |
|             | T $\sigma$      | $\rightarrow$ | E $\sigma$ | E1           |
|             | $\sigma$        |               |            | ERROR EXIT   |
| P1:         | T*P $\sigma$    | $\rightarrow$ | T $\sigma$ | T1           |
|             | P $\sigma$      | $\rightarrow$ | T $\sigma$ | T1           |
|             | $\sigma$        | $\rightarrow$ |            | ERROR EXIT   |

図9 Production language の recognizer

スタック上にある文の最初の記号 $\perp$ を置くことから、その文の解析を始める。次に、スタックの一番上のものを第1の“ $\mid$ ”の左隣の記号 $S_1, S_2, \dots$ と比較しながらプロダクションをずっと進んでゆく。

等しいものが見つかり、そのスタック内の一致した記号 $S_1, S_2, \dots$ を記号 $S_1', S_2'$ によって置き換える（もし、置き換えを行わない場合には、矢印“ $\rightarrow$ ”と記号 $S_1', S_2'$ があらわれていない）。記号 $\alpha$ が $S_1$ として現われていれば、これはスタック上のどの記号とも一致することを示している。第2の“ $\mid$ ”のあとに“ $*$ ”があれば、次の入力記号（input symbol）がscanされ、スタックにpush downされることを示す。そのあと、そのproductionの右にある名前を付けられているプロダクション（ここでは $G_1$ ）に行き、再びスタック上の記号との比較をはじめめる。任意のプロダクションにラベルをつけることができる。Earleyは、適当な（総合的syntheticな）句構造文法からproduction languageで書いたrecognizerを作り出すconstructorを作成している。

図1の文法からgenerateされるproduction languageのプログラムを図9に示す。

セマンティックスは、あるプロダクションがgenerateされたあとで、そのプロダクションの任意の行の第2の“ $\mid$ ”の直後にあるEXEC  $i$ の形の“action”を挿入することにより導入される。ここで $i$ は、あるセマンティック・ルーチンの番号である。

production languageの記述は、言語についての決定論的(deterministic)な記述であることを理解することは重要なことである。それは実際に、ある言語の文を解析するrecognizerを作成するためのひとつの言語である。これは普通の句構造文法をもつ場合ではないのである。

production languageやその変形は、多くのシステムの中で用いられてきている。実際にこれを使ってみると、プログラムするのに大変に自然で柔軟性のある言語である。

プログラマはそれでコンパイラを書くのを比較的簡単に学ぶことができる。今までどんなコンパイラも機械的に組み立てられたrecognizerを用いて作成されたことはない。

EXEC actionは任意のプロダクションに入れることができるので、通常は文法にほとんど変更を加える必要がないのである。そしてこのrecognizerはより多くの文脈を用いることができるので、より多くの文法が他の4つのconstructorよりもこのconstructorによって受け入れられることができそうである。

我々は、これでTWSのひとつの部分がほぼ完全であると考える。限定された文脈を用いて、文を解析するための全く異なるたかだか有限個のleft-right recognizerを発明することは可能である。ここに述べたはじめの4個のrecognizerは、使用するプログラミング方法が異な

っているだけである。—すなわち、理論的には、これらは皆第2章Cの限定文脈 ( bounded context ) を用いている。

演算子優先の方法は、これらの中で最もよく知られているものである。これは言語の一部 ぶつ算術式や論理式を認識するのに 度々用いられているもので、IBM 360 ( H-level ) の FORTRAN コンパイラでも使用されている。この方法を用いている この他のコンパイラのドキュメンテーションについては [ Ar 66, Grie 65 ] を見ればよい。 [ Gall 67 ] にも、これについて述べてある。遷移行列の方法は ( といっても、その constructor についてではないが ) いくつかの ALGOL コンパイラを書くのに用いられてきている。特に NELIAC コンパイラや CO-NO テーブルを用いた ALGOR グループによるもの [ Hals 62, Mas 60 ] などである。上記の両方法のいずれも、まぎれもなく数多くのコンパイラの中で使われてきているのである。production language は ALGOL コンパイラで用いられてもいる [ EVA 64 ], けれども他の沢山のコンパイラを作成してきている [ Rov 67, It 66 ] 2つのコンパイラ・コンパイラで重要な部分を占めている [ Feld 66, Mond 67 ]。

他のふたつのコンパイラ・コンパイラのプロジェクトがこの言語を用いており [ Fie 67, Grie 67b ], その独立した変形が [ Che 65 ] その他で用いられてきている。precedence および extended precedence の方法は、主に Wirth [ Wir 66a, Wir 66b ] と McKee-man [ McKee 66 ] によって用いられたものである。

operator precedence, precedence, extended precedence の recognizer においては handle ( もしくは、prime phrase ) が認識されるごとに、還元されるべき handle および実行されるべきセマンティック・ルーチンに対する記号を見つけるために production をサーチしなければならない。同様に Floyd と Evans の analytic production language を用いる場合には、ある reduction が実行されうる前にそのスタックがいくつかの可能性に一致していなければならない。したがって これらの方法はすべて効率のよいテーブル・サーチングが実行されない限り 非常に遅いものとなる。

遷移行列の方法は、可能性のある reduction のそれぞれに対して別々のサブルーチンをコーディングし switching table を用いることによって この問題を解決している。このやり方の欠点は、行列とサブルーチンに費すスペースが多すぎることである。

理論的な考え方の読者に対して、次に もっと一般的で、強力で、かつ複雑なく ( したがって、余り効率のよくない ) left-right recognizer について議論をすすめて行くことにしよう。形式言語 ( formal language ) の理論についての基本的な参考文献は次の節の終りで与えることにする。

[ 参 考 文 献 ]

REFERENCES FOR 11.B

Operator precedence: Ar 66, Gall 67, Grie 65.

Precedence and extended precedence: McKee 66, Wir 66a, Wir 66b, Wir 66c.

Transition matrices: Grie 65, Grie 67, Hals 62, Mas 60, Sam 60.

Production language: Che 65, EvA 64, Feld 66, Fie 67, Grie 67b, It 66,

Mond 67, Rov 67.



## C. Formal Studies of Syntax

### C. 1. Bounded Context Grammars (限定文脈文法)

(Eickel [Ei 63, 64], Floyd [Flo 64a], Irons [Ir 64], WirthとWeber [Wir 66c])

任意の文法は、その handle がその左にある  $m$  個の記号とその右にある  $n$  個の記号によって常に一意的に決定される時 またその時にかぎり、 $(m, n)$  位限定文脈文法 ( $(m, n)$  bounded context grammar) と呼ばれる。解析の間  $\text{left-right recognizer}$  は、可能性のある handle の左については (スタックの中にある) たかだか  $m$  個の記号 右については  $n$  個の terminal symbol を各段階で考慮しながら、 $(m, n)$  位限定文脈文法の sentential form の handle を見出すことができる。

第Ⅱ章 B で論じたはじめの 4 種の文法は、本質的には  $(1, 1)$  位限定文脈文法である。ある意味では、 $(m, n)$  位限定文脈文法  $G$  に対する (bottom-up) recognizer を組み立てることは、 $G$  から同値な文脈依存文法 (context sensitive grammar)  $G_1$  を組み立てることである [Gin 66a, P9]。

文脈依存文法とは、 $xUy \rightarrow xuy$  の形の production (生成規則) を持つ文法である。したがって このような文法においては、 $U$  によって  $u$  を置き換えるということは  $x$  が  $u$  の左にあり  $y$  が  $u$  の右にある場合にかぎり 実行できるのである。その recognizer を作る場合には、その文法が どんな文脈の中で各 reduction がなされうるのかを 明らかに、そしてあいまいでなく述べるようになるまで文脈を各 production の中に左から右へとつけ加えていく (そして このためにより多くの production を組み立てることになる)。

$m > 1, n > 1$  に対する  $(m, n)$  位限定文脈文法の recognizer は、コンピュータの時間とストレージのスペースを法外に要求をするらしい。したがって  $(m, n)$  位限定文脈文法は、コンパイラの中ではほとんど用いられてはいない。

$(1, 1)$  位限定文脈文法に対する recognizer の作成に関しての初期の論文のひとつに [Ei 63] がある。この組立てのアルゴリズムの第 1 段階は、production の right part の長さを 1 つか 2 つにまで変えるために 中間的な production を挿入することである。従って、production

$$U_1 \rightarrow abcd, \quad U_2 \rightarrow abd$$

は機械的に次のように変更されるのである。

$$U_3 \rightarrow ab, \quad U_4 \rightarrow U_3 c, \quad U_1 \rightarrow U_4 d, \quad U_2 \rightarrow U_3 d$$

(第II章B, 4で用いた同様の方法とこれとを対照させてみる)。さて recognizerがスタックの一番上で handle を見つけた時には スタックの一番上の二つの記号  $S_i$  と  $S_{i-1}$  および入ってくる terminal symbol  $T_j$ が そとでとるべき段階を一意的に決めなければならない。そのため 各 triple  $(S_1, S_2, T)$  について次の4つの状態のうちのひとつが必ず成り立たなくてはならない。

1.  $S_1 S_2$  が handle であって,  $U ::= S_1 S_2$  というひとつの reduction が実行されるはずである。
2.  $S_2$  が handle であって,  $U ::= S_2$  というひとつの reduction が実行されるはずである。
3.  $T$  がスタックに push down されねばならない。
4.  $S_1 S_2 T$  は sentential form の部分記号列としては現われない。

triple とそれに対応する action を作り出すアルゴリズムは, [Ei 63] に 例を挙げて出ている。このアルゴリズムと作り出される recognizer は, プログラムされテストされているがコンパイラを書くのに用いられてはいないのである。

一般的な限定文脈文法については, 3つの主な論文が書かれている。そして それぞれの定義する“限定文脈”はわずかずつ異なっているのである。しかし それらの背景にある考えは同じなので, ここではその差異について論じようとは思わない。限定文脈に関する Floyd の論文 [Flo 64a] と構造的結合 (structural connection) についての Irons の論文 [Ir 64] は限定文脈の秘密をより深く探究することに興味ある読者には一読をおすすめするものであるが, いずれも実際に recognizer を組立てるためのアルゴリズムを与えてはいない。Eickel のもの [Ei 64] には, その recognizer と作り方について詳しく書いてある (そしてそのために この論文の方が前の2つよりも読みにくいものである)。

この recognizer は通常のスタックを用いる。[Ei 63] によれば その文法は, 先ず すべての production が長さ1か2を持つようなものに還元される。次に5つの組 (5-tuples)

$$(x; S; y, k, U)$$

を作り出す。ここで  $x, y$  は長さ  $|x| \leq m$  で長さ  $|y| \leq n$  のような記号列であり,  $U$  はひとつの nonterminal,  $k$  は正の整数である。例えばスタックが

$$S_0 \dots S_{i-1} S_i$$

を含み,

$$T_j T_{j+1} \dots T_{j+l-1} T_{j+l} \dots T_m$$

が残りの入力記号 (input symbol) で  $l > 0$  であるとする。個数  $l$  は入力記号列の最初の  $l$

個の記号がその点における右側への文脈として必要とされることを意味している。5-tuple は  $S = S_i$  で  $x$  が  $S_0 \dots S_{i-1}$  の tail,  $y$  が  $T_j \dots T_{j+l-1}$  の head であるようなものが見つかるまでサーチされる。

取られるべき段階は、対応する  $k$  と  $U$  とに依存するものであり、次のごとくである：

$k$                       action

0   stop ... シンタックス・エラー

1   handle  $S_i$  を  $U$  でおきかえる。(  $U \rightarrow S_i$  という reduction を行なう。)

2    $i \leftarrow i - 1$

$S_i \leftarrow U$  ( handle  $S_{i-1} S_i$  を  $U$  でおきかえる )

3   if  $l = 0$  then  $l \leftarrow 1$  else begin  $i \leftarrow i + 1$  ;

$S_i \leftarrow T_j$  ; 入力記号列から  $T_j$  を削除する ;

$l \leftarrow l - 1$  ; end

4    $l \leftarrow l + 1$       ( 右側にもっと文脈が必要である )

文法が  $(m, n)$  位限定文脈であっても 右に  $n$  個の記号が常に必要であるとは限らない、ということに注意すること。recognizer は、取られるべき action を決めるのに必要なだけの文脈のみ 用いるのである。これが前記の個数  $l$  を使った理由である。

Eickel はこの constructor と recognizer とをプログラムし テストしてきているが、この方法を用いてコンパイラを作っていない。この constructor は、 $x$  と  $y$  の長さを 1 に制限してスタートし すべての可能な 5-tuple を作り出している。

もし 同じ  $x$ ,  $y$  と  $S$  および異なる  $k$  (または、同じ  $x$ ,  $y$ ,  $S$ ,  $k$  と異なる  $U$ ) に対して二つ (またはそれ以上) の 5-tuple が存在すれば、その文法は  $(1, 1)$  位限定文脈文法ではない。このような 5-tuple に対しては、矛盾が解決するか 最大値  $m$  か  $n$  になるまで、前述のようにさらに文脈 (と 5-tuple と) を加えながら  $x$  と  $y$  の長さを増加するのである。

Wirth と Weber [Wir 66c] は、記号列  $x$  と  $y$  に対する記号  $X$  と  $Y$  の間の順位の概念を拡張した (II 章 B. 2 を参照)。このため、長さ  $(x) \leq m$  および長さ  $(y) \leq n$  において、 $x \ominus y$ ,  $x \otimes y$ ,  $x \odot y$  を得る。(  $m, n$  ) 位順位文法は もちろん我々の定義する  $(m, n)$  位限定文脈文法でもある。第 II 章 B. 2 によれば、順位文法は  $(1, 1)$  位順位文法である。

## C. 2 Deterministic Pushdown Automata

(決定論的プッシュダウン・オートマシ)

(Ginsburg と Greibach [Gin 66b])

DPDAとは、ひとつのスタックを伴って働く left-right recognizer の概念に関するオートマトン理論の形式化のことである。出発の状態 (start state)  $S_0$  を含む状態 (state) の有限集合  $S = \{S_0, \dots, S_n\}$ , 入力 (input) の集合  $a$  (terminal symbol), 出発記号 (start symbol)  $Z$  を含む集合  $\eta$  (nonterminal symbol に対応している) および写像 (mapping)  $\delta$  があるとする。ここで  $\delta$  は次のとおりである。

$$\delta : (states \times (nonterminal\ symbols) \times (input\ symbols)) \rightarrow (states \times (nonterminal\ symbol\ の記号列))$$

あるいは

$$\delta : (S \times \eta \times (a \cup \{\Lambda\})) \rightarrow (S \times \eta^*)$$

この写像  $\delta$  は (ひとつの値を持つ) 関数でなければならない。このほかに input のどこに現われてもよい空の記号 (empty symbol)  $\Lambda$  をとりあつかうために、制限が加えられる。各段階において我々は triple

$$\underbrace{(S_p)}_{\text{state}}, \underbrace{U_1 \dots U_i}_{\text{stack}}, \underbrace{T_j \dots T_m}_{\text{input の残りのもの}}$$

を得る。初期の triple は  $(S_0, Z, T_1 \dots T_m)$  である。各ステップで  $n \geq 0$  の時  $(S_p, U_i, T_j) \rightarrow (S_q, U_1' \dots U_i', T_{j+1} \dots T_m)$  という写像の助けを借りて、その triple を次のように変えることができる：

$$(S_q, U_1 \dots U_{i-1} U_1' \dots U_i', T_{j+1} \dots T_m)$$

(入力の) 記号列は 最終状態 (final state)  $S_m$  が最終状態の集合  $F$  の要素であれば、受け入れられる (accept される)。

任意の言語 (何らかの文法から導かれる入力記号から成る記号列の集合) は、それが DPDA によって受け取られる時 deterministic (決定論的) であるという。これは deterministic な言語が、あるマシン — DPDA — によって定義されるということであって、その言語を定義している文法の特別な性質によって定義されるということではないということに注意すること。Ginsburg と Greibach は、DPDA や deterministic な言語についていくつかの興味深い特徴を証明している。ここで我々にとって意味のあることは Knuth の LR(k) 言語との関係である。(第 II 章 C.3)

遷移行列を用いて DPDA をインプリメントできることに注意すること。この場合には可能な state  $S_p$  の各々を行で表わし、各 terminal  $T$  を列で表わす。各段階において行列要素  $M_{S_p, T_j}$  は、そのスタックの一番上にある  $U_i$  に従って適当な写像を実行するサブルーチンを決定するのである。

このように遷移行列の方法は、かなり一般的である。しかし Gries によって書かれた constructor (第II章 B.4) は、(1,1)位限定文脈文法のサブセットしか受け取ることができないのである。

### C. 3. LR(k) Grammars (LR(k) 文法)

(Knuth [Knu 65])

任意の文法は handle がその左の記号列全体とその右側の  $k$  個の terminal symbol によって、常に一意的に決定される時に限り LR(k) であるといわれる。これに対応する言語が LR(k) 言語である。したがって スタックを用いて文を解析する場合、left-right recognizer は完全なスタック (the complete stack) (そしてその中の一定個数の記号とは限らない) と、その文でそのあとに続く  $k$  個の terminal symbol を調べればよいのである。これは、それに対してその文法から機械的に作り出されうる効率のよい left-right, bottom-up な recognizer が存在する文法の最も一般的な形である。事実 論究された他の constructor のどれによっても受け取られる文法は、ある  $k$  に対する LR(k) である。したがって LR(k) の条件はまた、今有効であるあいまいさのないことについての最も強力な一般的なテストでもあるのである。

Knuth はある文法が与えられた  $k$  に関して LR(k) であるか否かを決定するために2つのアルゴリズムを与えている。第2のアルゴリズムもまた — もし その文法が LR(k) であれば — 本質的には (上述の) DPDA の形で recognizer を組み立てるものである。これは一見 不思議に思えるかもしれない。というのは スタック中に現われうる無限個の記号列  $S_1, S_2, \dots, S_i$  があり 従って 解析の決定を行なうのに用いなければならない無限個の記号列があることになり、一方 DPDA はスタックの一番上の要素と state と入力記号のみを必要としているからである。

しかしながら production の個数が有限であることから、解析の各段階において スタック中に存在する記号  $S_1, S_2, \dots, S_{i-1}$  には関係なく 有限個の可能な action が存在するだけである。このようにして 可能な記号列  $S_1, S_2, \dots, S_{i-1}$  を "state"  $S_p$  と呼ばれるクラスに分けること、およびある  $k$  に対してその文法が LR(k) であるならば 必要なひとつの値をとる (single valued) 写像

$$(S_p, S_i, T_j) \rightarrow (S_q, S_1', \dots, S_n')$$

を記述することのみ必要である。

もちろん 我々はここで、その概念をわかるようにその過程を非常に簡略化しているのである。

DPDAのスタック中の記号 $S_i$ はもとの文法の記号であるばかりでなく、それ自身その必要な single valued な写像を決めることができるような複雑な場合であるかもしれないのである。

Knuth はまた DPDA によって受け取られる言語 $\mathcal{L}$ のそれぞれに対して、 $\mathcal{L}$ を定義する LR(1)文法が存在することを construction によって証明している。したがって任意の LR(k)言語は、同じ文法を持たないにもかかわらず LR(1)でもある。Earley [Ear 67] は LR(k)文法に対する constructor を書いたが、そのアウトプットは Floyd-Evans の production によく似た形の production である。

#### C. 4. Recursive Functions of Regular Expressions

(Conway [Con 63], Tixier [Tix 67])

Conway は、多くの transition diagram をそのシンタックスが指定されるコンパイラについて論じている。その transition diagram とは、図10のように一定の nonterminal から導かれ得る記号列を認識するものである。

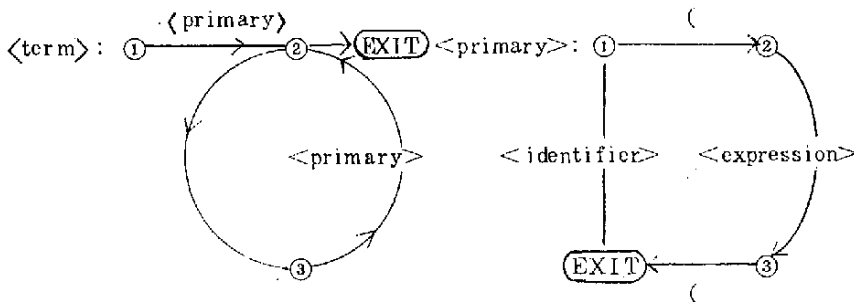


図10 2つの transition diagram

文の解析の各段階において、ある current diagram の current node が存在する。left-right recognizer によって実行される action は (diagram や push down スタックを用いて) 次のように決定される：

(a) 次の入力記号 (1 個の terminal) が current node から出発する線 (line) のひとつの名と一致する場合には、新しい current node までその線を進み 次の入力記号を scan する。

(b) (a) で一致しない場合で ひとつの線が non-terminal の名前を付けられている時には、

その線をたどり current diagram の名前と新しい current node とをスタック上に push down し current node を U と名付けられたその diagram の最初の node に変更する。

- (c) (a) でも (b) でもなく ラベルを付けられていない線があれば それをたどる。
- (d) current node が EXIT である時には、その current diagram で記るされた nonterminal から導きうる記号列が認識できたわけである。したがって その current node (と diagram) をスタックの一番上の要素によって指定されるものに変え、スタックからその要素を削除する。

これは 逆戻りを伴わない top-down left-right recognizer である。(b) の段階において、U から導きうる記号列が認識されるであろうという prediction (予言) に注意すること。この方法は、各サブルーチンに含まれる文字集合 (character set) が大変小さなものであるために、構文解析を効果的に小さな簡単な部分に分け スペースを節約するものである。各 transition diagram が回帰的に他の有限状態のオートマトンを呼ぶことのできる個々の有限状態のオートマトンを表わしているということは明らかである。

Tixier は、この概念を彼の学位論文の中で独立に形式化している。彼は production  $U_i \rightarrow x_i$  を  $U_i = x_i$  という regular expression であると考えた。ここでは union (+), product (=), closure (\*) という演算の集合が用いられている。

このようにして production

$\langle \text{identifier} \rangle \rightarrow \langle \text{letter} \rangle$

$\langle \text{identifier} \rangle \rightarrow \langle \text{identifier} \rangle \langle \text{letter} \rangle$

は、 $\langle \text{identifier} \rangle = \langle \text{letter} \rangle + \langle \text{identifier} \rangle \langle \text{letter} \rangle$

あるいは  $\langle \text{identifier} \rangle = \langle \text{letter} \rangle \langle \text{letter} \rangle^*$

と等しく書くことができる。

Tixier は Euler の 120 の production [Wir 67c] を 7 個の変数から成る 7 つの関数として書き直している。そのうちの 3 つをここに示すが、記号 "( " と " ) " は bracket set expression に対するメタシンボルとして用いている：

*program* = *block*

*block* = begin ( ( new id + label id ); ) \* ( id : ) \* *expr* ( ; ( id : ) \* *expr* ) \*  
end

*expr* = ( cut + if *expr* then *expr* else + id ( ( *expr* ) + ) \*  $\leftarrow$  ) \*  
( go to primary + *block* + catena )

Tixier の constructor は、main goal が有効な有限状態オートマトンを通して できるかぎり解析を行なう時、もし必要ならば最少数の方程式になるように任意の文法の方程式（もしくは production）を操作することができる。次に適切な文法に対して その constructor は その言語の文をあいまいさがなく、そして逆戻りせずに、互いに回帰的に呼び出しあう（各方程式に対してひとつの）有限状態オートマトンの集合の形において 解析できる（modified top-down）recognizer を作り上げるのである。ここで最も困難なことは、他の有限状態オートマトンが呼び出されるかどうかを current node と入力記号からあいまいでないように決定し、もし そういうものがある場合にはどれであるかを一意的に決めることである。受け付けられる文法とその言語とは共に regular context free (RCF) と呼ばれる。

この constructor は実際に効率のよい制限された DPDA を作るものである。したがって RCF 言語は LR(1) 言語である。

## C. 5. Summary

図 11 は、この報告書で論じた個々の constructor によって受け付けられる文法のクラスを inclusion tree の形で示すものである。この tree は いくつかの場所でまぎらわしくなっている。Floyd の production language のようなある種の metalanguage は 強力なものではあるが、このような metalanguage を用いる recognizer の特定の constructor は 受け付けることのできる文法を制限するかもしれない。参考文献が書いてあれば、その node はこの論文で定義されている言語（または constructor）を指すものである。図 11 について次のようなことに注意すること。

- (a) (1, 1) 文法と extended precedence 文法とはいずれも triple を用いるが、(1, 1) 文法の長所は本質的に より広範囲の文脈を許容する自動的な中間的 reduction (automatic intermediate reduction) を実行することである。もちろん 任意の constructor が、これらの中間的 reduction を含むように変えることができるであろう。
- (b) 遷移行列 (transition matrix) 文法は (1, 1) と (0, 1) 位限定文脈文法の間のどこかに落ちることになる。
- (c) 演算子順位の条件が、右側が同一であることを禁ずるよう増大されていると仮定する。そうでないと、この inclusion は成り立たない。演算子順位における行列の方法の利点は、(a) と同様に automatic intermediate reduction を用いることである。
- (d) Feldman は [Feld 64] の中で、各 DPDA が production language で書くことのできる何らかのプログラムに同値であることを示している。



ここでしばらくの間、recognizerをtop-downとかbottom-upとかに分類するという問題に立ちもどってみよう。Tixierのrecognizer(第II章C.4)は明らかにtop-downである、すなわち新しい有限状態オートマトンに切り替える際に、prediction(goal)が作られることがわかるのである。

Knuthの組み立てたLR(k)文法のrecognizerはtop-downであろうか？ある人々は「否」というであろう。このrecognizerは右側のk個の記号と左側に対応するスタック内の記号とを文脈に従ってreductionを行なうだけである。LR(k)自体そのrecognizerも含めて(m,n)位限定文脈recognizerを拡張したものであり、この(m,n)位限定文脈recognizerは通常bottom-upである。さらに、実行されるべきreductionを決めるための助けとなるように余分の“state”がスタックにつけ加えられている。しかしまたある人々はこれらのstateを用いることによってpredictionとかgoalとかが実際に導入されているとも言っているのである。このためにtop-downであるTixierのrecognizerも、Knuthのrecognizerも共にDPDAである。ここでこれらの概念がひとつになってしまい、ある場合にはどちらかになりうるということがはっきりと言えるのである。

コンパイラ作成に関係のあるシンタックスの研究においてこれまでにいくつかの概念を要約してきたが、沢山のものを省略してきている。Gilbert [Gil 66]について少し言及するならば、彼はreductionを行なう際に(もしいくつかの可能性があるならば)どのproductionを用いたらよいかを(sentential formに基いて)示すselection function(選択機能)を文脈依存文法につけ加えたのである。このようにselection functionを用いることによってsynthetic(総合的な)文法をanalytic(分析的な)文法として用い得るのである。Gilbertはこれらの文法とselection functionのいくつかの特徴を証明し、(宣言文の比較と変数の使用を含む)FORTRANやALGOLのような言語のシンタックスを完全に記述することができることを示している。

top-down recognizerについては[Cne 64c]や[Flo 64b]によく出ているので、ここでは詳しく述べないことにした。文脈独立言語(context free language)におけるあいまいさの問題については、それがTWSに関連する時だけ簡単にふれてきた。オートマトンの理論の分野についても関連が深いけれどもここでは述べなかった。

事実、我々は文脈独立文法、オートマトンの理論やマシンについてとり扱っている多くの論文をこの報告書の最後の参考文献から除外しなければならなかった。これらの多くのものや、それ以外のものに対するリファレンスは、J. ACM, Information and Control および[Gins 66a]に出ている。

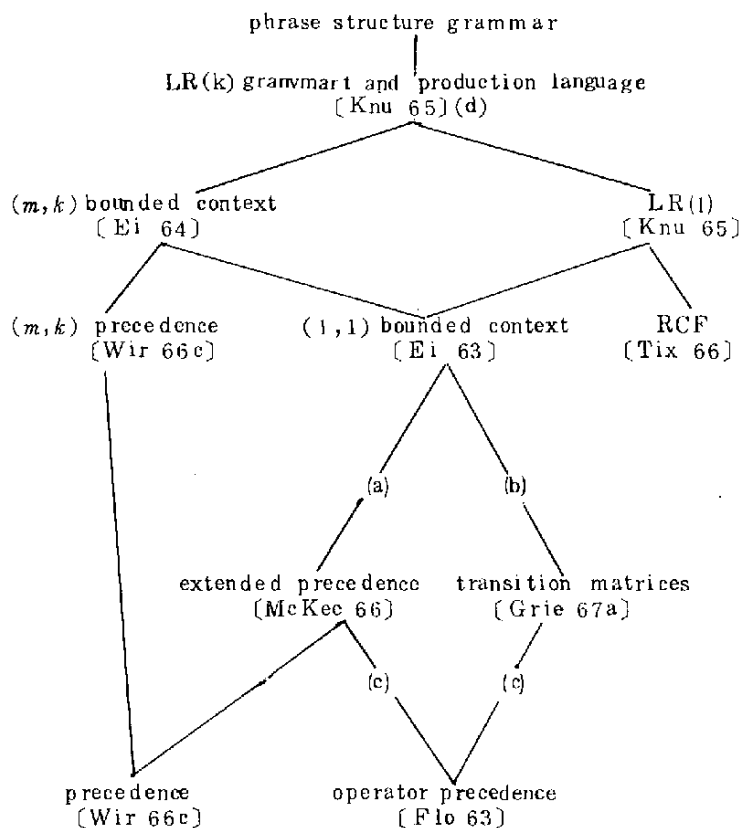


图 11 Inclusion tree

[ 参 考 文 献 ]

REFERENCES FOR II.C

- Introduction to the theory of formal languages: Bar 64, Gins 66a.  
 Pure or modified top-down algorithms: Barn 62, Br 62a, Che 64c, GraR 64,  
 Ing 66, Jr 63a, Kir 66, Kun 62, Rey 65, Scho 65, War 64.  
 Construction of efficient recognizers — sufficient conditions for unambiguity:  
 Ea 65, Ea 67, Ei 63, Ei 64, Flo 63, Flo 64b, (ii) 66, Gins 66b, Grie 67a, Ir  
 64, Knu 65, McKee 66, Paul 62, Wir 66c, Tix 67.  
 Surveys, tutorials on recognizer techniques: Che 64c, Flo 64b, GraR 64.  
 Ambiguity in context free languages: Can 62, Flo 62a, Flo 62b, Gin 66a, Gor  
 63, Lang 64, Ross 64,  
 Thirteen different ways to define languages: Gorn 61.

### III Sementics

#### A. Syntax - Directed Symbol Processor

この章で述べるプログラムは、正確には コンパイラ・コンパイラとは呼ばないものであるがコンパイラ作成に用いられてきたものである。記号処理の仕事としてのこれらの共通のコンパイラ作成のやり方は、これらのプログラムをTWSより以上のものにも 以下のものにもしている。TWSの初めのうちの努力の大部分は この種のものであり、最も特筆すべきものは[ Ir 61]である。このようなシステムは より一般的であるので、第IV章Aのいろいろな nontranslatorの仕事によく用いられる。実際その goal が 最初からもっと一般的であるために、AEDについて([ Ross 67 ])は その節に譲ることとした。

##### A. 1 TMG (McClure [McCl 65])

TMGシステムは Texas Instruments で、シンボリックなアウトプットを作り出す。簡単な1パスのコンパイラを作成する手段として開発されたものである。シンタックスの処理方法は逆戻りを伴う簡単な top-down scan である。しかし そのセマンティック・ルールの埋めこみは、セマンティックな面において いくつかのシンタクティックに可能な goal を削除することによって、その recognizer がより効率のよいものとなるようにしているのである。

基本的なTMGステートメントの形は、スペースによって分離された action の列である。任意の action の先頭に<label>をつけることができ、action の後には その action が失敗であった場合に取りられる<destination>をつけることができる。<action>は syntax recognizer に対する中間的な goal、入力に関する記号列の計算(string computation on the input)、組み込まれたステートメント(built-in)のいずれかである。

これらの action は すべて 中間的な tree を作り上げる際に、トランスレータによって実行されるものである。

実際のコードの出力は、以下に述べるような別のルーチンの集合によってなされる。

primitive MARKS と INSTALL を用いて、入力記号列から作り出される文字をベースとしたシンボル・テーブルがあるとする。このとき 次の例について考える。

INTEGER: ZERO \* MARKS DIGIT DIGIT \* INSTALL

action ZERO \* はすべての leading zero を scan することであり、MARKS は 入力記号列のポインタの現在の値を記憶することである。action DIGIT DIGIT \* は<digit>というクラスに属するすべての文字を scan することである。INSTALL の実行によって、

MARKSのポインターで始まる記号列がシンボル・テーブルに入れられ、それへの reference が中間的な tree(intermediate tree)に入れられる。このテーブルに許されているこのほかの情報は、identifierの属性を記述するのに用いられる宣言されたFLAGS(論理変数)の集合だけである。

組み込みルーチンは 条件付きの算術式、数値変換 そして2, 3の入出力関数を含んでいる。また 入力ポインタを入れておくJ, 入力された最近の記号列の長さを入れておくSYMNRMのようなsystem cellもある。出力もcharacter-orientedであって、例えば次のようである。

LABELFIELD: LABEL=\$(\$P1/BSS/0//)\$

このステートメントは、ある言語のlabelの処理に用いるものであろう。“=”記号は、“\$( ”と“\$)”で囲まれる出力ルーチンを知らせるものである。この出力ステートメントの本体はアセンブラのコード

value(P1) BSS 0

から成る1行を形成する。

記号\$P1は=の左にある第1の要素、たいていは labelのシンボリックな名前を算定するコマンドである。/はタブ(tab)を入れることと BSSおよび0が、その文字自体を表わしていることを示している。最後に//は出力の中にキャリッジ・リターンを置くことを意味する。出力ルーチンはプログラムの中間的なtreeの上から下へと働きかける。したがって 出力ルーチンの中の\$Pnはsubtreeを参照してもよく、このとき\$Pnの算定は 他の出力ルーチンの帰納的な呼び出しを含むことになる。valueによってパラメタを内部のルーチンへ受け渡すことも可能である。この論文では これらの機能の例をいくつか上げており、また TMGのエラーのrecoveryの能力についても簡潔に述べている。

このTMGの仕事は 試験的なプロジェクトであり、その体裁のよくないシンタックスは 定着する(fix)のが容易であった。これはすでに多くのコンパイラを作成するのに用いられており 関連システムTROLは、Knuthによってコンパイラ記述を教えるのに用いられたものである。MULTICSで使ったEPL(Early PL/I)は、TMGの定義を用いて 一層よいコードを得るように2パスシステムとして作成された。TMGシステムは 以下で考えるようなシステムのあるもの程 筋の通ったものではないように思われるが、他のiterationから利益をうけているものである。

## A.2 The META Systems

(Schorre, [Schor 64], Schneider and Johnson [Sch 64])

METAシステムは、ロサンジェルスSIGPLANのsyntax directed compiler に関する

る作業グループの作り出したものである。そのもともとの仕事は多様性に富んだものであったが、現システムは大部分は Schorr が開発した META II として知られるモデルに基づくものである [Schor 64]、このモデルの中ではある言語についての解析と翻訳の過程は BNF のような規則の集合としてすべて述べられている。これらの規則はインプリメントされた時に埋め込まれる code generator と共に回帰的な recognizer となるものである。この規則は、その代りに prefix iteration operator \$ を用いているので、左回帰を禁止している。terminal symbol は引用符で囲み、システムの記号は " . " を先頭につける。そして特別な印のない記号はユーザーの nonterminal である。right part の選択するものをまとめるには括弧を用いる。次の規則は、論理式を翻訳するのに用いる。

1. UNION=INTER('OR'.OUT('BT' \*1) UNION .LABEL \*1|.EMPTY);
2. INTER=BPRIMARY('AND' .OUT('BF' \*1) INTER .LABEL \*1|.EMPTY);
3. BPRIMARY=.ID.OUT('LD'\*) | ('(UNION)');

最後の規則は、算法言語における論理1次子を認識するための手続きを定義するものである。" = " の前にある BPRIMARY という語は この規則の名前を定義し、規則の right part は入力ストリームが論理式として正しいかどうかをテストするアルゴリズムと ( .ID ) が見つかった場合の code generator である。

上記の規則は、多くの META コンパイラで用いられている3つの基本的な要素を含んでいる例である。他の規則の名前 この場合には UNION について、記述は 呼び出されるべき recognizer の回帰的な呼び出しをひきおこす。リテラル・ストリング" ( " が現われることは、入力ストリームが左括弧であるかどうかのテストされるはずであることを意味している。出力ステートメント .OUT は テキスト1行を作り出し、ここでは " \* " は常に primitive な non-terminal .ID によって認識された最近の項目 (item) を指している。

(上記の規則1のような)規則中にある最初の \*1 は label の generation とその label の出力の両方を意味している。同一の規則の中で次の \*1 は同じ label を出力することである、すなわち新しい規則が現われるごとに新しい label が generate され得るのである。これらの label は、その規則が active である時にかぎり存在するものである。他の規則の呼び出しがなされる場合には この label は、スタックに push down される。呼び出した規則から帰ってくると、その前の label が restore される。action .LABEL \*1 は \*1 に対応する label が書き出されることを示している。.EMPTY は 入力に関しては効果を持たないけれども、常に満たされているか 真 (.true) であるかのいずれかの primitive な nonterminal である。

入力ストリーム”(A OR B)AND(C OR D)”については、次のようなコードが作り出されるはずである。ここでLD, BT, BFはそれぞれ Load, Branch True, Branch Falseを示すニーモニック・コードである。

```

LD A
BT L1
LD B

L1
BF L2
LD C
BT L3
LD D

L3
L2

```

META-IIの有用性は 逆戻りの機能のないことと 出力を並べ直すこと(reordering)によって、著しく制限される。METAの方法を完全なTWSへと拡張しようとする試みが度々なされてきている。META-3 [Schor 64]は 基本的なMETA-IIの概念を拡張しようとしたものであり、これによってALGOL 60を7090でコンパイルすることができたのである。これはセマンティックなテストとレジスタの操作について いくつかの機能を与えているが、その付け加えたものは 必ずしも適当であるとは証明できない。META-5は既に多くのフォーマット変換やソース言語からソース言語への翻訳に用いられてきているが、コンパイラのためには使われていない。最も最近 開発されたものは TREE METAであり、これは中間的な syntax treeの複雑な処理を用いるマルチパス(multipass)システムである。METAコンパイラのスピードの遅いことと 効率のよくないことは それらの製作者も認めるところのものであるが、インプリメンテーションのしやすさ、boot-strappingの機能 それに とりあつかうことのできる言語のクラスの広いことなどが、彼らの発展しつつある この仕事を正当化するのに役立つものとなっている。

### A. 3 COGENT (Reynolds [Rey 65])

COGENTシステムはArgonne National 研究所で設計され CDC 3600に対してインプリメントされたものであって、BrookerとMorris(第三章B.3)や Irons [Ir 61, May 61]の考え方に強い影響を受けている。そして LISP COGENTは、大変によく考えられぬ

かれたもので この節で記する他のシステムよりも比較的わかりやすいものである。COGENT コンパイラは、完全にそれ自身の言語で書かれている。bootstrapping を 3 回行なうことによつて、それ自身のコンパイルのスピードは 6 の倍数で速くなっていったのである

COGENT で書かれたシステムは、2 つの部分から成っている。すなわち シンタックスと generator とよばれる処理ルーチンの集合とである。そのシンタックスは、総合的な句構造文法で与えられる。左回帰の文法を含めてほとんどすべての文脈独立文法が受け付けられうる。わずかの制限として right part が空なものはだめである。この文法を用いる recognizer は、並行に処理していく各段階で代わりのもの (alternative) を持っている。修正された top-down の形である。任意の記号列は、recognizer が それに対する一意的な syntax tree を見つけられれば、受け付けられ得るのである。

構文解析 (syntactic analysis) は、中間的な tree を表現するためのリスト構造を作り出す。例えば 次の production

$$\langle \text{term} \rangle ::= \langle \text{term} \rangle + \langle \text{factor} \rangle$$

を用いることによって、 $\langle \text{term} \rangle$  と  $\langle \text{factor} \rangle$  についての subtree に対するポインタを持つリスト・エレメント  $\langle \text{term} \rangle$  が作り出されるのである。

generator (semantic routine) の名前によって任意の production を処理することができる。この generator は その production が tree を作り上げるのに用いられる時、次に実行されるものである。(alternative を並行に処理しているために) 2 つ以上の可能な syntax tree がある場合には、この局所的なあいまいさが解決し 唯一とつての tree が残るまで、これらの generator の実行を遅らせ 構文解析を続ける。そのあと すべての generator が、正しい順序で呼び出される。

例として、次の label をつけた production を考える。

$$\text{Processterm} / \langle \text{term} \rangle ::= \langle \text{term} \rangle + \langle \text{factor} \rangle$$

$\langle \text{term} \rangle$  に対する subtree が root (根) として完全に形成される時、アーギュメントとして  $\langle \text{term} \rangle$  と  $\langle \text{factor} \rangle$  についての subtree を持つような generator  $\nabla \text{processterm} \nabla$  が呼び出される。 $\nabla \text{processterm} \nabla$  は これらの subtree を操作したり、それらを削除したり、それらに対応するコードを作り出したり などするのである。

generator の言語は、リスト処理のオペレーションと失敗のメカニズム (mechanism of failure) を基礎としている。リスト・エレメントは 他のエレメントへのポインタを持つことができ、そのポインタの個数は可変である。リスト・エレメントのタイプは 数値 (固定小数点もしくは浮動小数点)、generator の入口のポインタ、dummy element, identifier element

およびパラメタ・エレメントである。固定小数点数は 任意の大きさであってよく、その大きさを表現するのに十分なワード数を取ることができる。この性質は COGENT の記号数学的な応用に (symbolic mathematics applications) 便宜を与えている。

普通の代入ステートメントに加えて、generator では リスト構造の総合と分析を記述するために、総合的および分析的な代入ステートメントを用いることができる。総合的代入ステートメント (synthetic assignment statement) は、

$\langle \text{identifier} \rangle / = \langle \text{template} \rangle, \langle \text{expression list} \rangle$

の形で、ここで  $\langle \text{template} \rangle$  は パターン・マッチングのために用いるものであって " $:: =$ " のかわりに " $/$ " を伴う括弧内の production と思われる。このステートメントは  $\langle \text{identifier} \rangle$  が  $\langle \text{template} \rangle$  のコピーであるようにし、この中では第  $i$  番目のパラメタ (nonterminal) が  $\langle \text{expression list} \rangle$  内の第  $i$  番目の式の値によって置き換えられるようにする。

例えば 総合的な代入ステートメント

$Z / = (\langle \text{term} \rangle / \langle \text{factor} \rangle * \langle \text{factor} \rangle), X, Y$

を実行する。ここで  $X$  は  $(\langle \text{factor} \rangle / ABE)$  の値、 $Y$  は  $(\langle \text{factor} \rangle / BED)$  の値を持っているとすると、この実行により  $(\langle \text{term} \rangle / ABE * BED)$  のコピーが  $Z$  に代入される。

同様に 分析的代入ステートメントの形は

$\langle \text{test expression} \rangle = / \langle \text{template} \rangle, \langle \text{identifier list} \rangle$

であって、これは式を分解するのに用いるものである。 $\langle \text{test expression} \rangle$  を  $\langle \text{template} \rangle$  に対して照合させる。もしそれが一致すれば、その  $\text{template}$  の第  $i$  パラメタ (nonterminal) に対応する値が  $\langle \text{identifier list} \rangle$  の第  $i$  番目の identifier に代入される。したがって いま  $Z$  が  $(\langle \text{term} \rangle / ABE * BED)$  という値を持つ時、次のステートメント

$Z = / (\langle \text{term} \rangle / \langle \text{factor} \rangle * \langle \text{factor} \rangle), X, Y$

は、 $X$  に  $(\langle \text{factor} \rangle / ABE)$  という値、 $Y$  に  $(\langle \text{factor} \rangle / BED)$  という値を与えるのである。

もし  $\langle \text{test expression} \rangle$  と  $\langle \text{template} \rangle$  が一致しなければ、その分析的代入ステートメントは失敗である (fail), 失敗 (failure) というものは COGENT においては、分岐の方法である。もし すべての条件ステートメントが失敗するような action を含んでいない場合には、その generator 全体が失敗である。この失敗は、何らかの条件ステートメントが現われるまで generator の呼び出しの chain をたどって進んでいくのである。

COGENT のプログラムは、任意個のシンボル・テーブルを用いることができる。action label \$IDENT n は、その production の結果 (文字ストリングのはずである) が シンボ



ル・テーブル  $n$  に入れられなければならないことを指定するものである。それがすでにそこに在れば、その古いコピーに対するポインタが返される。すなわち、任意の指定されたテーブル内で、すべての identifier は、一意的な文字ストリングを持っている。テーブルの各エントリは、その identifier および、通常その identifier の属性を指すポインタ・エレメントから成っている。

出力は 1 個のパラメータ——出力行に置かれるべき文字内部コード——を持つ PUT プルーチンを呼び出すことによってなされる。1 行が一杯になると、それは外に書き出され新しい行が始まる。しかし OUP を呼び出して、ひとつの行が完全に一杯になる前に外にプリントすることも可能である。このほかの primitive な generator STANDSCN( $X$ , PUTP) は、リスト構造  $X$  をその end node によって表わされる記号列  $S$  に写像する。STANDSCN は  $S$  内の記号を見つけ、それらを 1 度に 1 つずつ PUTP に送るものである。

COGENT は明らかに実験的なものであり、いくつかの欠点も持っている。generator に対する言語の構造には、ALGOL が持っているような言語としてのスマートさはない。入力中のシンタックス・エラーがあれば、それは致命的なものとなる。リスト処理は、plex が単にそのシンタックスに基づくだけでなく 任意の plex の創造を含むように一般化されるべきである。COGENT は、記号的な数学における諸問題に適用されてきた。

Reynolds は データ構造のよりよい理論の展開を未解決のままにして COGENT に関する仕事を一時中止していたが、現在は 他の人々と共に研究を続けている。

(第 IV 章 C を参照のこと)

#### A. 4 ETC (Garwick [Gar 64], Gilbert [Gil 66], Pratt [Pra 65])

この章では、syntax-directed symbol processor として記すのによい 3 つの仕事について記述してみようと思う。いろいろな理由で これらのシステムは、前述のシステムの影響を受けてはおらず 詳しくは述べられていないのである。

GARGOYLE システム [Gar 64] は、ノルウェーの Defense Establishment において、Jan Garwick が Control Data 3600 に対して開発したものである。内部報告書としては かなり完全なものがあるが、論文として発行されたもの——これが初めて書かれたものであるが——は、GARGOYLE の大まかな輪かくしか与えていない。

GARGOYLE は、多くの面で TMG (第 III 章 A-1) に似ている。すなわち その recognizer とは、サーチの方向づける機能を備えた top-down である。シンタクティック、およびセマンティックなステートメントはすべて  $\langle \text{label} \rangle \langle \text{else} \rangle \langle \text{next} \rangle \langle \text{link} \rangle \langle \text{action} \rangle$  の 5 つのエンドリからなるテーブル形式で記述される。(逆戻りが暗黙のうちにこなわれたりす

るので), 次にどこに進むかという規則 (sequencing rule) は 大変に複雑であって, 通常は "steering routine" によってなされるのである。

逆戻りのメカニズムもいくつかの変数をスタックすることや, そのほかのものをコピーすることなどを含んだ複雑なデータのとりあつかいを要するのである。5つのエントリはすべて いくつかの方法で用いられる。例えば label は, その行がどのように翻訳されるかをコントロールするコードの代わりになるものである。

次の例は, <代入ステートメント>の処理に用いられるものである。

| (label) | (else) | (next)   | (link) | (action)                                    |
|---------|--------|----------|--------|---------------------------------------------|
| R       |        | ASSIGN   |        |                                             |
| L       |        |          |        | VAR                                         |
| 0       |        | VARIABLE | 1      |                                             |
| 1       |        | NEXT     | 2      | if SYMB=COLEQ then<br>VAR←WORD              |
| 2       | 3      | ASSIGN   | 4      |                                             |
| 3       |        | EXPR     | 4      |                                             |
| 4       |        | RETURN   |        | OUTTEXT ("STO", 10);<br>OUTFIELD (VAR, 20); |

第1行目は このルーチンASSIGNのheaderであり, VARのそのローカルな変数であって, 第2行目で宣言されている。VARIABLEが<変数>を見つけると, その次の記号が, COLEQ (" := ") であるか, このルーチンが失敗であるかのいずれかである。2という行は, 多重に出現した左辺を抜かうためにASSIGNを回帰的に呼び出すものである。最初の逆戻りは (右辺をコンパイルする) EXPRに導き, そして続いて起るリターン (return) は 多重の代入 (multiple assignment) において, 各変数について1回だけ ステートメント4を実行する。ステートメント4の中の<action>VARの続きの値の各々にストアを行なうテキストを作り出すものである。

<action>の列にある言語は, 部分的なword operatorや テーブル・サーチや入出力などのいくつかの決まったルーチンを含むものである。GARGOYLEのユーザーは, しばしば アセンブラ・ステートメントを組み込むことも予想されるし, <action>言語全体は [wir 66] の高レベルの機械語にもよく似ているように思われる。GARGOYLEは その製作者が 複雑な言語 [Gar 66] をインプリメントするのを助けるために用いられてきているが, 広く利用されるのは もう少し何かすっきりした設計と もう少しよい出版物が必要である。

GilbertとMc Lellan [Gil 67]によるTWSは, Gilbertのシンタックティックなアイデアを基として, UNCOL [Ste 61]の概念を復活させようと試みたものである。ソース・プログラムは, 機械に独立な言語BASEに翻訳される。このBASEという言語は, 多くの機械語に翻訳できるものである。最初の翻訳は いくつかの決まったデータ構造の約束を持つマクロ

・アセンブラ言語である中間言語で記述され、第2の翻訳は TMGのようなストリングの形で書かれる。この論文には 沢山のよいアイデアが含まれているが これらはこの章の他のものとそれ程異なっていない。しかし いくつかの悪いアイデアがユニークに説明されている。

machine oriented でありながら、機械には 独立な言語であるという UNCOL の概念は、いつも 言語とコンピュータの多様性において 挫折してきている。Gilbert と Mc Lellan は、BASE において新しいオペレータを定義することと BASE から機械語への各トランスレータにそれを無条件に送ることを許すことによってこれを避けようと試みている。このことは 機械に独立なことを示しているが、どのマクロを選ぶかという基本的な問題は 触れずに残している。この製作者たちは、また 彼らのシステムは“厳密に立脚している”という点を強調している。このことから 彼等は、そのシステムの姿全体を説明するプログラム言語の例として “ $A^m B^n A^m B^n C^m C^n$ ” という記号列の集合を用いる気になったのであろうと思われる。結局 まちがって置かれた厳密の典型的な例として、彼らは ループと go to ステートメントを禁止することによってそのシステムから無限ループを除外したのである。この論文[Gil 67]の中で唯一とつ参照されている他の TWS の文献は [Ir 61] である。

Pratt と Lindsay によって開発された AMOS システム [Pra 65, Pra 66] は、TWS のリスト処理の概念をそのまま応用したものである。ソース言語は、中間言語 (I 言語) に変換 (translate) され、この中間言語が AMOS によってインタープリートされる。この I 言語は 簡単にリスト処理語であって、記号列処理演算と荒っぽい形の算術 — 例えば、 “\* ADD \* P1, P2, P3” は “P1 と P2 を加え、その結果を P3 に入れる” ということを意味するを含んでいる。

この I 言語は 最小の形に設計され、マクロの流儀で拡張できるように設計されたものである。シンタックスは 回帰的呼び出し (recursive call) に重点をおいて production language (第 II 章 B.5) の変形によってとり扱われ、セマンティックスは I 言語で書かれる。AMOS の最も面白い特徴は、プログラムのみならずデータ構造の変換も与えようとしていることである。AMOS は リスト構造のとり扱いかには余り成功しているとは言えないが、この問題は もっと多くの注目を受けるのに値するものである。

#### [ 参考文献 ]

##### REFERENCES FOR III.A

Ab 66, Gar 64, Gar 66, Gil 66, Gil 67, Kirk 65, Ir 61, McCl65, Met 61, Pra 65, Pra 66, Rey 65, Sch 64, Schm 63, Schor 64.

## B. Compile-Compiler

TWSの中でこの集合に属するもののきわ立った特徴は、トランスレータの作成の postsyntactic な面の多くを自動化しようと試みている点である。このようなシステムは compiler writing system と呼ぶのが適当であろう。というのは metalanguage のプロセッサと同様に、翻訳時にも 実行時にも常駐する有意味なプログラムを含んでいるからである。この節のプログラムは 前に述べたものの大部分のものよりも もっと複雑であって、同じようなタイプの仕事に前にたずさわっていたことのある人でないとインプリメントすることに成功したことのないものである。

次に FSL (Formal Semantics Language) [Feld 66] に関する論文からの抜粋を載せる。これは 基本的な原理の概観を述べ、コンパイラ・コンパイラに関する我々の議論に対する適当な序論として役立つであろう。

ある言語 L に対するコンパイラが要求された時には、次のような段階がふまれる。最初に syntactic metalanguage で述べられた  $\mathcal{L}$  の形式的シンタックスがシンタックス・ローダーに与えられる。このプログラムは、言語  $\mathcal{L}$  中のプログラムの認識と解析をコントロールするテーブルを作り出す。次に semantic metalanguage で書かれている  $\mathcal{L}$  のセマンティックスがセマンティック・ローダーに与えられる。このプログラムは もうひとつのテーブルを作り出すが、こちらのテーブルは  $\mathcal{L}$  中のステートメントの意味の記述を含むものである。さいごに 図12 の二重の線の左側のものがすべて捨てられ、 $\mathcal{L}$  に関するコンパイラだけが残るのである。

結果として得られたコンパイラは、1個の LIFO スタックを用いる recognizer に基づく table-driven translator である。この main stack 中の各要素は2つのマシン・ワードから成っている。ひとつはシンタクティックは構成要素であり、もうひとつは 対応するセマンティックな記述を持つものである。

特定の構成要素が認識されるとその semantic word と semantic table が、トランスレータが何をしたらよいのかを決定するのである。

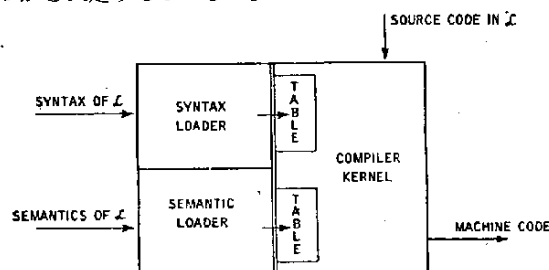


FIG. 12. A compiler-compiler

図 12 コンパイラ・コンパイラ

コンパイラの核心は、入出力 code generation ルーチン およびすべてのトランスレータで使われる  
そのほかの機能を含んでいる。

#### B.1 FSL and Its Descendents (Feldman [Feld 66])

もともとのFSLの仕事が直面した問題は、postanalytic(セマンティック)な処理を記述するための  
言語を開発することであった。適切なセマンティックmetalanguageによって、ソース言語の記述が  
できるかぎり自然であるようにできるはずである。他の人々が定義されているソース言語の意味を理解  
できるためには、それが読みやすくなければならない。効率のよい自動コンパイルができるためには、  
記述が十分厳密で完全でなければならない。そしてさらに そのmetalanguageは、特定のコンパイ  
ラの特徴に依存してはならないのである。

FSLで用いられる syntax metalanguage は、第Ⅱ章B.5で論じた production languageに近いもの  
である。このシンタックス言語の中のステートメントは、nとラベル付けられたセマンティック・ステ  
ートメントを呼び出す "EXEC n" という命令を含んでいる。このほかには シンタックスとセマ  
ンティックスの相互作用は、main stack上のシンタクティックおよびセマンティックな記述の対を作っ  
ていることだけである。

semantic metalanguage FSLは、この仕事の主な焦点であった。ここでは そのいくつかの点を  
詳しく論ずることにする。FSLにおける考え方は あとで述べる TRAGENがcodeの optimization  
を目的としたのと逆に machine independence (機械独立)であるということであった。その計画は、  
コンパイラの核心にうめこまれた primitives の大きな集合によって操作される translation の機械に  
依存する面を持たせることにより機械独立であるようなmetalanguageを作ることであった。この meta-  
language のステートメントは、(コンパイラ・コンパイラで)主に primitive な ルーチンの呼び出し  
から成る機械のコードにコンパイルされる。このアプローチを説明するのに役立つように例をいくつか  
挙げることにする。

シンタックス・フェーズがREALの宣言を処理中であり、スタックの第2番目の位置(LEFT 2)  
の中で宣言された identifierと共にセマンティック・ルーチン1を呼び出すものとする。

```
1: ENTER(SYMB; LEFT 2, (STORLOC+DOUBLE), REAL, LEV);
```

```
STORLOC←STORLOC+2
```

ここで その変数についての記述は、シンボル・テーブルSYMB内に置かれている。その変数につ  
いての項目はその名前、tagをつけたアドレス、REALという語および現在のブロックのレベルであ  
る。最後に STORLOCは、倍精度の変数に対しては2個のcellを割り付ける(allocate)ために  
2だけ増加する。

identifier が算術ステートメントの中で scan されると、セマンティックルーチン 2 が呼び出される。

```
2: IF NOT LEFT1 IS CONSTANT THEN IF SYMB
 (LEFT1, TYPE)=REAL THEN
 RIGHT1←SYMB(LEFT1, SEMANTICS)
 ELSE FAULT1
```

セマンティック・ルーチン 2 では、その identifier (LEFT1 の中にある) が定数であればこのルーチンを終了する。もしそうでないときはその identifier は変数であり、シンボルテーブル内から探し出されねばならない。テーブルの参照 (table lookup) は FSL では、

<table name>[<operand>, <position name>]

の形の独特なテーブル・オペランドを通してなされるものである。テーブル・オペランドについてのこの例では、その第 1 欄が LEFT1 の内容に等しい項目 (もしくは row) を探すためにテーブル SYMB をサーチすることから始まる。次に等しかった row の指定の位置 (TYPE) が選び出され、string construct REAL と比較される。これらが等しい時には、その変数は REAL として宣言されたものですべてがうまくいくのである。この場合には SYMB 内の一致した row の (tag をつけたアドレス) SEMANTICS が、その real 変数のセマンティックスとして代入される。その変数が REAL 型でないか、テーブル中に全くない場合には、ステートメント FAULT1 が実行される。これはコンパイル中のソース言語のプログラムのリスト上にエラー・メッセージをプリントするものである。

最後にそのシンタックスが、コンパイルされる加算 (addition) を認識しセマンティック・ルーチン 3 を呼び出すものとする。

```
3: RIGHT2←CODE(LEFT4+LEFT2)
```

code bracket "CODE( "および" )" は、それらの中のステートメントが翻訳時に実行されるのではなく、オブジェクト・コードにコンパイルされるべきものであることを示している。このステートメントの実行は、加算に対する一連のコードをコンパイルするために、スタックの第 2 と第 4 番目の位置にあるセマンティックな記述 (semantic discription) を用いる code generating routine への呼び出しを作り出すものである。このセマンティックな記述とはデータの型、符号、index 属性、オペランドの現在の location を含むものであって、これらはトランスレータの状態と相まって局所的によいコードを作り出すのに十分である。加算の結果はそれ自身式であり、そのシンタックスはシンタクティックな記号 E をスタックの第 2 の位置においてあると仮定している (図 9 の T 1 + 1 行目を参照)。

R I G H 2への代入は 結果のセマンティックに(例えばアキュムレータ内のDOUBLE) そのセマンティックな記号を結びつける。

F S Lシステムでは、ほとんどすべての要素(construct)をcode bracketの内側(実行時に行なうため)にも code bracketの外側(翻訳時に行なうため)にも、書くことができるのである。

semantic metalanguage F S Lは、コンパイラの作成者が トランスレータを作る際に、さまざまなデータ構造を宣言し用いることを許している。例に示したtableやcellのほかstack, mask, stringがある。このシステムは 翻訳時にも、実行時にも利用できる数多くの補助的なルーチン(例えば、フォーマットや、ファイル処理)を含んでいる。Formula A L G O Lコンパイラは 大部分F S Lで書かれたもので、このインプリメンテーションの記録[It 66]は F S Lの長所と短所のよい研究になるものである。

F S Lの弱点としては、いくつかの便宜の不足と多くの基本的な構造上の欠点のあることが挙げられる。便宜の不足とは index変数、回帰的なサブルーチン、アセンブラ言語の埋めこみ、デバックの手段などのないことであり、その開発の際の主題ともなるひとつの課題である。後のシステムでは 矯正されてきているものである。構造上の欠点は、主に機械独立であることを保とうとすることの結果として起ったことである。

F S Lシステムは、コンパイラ作成者の必要とするものが semantic languageの機能で満たれる範囲まで有効である。しかし これは、適切な概念の妥当な公式化が存在するところにおいてのみ可能である。したがってIV章Cに列挙したresearch problem(例えば データ構造、ページング、並行処理)はすべて どのF S Lの中でも問題である。F S Lは あるconstructを認識すると ただちに作り出すべきコードを要求するものであるというのは、ひとつの一般的なまちがいである。code generationは 無期限に引き延ばすことができるが、現在動いているシステムは グローバルなコードのoptimizationやマルチ・コンパイラに対して特に優れた機能は備えていないのである。

これらの問題は、現行のF S Lのような いくつかのプロジェクトでとり組んでいるものである。しかしながら 機械独立な方法で得ることのできるコードのoptimizationには限度がある。どんなF S Lシステムであっても失敗するであろうという観念がある。すなわちいつも方法が形式化されるということを十分よく理解される以前に、用いられてしまうのである。それでもなおこのようなシステムは 大変に役立つものであり、metalanguageの表現をサーチすることは新しい方法を注意深く研究することにつながるのである。さらに 特定のインプリメンテーションは、通常まだ形式化されていないconstructを扱うための非公式な方法(例えばアセンブラ言

語)を含んでいるのである。

このほかに完成し終っているFSL型のシステムは、Lincoln LaboratoryのVITAL [Mond 67]だけである。VITALタイムシェアリングのもとで動くもので、主にシステムの特徴においてFSLとは異なっている。これらのことは(この記述の中で用いられている)多くの著しい改善とあいまってVITALをより使い易いものにしてはいるが、反面理論的な面白さの少ないものにしてはいる。より有意義な変更の中に、シンタックス・テーブルの大きさを約4分の1に減らしスピードを増す実行可能な(executable) syntactic class nameがある。テキストはすべて、dictionary pointerでつながったブロックに退避される。これは行単位の編集や、再コンパイルの時間を約半分に減らすのに役立つものである。persistent storageとコンパイル時における実行の組合せは、インクリメンタル・コンパイラを作成する時に助けとなるものである。ユーザーはレジスタの割り付け(register allocation)でかなりの融通性を与えられるが、もともとのシステムにおけるようにこの責任を回避することを選ぶこともできる。小さいけれども賢明で重要な変更は、production languageにTESTというsyntactic <action>を加えたことである。このTESTは、セマンティックスによって変数の集合に依存するものである。これはBNFの伝統に違反するものであるが、あるトランスレータには必要なものであることがわかっており、その他のものにも多大な便宜を与えるものである。

FSLシステムはG-20やTX-2などという一般的でない機械でインプリメントされたために、明らかにハンディキャップをつけられている。このことを補うため、現在IBM 360シリーズで別々に3つのインプリメンテーションが進められている。カーネギーのCABALグループ[Fie 67]は、最小限度ALGOLを拡張したセマンティック言語を用いてマルチパス・コンパイラに対するシステムを設計中である。スタンフォードのGries[Grie 67b]のもとでもマルチパス向けであるが、特殊問題向きセマンティック言語を用いる仕事を行なっている。J. CurryによるLincoln Laboratoryの仕事は、多分、VITALによく似たものとなるであろう。これらのプロジェクトはすべてFSLと次の節のTGSの良いところを組みつけようと努力しているように思われる。

## B.2 TGS

(PlaskowとSchuman[Plas 66], Cheatham[Che 65])

TWSの研究における最も生産的なグループのひとつは、現在はApplied Data Researchの一部となっているMassachusetts Computer Associates(COMPASS)という小さな



コンサルティングの会社であった。彼等のTWSは 幾多の変更を受けただけでも、その基本的な観念と彼等の仕事の目的は ほとんど変わることなく存在している。彼等はコンパイルすることを6つの過程として定義した。すなわち 語句 (lexical analysis), 構文解析, 解析列の翻訳, optimization, コードの選択 および出力である。彼等の仕事の主な動機は、実行時の効率を追求することであった。ただし そのほかの考えも その時々で大切な役割を果たしてはいたのであるが。Computer Associates の現在のTWSの関するこの仕事は、コンパイル時の各段階に対して 1個の言語TRANDIRを用いている。TRANDIRは、基本的には 算術部門、パターン・マッチングの部門 (第II章B.5) および多くの組み込み関数から成っている。彼等の仕事の別の面は、TGS内の拡張可能なコンパイラの計画に関して第II章C.5で論じている。

COMPASSにおけるTWS問題の初めの試みは CGSと呼ばれるもので、彼等のこの仕事とは全く異なったものであった。彼らは このアプローチをやめてしまったけれども、これは周期的に再認識されているように思われるので、ここで簡単にCGSについて述べることにする。CGSシステムは、あとの解析で用いるような syntax tree を作り出す top-down recognizer をもとにしている。このフェーズに対する入力は、基本的には BNFであった。第2のフェーズは、GSLという tree-matching の言語を用いて中間的なコードを generate することであった。実際のコードを選択する過程は、第3の言語MDLで書かれた。この仕事が放棄された理由は tree を作り出すのに時間がかかり、それに対してパターン認識を行なうのがむずかしいことがわかったためである。

CGSとは異なり TGSシステムは、この節の他の言語と同じように コンパイラのすべてのフェーズを記述するのに1つの言語を用いるものである。この言語TRANDIRは、TRANGENインタープリタによって処理される interpretive code にコンパイルされる。もし図12のシンタックスとセマンティックスのローダーを組み合わせると、FSLのモデルが TGSにうまくあてはまる。事実 これらの2つの仕事の間には緊密なコミュニケーションがあり、それらは 著しい程度に一致してきた。しかし そのコミュニケーションは完全とは言えなかった。2つの同時期のインプリメンテーションであるTGSとFSLとは、2,300 ヤードの間で 何の接触ももたずに実施されてきたのである。

TRANDIR言語は、FSLで用いられたシンタックス言語 (第II章B.5) と本質的には同じ パターン・マッチングのサブセットを含んでいる。このTGSの version は さまざまのスタックを用いることができるか、記号の同一性とは別な特性を照合することができるなどの一層柔軟なものである。パターン・マッチングの特徴は、いろいろなコードの optimization や

構文解析で用いられるのである。

TRANDIRのそのほかの特徴は、FSLのsemantic languageに酷似している。FSLにおける“semantic word”に類似したものとして、こちらには各シンタクティックなconstructに関連づけられた“symbol description”(SD)がある。そのトランスレータによって用いられるために宣言のテーブル、マスク等々に関して、かなり精巧な機能がある。operatorに結びつけられたこれらのいろいろなstorage methodは、効率のよいコードをコンパイルするために必要な情報を記憶しアクセスするのに非常に柔軟性のある手段を与えている。code bracketについてのFSLの概念は、TGSではコンパイラ作成者が出力コードを作り出すのを助ける一連のsymbol manipulation primitivesによっておきかえられている。TGSコンパイラのオペレーションは、ほぼ完全な例を用いて記述するのがよいであろう。

次の例は[Plas 66]の中に述べられているミニチュア的な算法言語はLtoから取ったものである。ここで選ばれた基本的なコンパイル方法は、COMPASSコンパイラ[Che 66]に共通なテーブル状の中間コードを用いることである。

$$Z \leftarrow X * Y$$

に対する典型的な中間コードへの変換は

① TIMES XY

② STORE Z①

である。中間コードは、後でアセンブルする時のために最終的な出力を生み出すcode selectionフェーズで処理されるのである。

最初に、次のTGSステートメントを考える。

```
...VARAE//EMIT(STORE, COMP(1), COMP(0));
```

```
EXCISE;TRY(ENDST).
```

このステートメントの(//までの)左の部分は、<variable>(変数)と<expression>(式)というタイプのパターンであり、これはmain stack(SYMLIST)と比較される。一致するものが得られた時には、このステートメントの残りの(action part)が実行される。EMITというactionは、一致したスタック内の第0番目と第1番目の要素をオペランドとしてSTOREという中間的な命令を作り出すものである。結果として生じるsemantic description(SD)がないので、EXCISEというactionはスタックから既に一致した2つの要素を削除するのに用いられる。最後にTRY(ENDST)というactionは、TRANGENにENDSTというラベルのついたパターンと一致するかを試すようにさせるものである。

乗算を認識するためには、もう少し複雑なルーチンが用いられる。

```
...VALS*VAL//PHRASE(SYMRES(TIMES,COMP(2),
COMP(0)));
```

```
AESSET;SYNTYP(COMP(0))=AE;TRY(AE1)
```

"\$\*"がterminal symbol "\*"を示すことがわかれば、このステートメントの左の部分は明らかである。SYMRESというactionは、同じパラメタについてEMITを実行し、SDとしてその値を返すようなルーチンへの呼び出しである。このSDはPHRASEのパラメタと成ってそのスタックの一致した部分を置き換えるために用いられる。AESSETというラベルのついたactionは、新しい一番上の要素のシンタクティックなタイプにAEという値を代入するということを生じる。最後にTRY(AE1)というステートメントは、さらに式の処理へ進むものである。

以上2つのTGSステートメントは、逆の順序に実行されれば  $Z \leftarrow X * Y$  を中間言語にコンパイルするものである。実際の世の中では典型的なステートメントは、テーブルのオペレーション、string commands、条件およびもっと複雑なその他のTRANDIRの構成要素を含んでいるのである。さらにトランスレータを読みやすく、また書きやすく改善するかなり複雑な<procedure>という特徴もあるのである。

どのような場合でも中間コード自体は、code selector と呼ばれるTRANGENルーチンのほかの集合により処理される。これらは、前述のsyntax routineと同じ形で書かれるものである。

例えば

```
//TIMES INMEM INMEM...
LOADMQ(XM+1)
```

このステートメントは照合するための記号ではなく、スタックのエントリのpredicateのINMEM(in memoryを意味する)を含むパターンを持っている。(delimiter "//"と"..."は、そのパターンが中間コードのスタックに対して照合されるべきであることを示している。)LOADMQというサブルーチンは、パラメタとしてスタックの第2番目のオペランドに対するポインタを付けて呼び出される。このユーザーの書いたルーチンは、必要ならばLOADMQ命令をアセンブルし、現在MQレジスタにひとつのオペランドが在ることを考えてスタック上のSDを調整する。同じようなルーチンが適当な乗算をコンパイルするのに用いられる。その結果はアキュムレータにあり、TRANGENは結局はこのステートメントに一致するのである。

```
//STORE** *INAC...
IF.SIGN(SYMBOL(ACHOLDS))THEN
```

```

EMIT(CHS);
EMIT(STO, ARG(1));
LINE(TEMPS) = 0;
ACHOLDS = 0; MQHOLDS = 0;
TO(STEP).

```

このパターンに含まれる“\*\*”は常に一致するものであることを示し、“\*”は間接的な参照を意味している。ACHOLDSと記されるアキュムレータ内のオペランドが負である場合は“補数を取る”(CHS)命令が出される。置き換えるべき変数に関してテストをしない場合はいつもストア命令が出される。これに続くactionは、temporariesを取り返しACとMQのレジスタを解放してトランスレータの状態に影響を与える。最後にSTEPというラベルのあるactionに制御を移し、中間コードを通して続けていく。

TGSシステムは既にいくつかのコンピュータについてインプリメントされ、多様なコンパイラの製作に用いられている。そのコンパイラの作成者たちは既に専門家となっていて、形式的なシステム内に留まっている必要がなくなっているのである。TGSの使用は、かれらがそれを商業用コンパイラにひきつづき用いていることから、COMPASSにとって十分価値のあるものである。つい最近の[Che 66]によれば、Cheathamは多分(meta+meta)プロセッサによってTRANDIRのprocedureに変換されるべきものであることを意味する宣言的(declarative) metalanguage  $\mathcal{L}_D$  を用いることを提唱している。

言語の変換は、WirthとEarlyの機械的なconstructorの組み合わせた概念に基づいている(第II章Bを参照)。より強力な言語を考慮して(タイプのチェックなどの) predicate を添え加えたり、宣言的なシンタックスに対しても任意の計算を行なうことができる。最後にシンタックスの規則に付いている中間コードを出力するための規則が存在する。宣言的な言語(declarative language)はまだインプリメントされていないが、cheathamはそれがTRANDIRコンパイラの最初の形式化に役立つことが証明できたと主張している。このことは多分事実であろうが、proceduralな形への変換が現在機械的な処理でないということが予想される。さらにユーザーとしては、 $\mathcal{L}_D$ の方がTRANDIRよりも多少とも複雑でないとは考えられない。

TGSとFSLの主な違いは、企画目的において端的に現われている。すなわちTGSはコンパイラ作成者からより詳しい情報を得ることによって、より柔軟性のあるものとなっている。スタンフォードのGriesとカーネギーのFierst[Fie 66]の仕事は、簡単なcode bracketステートメントとマルチフェーズ処理を許すことによって双方の長所を備えようと努めている。

VITAL [ Mond 67 ] と最も最近の TGS [ Plas 66 ] は相互作用的で、複雑なトレース、編集およびデハッキングの特徴を持っている。

### B. 3 CC

(Brooker, Morris 他 [ Brook 67a, b, c ])

マンチェスタ大学で始められた CC ( Compiler-Compiler ) のプロジェクトは最も古く、また最も孤立した TWS の仕事のひとつである。CC システムは 既に時々動かされ、いくつかの算法言語をインプリメントする [ Cou 66, Kerr 67 ] のに用いられているけれども、発行された書類が不十分で CC は一般的には理解されていない。

この CC の仕事は、セマンティックスの問題に中心を置いている。構文解析は、memory と 1 記号前を見ること ( one symbol look-ahead ) を伴った top-down である。構文解析の結果として、セマンティックス・フェーズで用いられる完成した syntax tree を得る。もちろん これは処理が遅く、非公式には他の方法の用意がしてある。我々はここで、彼等の記法を勝手に変えながら正式な操作を示すことにする。

シンタックス・フェーズへの入力 は BNF に類似しており、これに ( option であることを示す ) angle bracket 内に出現しうる " ? " と ( 左回帰を置き換える ) repeat operation " \* " をつけ加える。次のステートメントは、算術和を取る代入ステートメントのシンタックスを指定するのに用いられるものである。

1. FORMAT [ SS ] = < variable > ← < sum >
2. PHRASE < sum > = < sign ? > < term > < terms >
3. PHRASE < term > = < variable > | < number > | < sum >
4. PHRASE < terms > = < sign > < term > < terms > | < empty >

行 1 は format definition と呼ばれ、行 2 ~ 4 の補助的な phase definition を利用するものである。SS は この format のクラスを指定するもので、これについては後で述べることにする。format definition も phase definition も共に top-down recognizer の " production (生成規則) " として用いられる。そのちがいは、format definition に対応する中間的な tree が完全に作られると 関連する format (セマンティック) routine が それを処理するために呼び出されることである。行 1 に関連する format routine は 次のように書かれうる。

with Line 1 would be written as follows :

5. ROUTINE [ SS ] = < variable > ← < sum >
6. Let < sum > = < sign ? > < term > < terms >

7. ACC ← <sign ?> <term>
8. L1: GO TO L2 UNLESS <terms> = <sign> <term> <terms>
9. ACC ← ACC <sign> <term>
10. GO TO L1
11. L2: STORE ACC IN <variable>
12. END

行5は、この routine が行1の format に関連し、そのシンタックスが行1に一致した時に呼び出されて 適当な中間的な tree を作り出すということを示している。行6は<sign ?>、<term>および<terms>の値として、中間的な tree から descriptor に代入する。アキュムレータを load するコードがコンパイルされたあとで(行7)、その<sum>が二つ以上の<term>を持つかどうかを知るために この tree を調べる。もしなければ このルーチンはストア命令をコンパイルし(行11)、そしてこのルーチン終了する。もっと複雑な<sum>については 行8, 9, 10でとり扱かう。実際のコードの出力は インプリメンテーションに依存し、通常は 簡単な記号列の処理ルーチンでとり行なわれる。

CCにおいて用いられるステートメントには、basic (BS), master (MP) および source (SS) の3つのクラスがある。BSという部分言語は FSLやTGSの semantic sublanguage に対応するもので、これは算法定言語における code generation, リスト処理, 語句の解析ルーチンを含むものである。これらのBSステートメントは、さらに前にコンパイルされているステートメント(precompiled statement) (例えば 上の行6, 8, 10, 12)と、FORMAT[BS]ステートメントとそれらに関連するROUTINEによって定義されるBSステートメントのトランスレータ指定の compilation (例えば行7, 9, 11)の2つに分類できる。BSステートメントは format routine 内のみに出現できる。

MPのクラスのステートメントは FORMAT, PHRASE, ROUTINEステートメントなど自体(行1~5)および編集やシステム・ダンプ命令を含んでいる。これらの構成要素は いずれも format routine 内に出てきてはならない。最後のステートメントのクラスSSは、ソース言語のステートメント自身を含むものである。これらは組み合わせられて、BSとMPステートメントで 事実 第三章Cの意味で強力な拡張可能なコンパイラにCCをするのである。CCシステムはもともとこのようにして働くように設定されているが、実際には 多少異なっている。

例えば FORTRANの定義をFORMAT[SS]とROUTINE[SS]のステートメントの集合として書き、CCがこれらをテーブルにコンパイルするとする。次に CCのこの

更新されたコピーを 伝統的な F O R T R A N コンパイラを生み出すような S S ステートメントだけを受け取るように C C をセットしたスイッチと共に記録する。この方法を用いると C C は、スタックではなく中間的な tree によって、シンタックスとしての F O R M A T および P H R A S E ステートメント、そしてセマンティックスとしての R O U T I N E を結びつけた形として 図 12 のモデルのようになりうるのである。C C が F S L や T G S のように中間的な記号に対して descriptor をとり扱う機能は持っていないということに注意すること。C C は top-down recognizer を用いているので、構成要素は それらが処理されるにつれて用いられる。これは中間的な descriptor を削除するが、これによって A L G O L コンパイラをひとつの大きな R O U T I N E として書かねばならないようになるのである。

C C グループは最近、かれらのシステムの使用と実行に関する報告書を作っている。これらは T W S を手書きのコンパイラと比較するという初めての試みを含んでいる [ Brook 67b ]。

Brooker は 1 年以内で、Atlas Autocode compiler の手書きによるコーディングのスペースに関して 1.6 倍、時間に関して 1.7 倍まで減らすことを可能にした。これらの結果は より多くの情報がないと解釈することはむずかしいけれども、多くの人々が想像している程 コンパイラ・コンパイラが 必ずしも莫大なスペースと時間とを要するものではないということを示唆するものである。このことは Kerr の結果によっても言うことができる [ Kerr 67 ]。

C C は、A L G O L のような言語である Atlas Autocode の中に首尾よくうめこまれてきた [ Brook 67a ]。S P G ( System Program Generator ) と呼ばれる。このシステムの応用はマンチェスターの Morris によって 現在開発中である。S P G は、その基礎となるメカニズムに関する知識を持っているシステム・プログラマを対象としている。C C のインプリメンテーションは 現在 Atlas K D F - 9 , G - 21 に在り、Carnegie-Mellon では I B M 360/67 で仕事が始まっている。

#### [ 参考文献 ]

##### REFERENCES FOR III.B

Design: Brook 60a, 62a, 63, 67b, 67c, Che 64a, 64c, 65, Cou 67, Feld 64, 66, 67, Fie 67, Gric 67b, Mond 67, Mor 67, Nor 63, Plus 66, Ros 64a, War 61, 64.  
Uses: Brook 67a, 67b, Cou 66, It 66, Kerr 67, Nap 67, Rov 67.

## C. Meta Assemblers and Extensible Compilers

これらの形のTWSは、これらが共にマクロの概念を高いレベルのプログラミング言語に拡張しようとしている点で類似している。マクロ・プロセッサの基本的な考え方は、ある記号をそれらに結びつけられているテキストの一部で直接おきかえるということである。ほとんどすべての最近のアセンブラは 大変複雑なマクロの機能を備えているけれども、この考え方について最もよく書かれたものは Strachey [Str 65] および Mooers と Deutsch [Moo 65] による一般的な論文に載っている。メタ・アセンブラと拡張可能なコンパイラ (meta assembler と extensible compiler) は どのようにしてマクロを高いレベルの言語に拡張するかという点で、異なる二つの概念に基づいているのである。すなわち メタ・アセンブラは コンパイラをアセンブラの特殊な場合と考えると近づいて行き、一方 拡張可能なコンパイラは 標準的なコンパイラにテキストの置き換え能力を加えたものとして近づいて行く。これらのアプローチは共に 広く普及しているものであって、ここで考えているものよりも後で多くの論文が発表された。そのためこれらは、この節の終りに列挙しておくことにする。

### C.1 General Discussion and METAPLAN (Ferguson[Fer 66])

Ferguson の論文は カルフォルニア州サンディマスにある ACM Programming Language Conference から 1965 年に提出されたもので、メタ・アセンブラについての優れた紹介を含んでいる。その基本的な考え方は、すべてのアセンブラが 多くの性質を共通して持っていることを観察することから生まれた。

シンボル・テーブルや location counter やマクロのようなものを扱うために、procedure を作ることによって 特定のアセンブラの作成のスピードを上げることができる。特定の機械に対するアセンブラを組み立てるためには ワードの大きさ、数値の表現 その他を指定するはずである。各機械についての出力は format statement を用いてプログラムされ、それは relocation やシンボリックなデバッグ情報を容易に含みうるのである。このようなシステムはアセンブラを書くのに便利で 全く有用なものに思われる反面、どのようにしてこれをTWSに拡張するのかは 明らかではないのである。

TWS としてのメタ・アセンブラの使用は、前に述べた コンパイラというものが マクロ・アセンブラの特殊な場合であるという前提に基づくものである。この前提に対する論議は、マクロ対高レベル言語の論争の生まれ変りのような感がある。マクロ・アセンブラ側は、保守の側に立ち 数においてまさり そのために議論上ではより激烈なのである。そしてこの議論は、アセ



ンブラの定義が定まっていない事から さらに複雑になるのである ([ Fer 66] の論文を参照されたい)。この例は、我々の論文の目的からすれば ひとつ掲げるだけで十分であろう。

Ferguson は、メタ・アセンブラが コンパイラに似たステートメントを いかに操作するかを次のように表わしている。

```
IF F(A) PLUS 5 EQ G(B) GO TO L
```

彼は、many-many macro とよばれる方法を用い IF, PLUS, EQ, GO, TO を(prafix) オペレータとして定義している。このmany-many macro は テキストの置き換えが行なわれている間、state 情報を用いたり 更新したりする性質を持っている。さらにこれは、マクロの一般的用法における主な問題——すなわち、アセンブルの過程におけるグローバルな性質 (state 情報) をいかにして効果的に使用するかという問題——に取りかかろうとするものである。many-many macro はかなり柔軟性に富んでいるので 現在知られているいかなるコンパイラをもインプリメントすることができるが、本当の問題は このmany-many macro がそうすることがよい方法なのかどうか ということである。これに対する答は state 情報を記録し、利用するメカニズムに依るもので、これらはこの論文には論じられていない。

## C.2 PLASMA

(Graham と Ingerman [Gra M 65])

Graham と Ingerman の メタ・アセンブラの仕事は、置換および接合(binding)の問題にその中心を置いている。彼等は、Halpern (次節で述べる) よりはシンタックスに重きをおいていない。なぜならば 彼等は、syntax-directed を彼等のシステムにおけるコンパイラ作成のための糸口と仮定しているからである。([ Ing 66] を参照)

このメタ・アセンブラに対する基本的な入力、リストのリストである "line" である。第1のリストは一般化された label で、これには任意の記号、それらが active である高いレベルの数値、同じくそれらが active である低いレベルの数値が含まれる。第2のリストはオペレーションを含み、第3のリストはオペランドを含むものである。入力は tree に変換され、tree 構造のもとで置換が行なわれる。記号的 (Symbolic) あるいは 数值的 (numeric) な置換を行えるようにするようによって、テキストの置き換えとアセンブラの機能とを結びつけているのである。

彼等は この研究を Cherry Hill の RCA で続けており、恐らくこれに関する論文を 再度発表するであろう。彼等の現在の仕事は、さらに精巧な置換処理を含んでいる。しかしまだ彼らのシステムをどのようにしてコンパイラの基礎として使用するかという点を明らかにするには到っ

ていないのである。

### C.3 XPQP (Halpern [Hal 64, 67c])

XPQPシステムは IBM 7090 でインプリメントされたもので、よい内部仕様書が作られているものである [Hal 67c]。少なくともひとつの言語ALTEXT [Star 65] は、このシステムでインプリメントされている。マクロの呼び出し形式が可能であることを示すために、ALTEXTの中で書かれているプログラムから2, 3の例も掲げることにする。

```
DO THRU LAB1 I=1 TO T-1 BY 1
IF 1 CHAR AT NAMES I IS EQUAL TO NAMES I+1, GO
 TO LAB2
IF U IS LESS THAN 3, GO TO PUT
IF MATCH GO TO (LAB1, LAB2, LAB3), I
COPY 11 CHAR CARDS 1 TO REJECT 1
```

この言語のステートメントはどれも、IBM 7090のアセンブラ言語で書かれたマクロに対する呼び出しである。従って プログラムは自由に、アセンブラ言語を高レベル言語のステートメントの中に使うことができる。XPOPの次のような性質から、上の例のようなステートメントをとり扱うマクロを定義することができるようになるのである。

1. マクロは普通 ステートメントの最初の語(DO, IF, GO, TO)によって認識されるのであるが、ある場合には マクロの名前がline のどこにでも出現しうるのである。
2. 各々のマクロの定義の中で そのステートメントの残りの部分を処理する際に用いられる句読点(punctuation)をそのマクロのパラメタとして定義する事ができた。
3. 各々のマクロの定義の中で そのステートメントの残りの部分の処理のために, noise word (無視されるもの)とkey word (パラメタを決定するのに用いられるもの)を定義することができる。
4. 特定のlabel が現われるまで、ひとまとまりのコードをアセンブルするのを延期することができる。(これは上記のDOステートメントに対する命令をgenerate するのに用いられる。)
5. マクロの定義の中の命令は、たとえばnameのグローバルな属性や条件付きのアセンブラの命令をチェックするような場合はアセンブル時に実行され得る。このことによっては その命令が実行される度ごとにそれらをアセンブルしなければならないということを除けば、(通常のごとく)コンパイラをアセンブラ言語で作成することが可能になるのである。

6. このシステムは多くの有用なトレースやデバッグを助ける機能が備わっている。

XPOPはコンパイラ作成の道具としての見ると、いくつかの短所を持っている。そのひとつは、すべてがアセンブラ言語かあるいは前以って定義されたマクロで書かれねばならないことである。第2に、変数の属性を入れておくシンボル・テーブルなどをインプリメントする便宜がなくアセンブラ言語でexplicitに書きあらわさなければならないことである。ALTEXTはnameを正しく使っているかどうかというチェックはほとんど行なっていないのである。第3にはXPOPシステムは、算術式に対して組み込みのコンパイラを持っているけれどもこれはすべてが浮動小数点命令であるかあるいはすべてが固定小数点命令の場合(スイッチに依存する)にのみコンパイルを行なう。したがって型のチェックも行なわないし、mixed modeの式を用いることはできないのである。最後に高水準の構造を持つALGOLのような言語は容易にはインプリメントできない。というのはマクロの呼び出しは、簡単なやり方ではnestさせることができないからである。十分なマクロの特性が用意されているかどうか確信できないと言っても過言ではないのである。ある言語に対して新しいステートメントを付け加えることは必然的に他のマクロの特性をもたらしうる。丁度上記の性質4がDOステートメントの処理のためにインプリメントされると同時に。

Halpernは、彼がALGOLのような言語をインプリメントするために用いる他のコンパイラ作成の道具を置きかえるつもりはないと述べている。彼はそのプログラムが英語のようである処理言語に興味を抱いており、彼のシステムはこのためには適当なものであると信じているのである。

Halpernは最も楽観的でありメタ・アセンブラの支持を言い続けている。メタ・アセンブラにおける彼の仕事は、XPOPによって自然言語に“近い”ものをインプリメントすることができるという彼の声明によって、自然言語プログラムに関する彼の議論の余地のある主張に結びつけられている。

Halpernの論文[Hal 67b]は、XPOP型のシステムを念入りに擁護するものである。彼はマクロシステムの<オペレータ><オペランドの記号列>という記法は自然言語や数学的な言語とは反対であるが、プログラミング言語の標準的なシンタックスであると主張している。

Halpernはまた、プログラミング言語の研究を3つの部分に分けている。すなわち関数的(マクロ)、記法的(句読点命令の変更など)そしてモードに関するもの(アセンブル時の実行)である。

#### C.4 Extendible Compiler — Basic Concepts

(McIlroy [Mc Il 60] 以来) コンパイラ・システムにマクロの特徴を組み込むという試みが幾度もなされてきている。そのひとつのアプローチは、マクロ・シンタックスの形を保持しながら コンパイラに似た 組み込みの特性を数多くつけ加えるというものである。例えば SET システムは 入出力、記号の操作、テーブルのとりあつかい それにリスト処理といった特徴を備えた skelton compiler を含んでいるものである [Ben 64a]。これらの組み込みルーチンは、TWS を作ろうとする時には 翻訳時のオペレーション (action operation という) に結びつけられるのである。もっと上首尾のアプローチとしては、拡張への基礎として高レベル言語の structured syntax を用いるものがあげられる。

多くの既存のコンパイラ (PL/I を含む [IBM 66]) は 簡単な形のマクロ展開を組み込んでおり、その最初のものは JOVIAL であろう [Shaw 63]。その最も原始的な形は、パラメタの代入を行わず、純粋にテキストを置き換えるものである。例えば B5500 ALGOL では、次のようなステートメントでマクロを定義することができ

```
DEFINE LOOP1 = FORI ← 1 STEP 1 UNTIL N #
```

あとで 次のようにステートメントを書くと

```
LOOP 1 N DO A[1] ← 0
```

展開されて次のようになる。

```
FOR I ← 1 STEP 1 UNTIL N DO A[I] ← 0
```

次の段階ではパラメタを持つマクロの定義が許される。この便宜は、数あるものの中で特に AED-0 コンパイラ [Ross 66] に含まれている。AED-0 においては、次のようなステートメントでマクロを定義できる。

```
DEFINE MACRO LOOP(P1, P2) TO BE
```

```
FOR P1 ← 1 STEP 1 UNTIL P2 DO ENDMACRO.
```

この場合 より短いステートメント

```
LOOP(I, N) A[I] ← 0
```

で、上記と同じ結果を得ることができる。

これらふたつの簡単なマクロの形は、いかなる高レベル言語につけ加えても有用な形であり一層複雑なマクロの方法に匹敵するメカニズムの展開を考えることもできるのである。AED-0 では パラメタとして任意の記号列を認め 入れ子になった定義 (nested definitions) を許しているけれども、条件付きアセンブルのような性質は 高レベル言語においては広く用いられるとはいえないようである。この理由のひとつは、コンパイラが 通常 テキストの構造に非常に依

存しているためである。次の2つの節では、コンパイラをマクロのやり方で拡張しようとする時に生ずる複雑な問題について述べることにする。

## C.5 Definitional Extensions

(Cheatham [Che 66])

高レベル言語の定義的な拡張 (definitional extension) は、Computer Associates のグループがとり組んでいる TWS における最近の問題である。この仕事は、第Ⅲ章 B.2 に詳細にわたって述べた TWS の仕事に照らして考えるとよく理解できる。

ここで主題となっているこの論文は、性急に書かれたもののようで 準備のためにいくつかの内部仕様書を参照している。これは明らかにこの線に沿った最初の試みでありこれに続く論文で拡張され明らかにされるであろう。ここで述べているコンパイラへの拡張は記述的な (descriptive) metalanguage  $\mathbb{E}_D$  (第Ⅲ章 B.2) と一連のマクロの機能という二つの大きなカテゴリーに分かれるのである。

マクロの方法を用いる言語の拡張は text macro, syntactic macro, computational macro の3つのカテゴリーに分けることができる。text macro はすでによく理解されているものと考え、さらに前述のものに類似したものであると仮定する。cheatham は、syntactic macro をとり扱う時、コンパイラに対してマクロの概念を適用させるという問題にまじめに直面しはじめるのである。

ひとまとめに考えることのできる2種の syntactic macro があって、これら二つの基本的な特徴は free format とパラメタに対するタイプの指定 (type specification for parameters) である。

```
LET N BE INTEGER
```

```
MACRO SQUARE N MATRIX MEANS 'ARRAY[1:N, 1:N].'
```

伝統的な<オペレータ><オペランドのリスト>の形よりも free format の方が優れていることは明らかである。パラメタに対するタイプの指定は、条件付のアセンブルとよりよいエラー検出ができるようにするものである。この syntactic macro の呼び出しは、特別な区切り記号 (例えば%) によってはじまり 検出可能な終り (detectable termination) を持つことができる。特別な区切り記号の使用を避けるという第2のアプローチは、そのマクロの形をそのままトランスレータのシンタックス・テーブルに加えることである。このやり方でいくと上の例は

```
LET N BE INTEGER
```

```
SMACRO SQUARE N MATRIX AS <attribute> MEANS
```

'ARRAY(1:N, 1:N)'

のようになる。ここで < attribute > は、その強調している言語の定義におけるシンタクティックなタイプである。MACROS と SMACROS は共に、適当な実パラメタの descriptor を代入することによってインプリメントされる（第Ⅲ章 B.2 を参照）。これらの方法は いずれも TRAGEN のインプリメンテーションでは問題にならなかったが、いずれも もしまちがって用いられると最悪の結果をもたらすものである。

syntactic macro を論じる際に、cheatham は マクロの定義に "セマンティックス" を加えるという問題にふれている。これは、many-many macro やメタアセンブラで用いられるアセンブル時の action の類型である。このアプローチは不可能であるとする Cheatham の結論は、ほとんどすべてをこの問題にかけてきたメタ・アセンブラのアプローチと比較すべきであろう。彼の解答は、多くの primitive なオペレーション（例えばテーブルの拡張）を与え、拡張可能言語の影にある 完成された metalanguage (TRAGEN) の存在を指摘することである。

マクロの拡張の第3のタイプは、computational macro と呼ばれるものである。このやり方によれば、宣言されたマクロから生ずる中間コードにおいて代入がなされるのである。マクロの本体に対する中間コードは 前以って（仮パラメタを用いて）作り出されているので、このやり方は それに対して中間コードが文脈に独立にコンパイルされるはずの構成要素 (construct) に限定されるのである。この条件に合う場合には、computational macro は有用であり 効率のよい道具である。computational macro の簡単な例として最大  $4 \times 4$  の左三角行列に対する次のような機能をあげることとする。

TAKE I, J AS INTEGER

MAP M(I, J) = (11 - I) \* 1/2 + J - 6 ;

ここで TAKE と MAP は、この言語における宣言である。

このコードは配列のアクセスのためのものであるから、これをソース・テキストの中に挿入することはできない。この computational macro の形はほぼ適切なものである。Cheatham が指適しているように computational macro は、コンパイラ作成者によって 配列をアクセスするコードを作り出すために用いられてきた。この論文では次節で述べる Perlis と Galler の論文に関する考察でも再び現われる問題である。アクセスの機能 (accessing functions) についていくつかの例をあげている。ここで重要な点は、Cheatham は アクセス機能を記述する procedural な方法を与えているのに対して、Perlis と Galler は procedural でない記述から コードを generate しようとしている点である。

## C.6 ALGOL C

(Galler と Perlis [Gall 67])

これは、プログラミング言語の分野における二人の傑出した人物による 非常に長く 難解なそして重要な論文である。この論文には沢山の有意義な面があるけれども、ここでは 拡張可能なコンパイラについてのみとり扱うこととする。その他の事項については、新しい研究領域の第一歩として 第IV章Cでとり扱うこととする。

基本概念は、ここでも 高レベル言語にマクロのような機能をつけ加えることである。この目的のために彼等は、拡張するのに適当であると思われるALGOL C と呼ばれるALGOLの version を定義している [Naur 63b]。ALGOL C を拡張したものはどれもALGOL D と呼ばれ、そして これらの中の任意のプログラムは 機械的にALGOL C のプログラムと等しいように reduce できるのである。その拡張は、トランスレータのシンタックス・テーブルに付け加えるCheathamのSMACROSにやや似ている構成要素を通じて 成しとげられるものである。Perlis とGaller は非常に複雑な方法で マクロ処理を行ないたいと思っているので、彼等は2, 3のきまったカテゴリーにおいてしか 再定義を許していない。基となっている言語のALGOL C は、配列をとり扱ったり 直接に拡張可能性に関係していたりする多くの特徴を備えている。

拡張可能性に関する特徴の中に、location と value の間の変換に対する演算子がある。すなわち

(a) 整数型の結果を伴う単項演算子：

loc of x

ここで、xは<procedure identifier>, <variable>, <array identifier>のいずれかである。loc of xは、基本的にはxという値を含む語のアドレスである。

(b) ある<procedure>, <variable> もしくは<array>の“アドレス”を表わす、左のオペランドが<type>で 右のオペランドが整数の式であるような 二項演算子：

<type> vc of x

<type> pic of x

これらは、それぞれ“value contents of”と“procedure identifier contents of”を表わしている。したがって もしxが<type> realの変数であれば

real vc of (loc of x) = x

である。

(c) location と value の概念は、適用演算子(application operator)  $\oplus$  を用いて

<procedure>に拡張される。厳密なシンタックスの変更は 配列の約束と密接な関係があるが、<primary>と<function designator> の修正した定義はその意味を示しておかなければならないであろう。

$$\begin{aligned} \langle \text{primary} \rangle &::= \langle \text{unsigned number} \rangle \mid \langle \text{variable} \rangle \mid \langle \text{function designator} \rangle \mid (\langle \text{arithmetic expression} \rangle) \mid \text{loc of } \langle \text{procedure identifier} \rangle \mid \\ &\quad \langle \text{type} \rangle \text{ vc of } \langle \text{arithmetic expression} \rangle \\ \langle \text{function designator} \rangle &::= \langle \text{procedure identifier} \rangle \oplus \langle \text{actual parameter part} \rangle \mid (\text{pic of } \langle \text{arithmetic expression} \rangle) \oplus \langle \text{actual parameter part} \rangle \end{aligned}$$

このようにして address variable とほとんど同じ方法で procedure の名前をとり扱かうことができ、例えば procedure array などを作ることが可能である。ALGOL C を作るための ALGOL に対するこれらの追加は 余分なメカニズムのほんの一小部分を構成するにすぎず、その大部分は ALGOL D の多様性のある形の中に埋めこまれている。

すべての ALGOL D という言語は、ほとんど同じシンタックスを持っている。すべての ALGOL D に対する共通のシンタックスは <type>, <算術式>, <論理式>, <代入ステートメント>を、それらのシンタクティックなタイプに対する特別な形の定義を可能にする規則の集合で置きかえることを除いては ALGOL C に同じである。新しい定義の導入は、任意のブロックの始めの部分にある一連の宣言文として出てくる。この過程についての詳しい説明は非常に複雑なので我々は、例を使って概略を示すことにする。

基本となる趣旨は、新しいデータ・タイプとそれらに関連する演算子の定義を認めることである。これらの演算子に対する記号を見つけるという問題は、太字の英大文字を使うことで解決する。演算子順位文法を仮定する(第Ⅱ章B.1)ことによって、MAD [Ar 66]におけるように新しい演算子の順位を既知の順位をもつ演算子に関連づけて定義することができる。演算子に関する残っている問題には データのタイプがあるが、これについては数センテンスをさくことにする。

新しいデータタイプはALGOL C の表現 あるいはmeans ステートメントによって、これ以前に定義されたタイプを用いて定義される。これは 仮パラメタを含みうる。この仮パラメタは、もし それが存在すれば 後の処理で重要な役割を演ずるもので、例えばmatrix(u, v) means array [ 1 : u, 1 : v ] などのようなものである。

次に context ステートメントの集合の中では、演算子とタイプの情報を結びつけることができる。context ステートメントは、ある演算子に対して そのオペランドのデータ・タイプとそ



の結果とを記述する。それはまた（最終的には）適当な ALGOL C のテキストに還元できる  
 <記号列>を含んでいる。次の例は（擬似）LISPの定義を示すものであるが、これらの概念  
 を明らかにするのに役立つものと思われる。

基本的な LISP の predicate atom と eq は、論理的な procedure として定義されていると仮  
 定する。

```
Boolean procedure atom(x); list x ;
 atom := cdr x = 0 ;
Boolean procedure eq(x, y); list x, y ;
 eq := car x = car y $\wedge$ atom(x) $\wedge$ atom(y) ;
```

次の定義は、名前の構造としてのリストを作るのに用いられる。

- (1) list means integer array[1:2];
- (2) cons = \* ;
- (3) car > cons;
- (4) cdr = car;
- (5) of < cons;
- (6) list a cons list b  $\equiv$  list 'list(a, b)';
- (7) car list a  $\equiv$  list'a[1]';
- (8) cdr list a  $\equiv$  list'a[2]';
- (9) loc of list a  $\equiv$  integer
- (10) integer a := list b  $\equiv$  integer'a := loc of b' ;

ステートメント(1)は新しいデータ・タイプ list を 二つの要素を持つ整数型 の配列として定  
 義する。ステートメント(2)から(5)は、LISP の 4 つの演算子の間の相対的な順位を述べている。  
 ステートメント(6)~(9)は式を定義する。例えば (7)は list "a" の car がモデル化している配列  
 の第 1 要素であることを定義し、それがひとつの list としてとり扱われるべきものであること  
 を指定している。ステートメント(10)は、ある list に整数を代入する代入ステートメントを定義  
 している。この論文には、このほかにリスト構造に関する EVAL関数と各種の sequencing  
 operator の定義も出ている。

この例は、LISP に対しても Galler-Perlis のシステムに対しても正当なものとは言え  
 ない。というのは 彼等のシステムの全構想は、一番外ブロックに置かれた同じような定義の  
 集合によって定義される ALGOL C を備えているからである。それに続く各ブロックでは、  
 このトランスレータは 局所的な定義の集合を用いるため type table と context table とを作

り上げる。局所的なALGOL D のテキストの実際的な処理はきわめて複雑である。これは、文脈が回帰的であること、および ALGOL C のテキストが局所的に定義されたテキストに挿入できるということに起因する。この論文におけるその論議は、ALGOL C のコードの作成に加えて、計算を最適化したいという希望によって、さらに複雑になっているのである。

我々は Galler と Perlis の論文の意図を、首尾よくではないにしても慎重にゆがめてきた。彼等は、配列やとりわけ行列の計算におけるスペースを節約することにも 関心を抱いている。彼等が自分自身で論点を分類したのであれば、彼等にとってすべての面で好ましかったであろう。既に述べたとおり この論文は、拡張可能なコンパイラのほかにも重要な問題について論じている。我々の論題へのその貢献は、実用的であるというよりは理論的である。彼等は、非常に複雑なマクロ処理が可能であり、それが算法言語に実質のある変化を与えることを示している。しかし、翻訳時の効率のよくないこととプログラミング・エラーに対する感受性が、その適用を大幅に制限するのであるということが考えられる。さらに、どの位の頻度で 高レベル言語を変える要求があるであろうかという、一般的な疑問が存在する。これについては第IV章Cで再びとり上げることにする。

#### 〔参考文献〕

##### References for III. C

Benn 64a, 64b, Brook 60b, Brop 67, Che 64a, 66, Fer 66, Gal 67,  
Gar 65, GraM 65, Hal 67a, b, c, IBM 66, Lea 66, McIl 60, Mea 63,  
Moo 65, Star 65, Star 65, Wai 67.

## IV Related Topics and Conclusions

### A. Other Uses of Syntax-Directed Techniques

TWSの開発が始まってまもない頃から syntax-directedな方法が、広範で種々な問題に対して利用可能であったことが観察されていた。syntax-directed なアプローチは、プログラムに対する入力(form)が その内容(content)の重要な一部を含む場合はいつも 考慮されるのである。syntax-directed な方法についての個々の適用例は、詳しく書きたてられたいはしない傾向がある。ここで扱っている適用例は、主として個人の知識に基づいているものであり 代表的なものではあるが、確かに理解しにくいものである。

第Ⅲ章で述べたTWSは 他の用途に用いられる場合、一緒に使われるものに応じて いろいろに容易に変化する。syntax-directed symbol processorは、最も柔軟性に富み かなり広範な用途に用いることができるのである。このようなシステムのひとつであるAED [Ross 66, 67]は、最初から一般的な目的のプロセッサ (general purpose processor) として設計されたものであった。しかし それに対する姿勢とその術語が特殊であったために、AEDのプロジェクトは 他のTWSの仕事程 効果を上げることができなかったのである。

AEDのシンタックスのフェーズは precedence (順位)の方法に基づいており、第Ⅱ章Bにあるものに類似している。タイプのチェックや手書きのシンタックス・ルーチンを加えることができる機能を組み入れることによって、AED parserは その根底に流れる理論を犯すという犠牲をはらって さらに強力なものとなるのである。しかしながら 非常に興味深いものは、AEDステートメントの中間表現 (intermediate representation) なのである。これは plexの使用を基礎にしており、このplexとは その要素のそれぞれが多くlinkとデータとを持つことのできるようなデータ構造のことである。一方 "modelling plex" についての構成と処理とは、マクロ・ルーチンの集合によって成しとげられる。これらのマクロ・ルーチンは、code generationやコンピュータ・グラフィックスやプログラムされた tool commandのためのルーチンを含むものである。文献 [Ross 63, 67]は、いくつかの問題分野における詳しい使用例が出ているので AEDの紹介として好適なものである。

AEDシステムの本質的な特徴は、シンタックスにおける順位行列とセマンティックスにおけるplexのとり扱いかである。syntax-directed な世界に対する もう少し異なったアプローチは、第Ⅲ章Bで論じた一般的なコンパイラ・コンパイラのモデルから展開され得るのである。この計画においては データ構造の選択を含むセマンティックなメカニズム全体が、各適用分野によって異なることが可能である。たとえば、VITAL [Mond 67]の仕事では 二つの基本的

に異なるデータ構造の言語（共に VITAL で書かれている）が、それ自身 VITAL に基づいている syntax-directed な graphics package [Rob 66] の中で 比較されているのである。

その他の TWS の適用例の多くは、ある種類 あるいはその他の種類の記号操作の作業に現われている。最初の適用例のうちのいくつか [Scho 65] は、記号数学におけるものであった。ある TWS は、式の構造をモデル化する — おそらくは簡略化したり 分化したりする — のを助けるために用いられることもある。記号数学における TWS の使用（とりわけ COGENT や META の使用）は、現在 普及しつつあり、特にその目的のために作り上げられたシステム [Cla 66] を生じてきている。さらに syntax-directed なモデルが有用であることを見出している純粋数学の学者もいるのである。

TWS の適用例で format 変換の問題は、最も普及し余り意外とは言えないものである。これらは 大規模なデータ・ファイルと関連し、密接に関係のあるソース言語間の翻訳から生じてきた。もう一度述べておくが 第三章 A の syntax-direct symbol processor は最も度々用いられてきているのである。これらのシステムは 論理設計、geometric description の翻訳、シミュレーション、logging routine などのような多様な作業に利用することもできる。この他にも 特殊目的のデバイス（装置）に対する一連の命令を作り出すような TWS の適用例も沢山ある。例えば かなり複雑な TWS [Cas 66] は、人工衛星追跡システムのさまざまな構成要素に対して命令を翻訳するのに用いられた例もあるのである。

TWS は 幾多の分野に直接適用されているばかりでなく、ほかのいくつかの分野の仕事にも示唆を与えている。これらの分野のうち 特に盛んなのは、picture のシンタクティックな記述である。picture 処理に対する syntax-directed なアプローチは しだいに普及してきたように見えるが [ShaA 67, An 67 を参照]、その初期のある研究者 [Nar 66] は このアプローチを（[Nar 67] のある章の中で）否定しているように見える。新しいリスト処理言語に組み入れられた pattern matching の特性は TWS による示唆を一部受けており、artificial intelligence の関係分野でも 従位的計画として syntax-directed なプロジェクトを備えているのである。

コンピュータ言語学の分野は、その理論面からも実際面からも TWS の研究に密接に関連している。この分野における応用例は 意外に少ないのではあるけれども、大変重要なものである。コンピュータ言語と TWS のシンタクティックな理論は、いずれも Chomsky の初期の研究 [Chom 63] に基づいていて 共通な概念も多い。英語のシンタックスのインプリメンテーション（特に [Kun 62]）は top-down TWS と同時に開発されたが、自然言語の仕事は TWS の研究において開発された有効な進歩をなかなかその中に組み入れないようである。セマンテ

ィックスの応用では言語学に全く新しいアプローチをもたらしたDEACONプロジェクト〔Th 66〕が、TWSの方法の直接的応用であると言えよう（〔Nap 67, Col 67〕を参照）。そして言語学者たちがセマンティックを理論に検討を加えようとしているこれらの研究部門とnonproceduralな言語に対処しようとするTWSの研究者とが相互に影響を及ぼし合うことが期待される。

最後にこれは決してここで考えている用例が最も少ないというわけではないが、教育面に応用されたTWSにふれておく。以上に述べたTWSのうちのいくつかは、トランスレータ作成の課程の基礎として用いられている。これらは教養課程から学部ゼミにまで及び、大変注目されている。またこれらの課程は機械についての問題を除外して講義されているにもかかわらず学生が討論を通じてTWSに容易に近づく時、多大の成功を収めている。

〔参考文献〕

REFERENCES FOR IV.A

Ab 66, An 66, Brook 67a, Cas 66, Chom 65, Cla 66, Col 67, Gro 66,  
Hal 66, 67b, It 66, Kun 62, Mond 67, Nap 67, Nar 66, 67, Rob 66,  
Ross 63, 66, 67, Scho 65, ShaA 67, Th 66.

## B. Formal Studies of Semantics

コンピュータ・サイエンスは 今まで数学に負うところが非常に大きかったが、今やその負債を返す時が来ようとしている。プログラミング言語のシンタックス(第Ⅱ章C)とセマンティックスは共に、形式的なとりあつかいを示唆している。この節では TWSに關係の深い開発について簡単にふれ、プログラミング言語の形式的なセマンティックスに関する文献の概要を述べておきたい。

プログラミング言語のセマンティックスの 形式的研究にはいつも、シンタックスをセマンティックスからどのようにして分離するかという問題が起ってくる。プログラミング言語は、(問題を解く)論理から出た概念と(現在では軽視されている)自然言語とを結びつけるものである。そして プログラミング言語を取り扱う場合、シンタックスが 厳密には ここで問題となっている syntactic metalanguage で記述可能な言語の一面と取られることが多い。しかしこの習慣は、metalanguageでの定義が変わるとに シンタックスの定義が変わるという 好ましくない結果を得るのである。

論理になれているコンピュータ・サイエンスの研究者(例えば [Tix 67])は、そこで用いられる定義——すなわち、その形(form)の表現によって書きあらわされうる記号列の性質がシンタックティックであるという定義——を採ることを歓迎するであろう。このようにして記号列がある計算規則における定理であるか否かは、シンタックスにおけるとうてい結論の出そうにもない問題なのである。このアプローチは 自然言語に対して有効であることを証明しているのではなく、プログラミング言語における直接的な問題なのである。

例えば

$$X \leftarrow Y / 0.0$$

L 1 : go to L 1

というステートメントは、シンタックスの面から見て ALGOLとして正しい形であるだろうか。確かに データ・タイプを処理することのできるアルゴリズムは、これらのエラーを見つけるであろうが、問題は どこまでそれが及ぶかという事である。論理に携わる人々の要求を満たした上で、形式の整った決定可能な性質をもつプログラムを残すようなシンタックスの概念を作り出すことができるかどうかは明らかではないのである。

この問題をさらにむずかしくしているのは、総ての主要な言語が形式的シンタックス(formal syntax)だけでは解析できないステートメントを含んでいるという事である。このことについての良い例は、PL/I ([IBM 66])では、ラベルをつけたENDを使用可能であるが、ALGOLでさえも今だにこのような構成要素を用いることのできないことである。したがって 実

用においては、syntax-directed compiler は シンタクティックな面に “セマンティック” 性質を組み込まねばならないのである。このシンタックスとセマンティックスを分けるという問題についてのひとつの功妙なアプローチは、McCarthy の abstract syntax [Mc Car 62a] である。彼は主に セマンティックスについての研究を行っており、(分析的な) シンタックスは ソースの記号列の形から適切な情報を引き出すのに必要な predicate や function の集合にすぎないと考えている。このシンタックスを定義するという問題を「解決」はしないけれども、シンタックスとセマンティックスについて考えることを可能にするものである。

例のごとくセマンティックスの形式的研究は、プログラミング言語のシンタックスの研究に遅れをとっている。そしてこの問題に関する研究のうちで 最も優れた一般的研究は [Ste 66] であり、そしてこれは書物よりもむしろ討議によって形式的セマンティックスの概念規定を与えているものである。現在までに考案されたさまざまな形式化 (formalization) は、すべて procedural なものである。それらは 抽象的な機械のこともあるし、 $\lambda$ -calculus [Chu 51] のような命令的形式化 (imperative formalisms) のこともある。これは期待しても当然のことなのであるが、現在 存在している数学的モデルの選択をかなり制限することである。

形式化 (formalization) が procedural なことから、プログラミング言語の記述においては “meaning” という言葉よりも “effect” という言葉の方が好まれるようである。これはセマンティックスの概念を effect として支持するという立場をとるのではなく、ただ 物事を見るのに都合のよい方法として それを採用するのである。この見方によって “environment (環境)” に依存して異った effect を持つようなプログラムを期待でき、このことは我々の討論に有益であることがわかるのである。また このことは semantic metalanguage の選択が、形式的記述の意図的な使用によって影響をうけるという疑いを抱かせるようになるのである。

現在の形式的セマンティックスにおける研究は、プログラムに関する証明に関係のあるものと コンピュータによるプログラムの処理を解明するのに興味をもつものとの二つに大別できるとであろう。またこれは必ずしも必要ではないが、第Ⅲ章Bで述べた semantic metalanguage を 後者の中に入れて含めることができる。しかしこの中に含めてよいと考えられるものは、多少 抽象化された翻訳モデル ([Wir 66c] など) もあるのである。このようなモデルにおいては、ある言語はそのトランスレータがインタプリタであるか コンパイラであるかによって大変に異なった effect を持ちうるのである。これはプログラマにとっては理にかなったことのように思えるのであるが、meaning がプログラムではなく アルゴリズムに存するという見方を好む数学タイプにとっては ただ混乱を招くだけなのである。上記の研究に関連した開発を掲げると、すべてのプログラミング言語をひとつだけそれよりも高いレベルの言語 [Ste 66] あるいは機械的な言語 [Brat 61, Ste 61]

に対する reduction によって定義しようとする計画がある。

これまで述べてきた形式化に対するアプローチは TWS に関係が深いものであるが、これは余りにも複雑すぎるので これを利用して証明を行なうには適していないのである。証明が形式化の唯一の目的であると考えた人々（そしてそういう人々がこの報告書を読んでいるのであろう）にしてみれば、これまで述べてきたことなど まわりくどい言いまわしにすぎないと思うであろう。形式化に基づいている数学的試みの多くは 取り扱いやすいという点に重きを置いており、そのほとんどは 既存の数学に基づくものである。論理 (logic) における命令的なシステム (imperative system) は 2, 3 あるだけであるが、そのそれぞれはコンピュータ・サイエンスのある面を形式化するのに用いられているのである。形式的セマンティックスにおける研究の大部分は、church [chu 51] の  $\lambda$ -calculus や Curry の combinator calculus に基づいている。

これらの理論はいずれも 主として変数の役割に関係しており、また これらがプログラミング言語において成功しているのも 主にこの分野においてである。 $\lambda$ -記法は LISP において重要な役割を果たしており、いろいろな LISP に関する書物では プログラミングの概念として論じられている。これはまた セマンティックスを形式化しようとする試みに対する最もポピュラーな方法なのである。Landin と Strachey の業績は、彼らが CPL と呼ばれる ALGOL 60 の拡張型 [Burs 65, Cou 65] の開発に彼らの研究を組み合わせたことによって 特別に興味深いものとなったのである。 $\lambda$ -記法のセマンティックスへの適用は、Landin が最も苦心して追求したものであった。一連の研究論文で 彼は、プログラミング言語 (ALGOL) と増大した  $\lambda$ -記法との間に imperative applicative expressions (IAE) と呼ばれる関係を考えている。ALGOL における変数の宣言と結合が非常に明らかにモデル化されており、その形式化が ALGOL のいくつかの弱点を指摘するのに役立っている。(純粹の LISP と同じように) IAE システムは、純粹に functional なものであり そしてその環境に副作用 (side effect) を伴う O- $\lambda$ -adic function として ステートメントを表現しなければならないのである。実際 Landin の ALGOL についての記述の多くは、LISP における "プログラムの性質" [Mc Car 62b] の一般化として見るのできるものである。しかし この段階まででは、これらの研究は必要とされている表記上の明快さを達成しておらず、またオリジナル・プランに関連した  $\lambda$ -calculus の扱いやすさを保持するのにも到っていないのである。この研究による最もすばらしい利益は、極端な程までに文明化された (civilized) 言語 CPL [Cou 66] を持っていることである。そして MIT には、このアプローチをさらにできる限り推し進めて行こうとしているグループが 現在活動을続けているのである。



McCarthy は、彼がコンピュータに  $\lambda$ -calculus を導入したのではあるが、形式的セマンティクスへの彼のアプローチは多少異なっている。彼の使う "theory of computation (コンピュテーションの原理)" という言葉は、彼が算法定語よりも アルゴリズムに関心を持っていることを示している。彼のアプローチは state ベクトル、およびそれに対するオペレーション、抽象的なシンタックス、そして条件付きの expression を利用するものである。典型的な state 関数は、 $c(x, d)$  — state ベクトル  $c$  におけるシンボリックなポジション  $x$  の内容 — および  $a(x, z, d)$  — state ベクトル  $d$  において  $x$  に  $z$  を代入する時の結果の state — などである。次に 彼は 条件付きの expression を、機械語のコードのようなオペレーション および抽象的なシンタックスによって記述されるような高いレベルの構成要素の定義であるにとどめることができるようにしたのである。この結果として得られる形式化は かなり扱いやすいので、McCarthy とその研究員は 多くの証明 [Mc Car 67] を完成することができたのである。最近の成果としては、Painter [Pai 67] が "コンパイラ" の正統性を証明できたことが掲げられる。

Floyd は、最近 直観的によく満足できるアプローチを開発した [Flo 67]。彼は通常の (固定した) プログラミング言語で書かれたプログラムのフローチャートを考えた。その基本的な考え方は、フローチャートの各接続部分 (connection) に state ベクトルを適用して proposition を付けることである。この proposition は connection が実行中 (翻訳の間を考えると) に保たれている限り、成り立つものである。これらの proposition および これに関連したいいくつかのメカニズムを用いて、Floyd は「もし初期状態 (initial state) が  $R_1$  を満足する時 これが到達するものであれば、最終状態 (final state) は  $R_2$  を満足する」という形式の文章を証明する方法を開発したのである。さらに この結末の証明は、例えば そのプログラムが実行されるにつれて減少していく正の整数による何らかの関数を示すことによってとりあつかうことができるのである。また つい最近 限定された言語に対して proposition を生成することと、その正統性を証明することの両方を自動化しようとする研究も進められている。

以上 述べてきたのは 極端な例なのであるが、このほかにもこれらの中間をいくようなセマンティクスの形式化へのアプローチがいくつか行なわれている。そのひとつの例 [Don 67] をあげると、これは (タイプのマッチング等を含む) シンタックスとプログラムの compilation の双方を記述する post canonical なシステム [Pos 43] のひとつの version を用いるものである。ここでなされた定義は 一見合理的に見えるけれども、それらが compilation あるいは証明の中で役に立っているかどうかは疑問である。また 現在の数学を利用する例としては、van Wijngarten と de Bakker [Bak 65, 67, Wij 66] の研究があげられる。彼らは 単純で

はあるが 有効なルールを読み 翻訳することのできる universal な機械を用いて、セマンティックスをモデルの複雑さを解消しようとしたのである。そしてこれらのルールは、ソース言語についての Markov アルゴリズム [ Mark 59 ] の記述にあたるものを定義するのに累積的に用いられるのである。ここで問題となるのは その形式論が余りにも初歩的なものであるため、ALGOL の記述が逐語的に書かれたもの [ Bak 67 ] となるだけで 証明や洞察が効果をあげるには到らないという事である。

また プログラミング言語のセマンティックスを遂行する抽象的な機械を定義しようとする試みも 数多くなされている。このうちで最も意欲的な研究は [ Elg 64 ] の RASP であるが、この研究は現在では続けられていないようである。[ Nar 67 ] による非常に興味深い論文があるが、これは 以上に論じてきたいろいろな性質を統合した形式論を含むものである。

Narasimhan は フロー・チャート、state 変数についての関数、宣言文、演算子の選択やアドレッシングを含む関係の深い形式論において言語や機械を定義しているのである。そして この研究は有望なものに思われるのであるが 具体的な成果はまだ現われておらず、また彼の根本的前提にかなり問題があるのである。というのは トレンスレータを可能な限り単純化しようとするあまり Narasimhan は、シンタックスおよび recursion が価値を持たないような結論に到ってしまうのである。彼はまた (文献には載っていないが) TWS の研究は総て失敗しており、そこから得られる興味も衰退の一途をたどっているとも言っているのである。

形式的セマンティックスの研究について これまで述べてきた事は余りにも その内容が浅いため、おそらく誤解を招くのではないと思われる。これらの研究に携わっている人々は 相互に連絡しあっている事もあるであろうし、また その成果も我々の予想よりも豊かなのかもしれない。そこで この論文に満足できない方々は、この終りに TWS 関係のすべての傾向にわたる参考文献を掲げておいたので これを参照されたい。形式的セマンティックス 特に証明に傾いているもの (proof-oriented) は、それぞれ孤立したわずかの研究に影響を与えているにすぎない。たとえば Krohn と Rhodes のオートマトンの理論などは、全くといって良い程 どんな研究にも影響を与えていないのである。思うに既存の命令的な論理 (imperative logic) では、この問題を解明するには到らないのではなからうか。そして プログラミング言語は、数理言語や自然言語がつき当たったのと同じ問題に直面せざるを得ない時がやってくるのではないと思われるのである。Minsky と Papert [ Min 67 ] が、同様の観測を述べているので最後にこれを記しておく。

優れた理論というものは、良く理解された実際の問題の文脈から離れてしまつては 展開できることはまれである。そして例えば 回帰的な関数の数学的展開、オートマトン、形式言語学など

のような「Computation の抽象理論」を作り出す事だけを目的とした研究が、数学という面から見れば いかに優れたものであろうとも 実行という段になると全く惨めな結果しかもたらさなかったとしても 別に驚くにはあたらないのである。

〔参考文献〕

REFERENCES FOR IV.B

Bak 65, 67, Braf 63, Burg 64, Burs 65, Cal 62, Chu 51, Cul 67, Cur 58,  
Don 67, Elg 67, Flo 67, Ir 61, 63b, Knu 67, Landi 63, 65, 66, Luc  
65, McCar 62a, 67, Min 67, Nar 67, Org 67, Pai 66, Rig 62, Ste 64,  
Tars 56, Tix 67, Wir 66c, Zar 67, Zew 66.

## C. Summary and Research Problems

この論文で取り上げたTWSは、多くの傑出したコンピュータ科学者たちの研究の長い歴史に支えられてきた最近の研究成果を代表するものである。また 第Ⅲ章で扱ったカテゴリーは、それぞれ独特の長所と短所を持ち、それぞれに都合のよい問題の範囲が示されている。そこで種々のカテゴリーの間の関係をまとめた後、将来の研究に対する多くの実り多き分野について述べてみたい。

第Ⅱ章Bで述べた recognizer の自動な構成要素 (automatic constructor) は syntax-directed なアプローチで解明できるような問題ならば、どのようなものにもでも使用する事のできる道具である。効率のよい recognizer を自動的に作り出すことによって、このようなシステムは syntax-directed な方法の有効範囲を広げることができる。ここで一番大きな問題は 総合的 (synthetic) なシンタックスに セマンティックスな定義をうめこむ手軽な方法を見つけることである。この問題が解決できれば、それは 第Ⅲ章Aの syntax-directed symbol プロセッサの能力も飛躍的に増大することになるのである。これらのTWSは総て、セマンティックスを導入するのにはかなり便利な方法を備えているのであるが、これらは比較的効率のよくない recognizer を用いているのである。このため、さらに効率の良い recognizer が開発されれば、現在もかなり進んでいるこれらのシステムの応用も拡充される事は明らかである。

第Ⅲ章C.1~C.3で述べたメタアセンブラは、現在のところ コンパイラの作成よりは アセンブラの作成に適しているようである。しかし、これらは広範囲に効果をもたらすマクロ言語にいくつかの有意義なものをつけ加えるものである。メタアセンブラの機能を、翻訳時の action に対して拡張したり、シンタックスなフェーズをつけ加えたりすることによって、それらを第Ⅲ章Aの syntax-directed プロセッサに匹敵するものとする事が可能であろう。

第Ⅲ章Bのコンパイラ・コンパイラは、専門的なTWSの発展したものとしては最高のものである。専門的であるということによって、それらはコンパイラの作成には、はるかに便利になったのであるが、その代償も大きなものである。すなわち、コンパイラ・コンパイラは、インプリメントすることがむずかしく、コンパイルとは異なる仕事には向かないということがしばしばあるのである。セマンティックス言語が、より複雑なプログラミングの構成要素を含む方向に進んで行くにつれて、このようなシステムはさらに専門化してくる事が予想される。そして、これは言うまでもなく、過度に専門化してくるという危険性をもっており、TWSの研究者の中には、COGENT (第Ⅲ章B.3) のような一般的な syntax-directed プロセッサの方が、将来においては役に立つのではないかと考える向きもあるようである。

拡張可能なコンパイラの研究 (第Ⅲ章C.4~C.6) は、ごく最近のものであり、いかなる高レ

ベル言語にも いくつかのマクロの機能を含むようにするべきであるということは明快なようではあるが、これを正確に評価する事は困難である。これよりもさらに風変わりなシステムは 利用範囲が狭くなると思われる。理論的に言えば、ある言語をマクロの方式で拡張してから その後その拡張をそのコンパイラに効果的に組み入れることも可能である。COシステム(第三章B.3)は この両方の便宜を備えており、これがその問題を解決するとは言えないが その解決法を試すのには好適な便宜である。

これまで述べてきたTWSは そのどれをとっても それだけで総ての問題が解決されるというわけには行かない。そして我々は そのようなものを期待するのは無理であるということを説明しようとしてきたのである。また いろいろな種類の universal programming システムについてのさまざまな試みのもたらした結果からも これが正しいという事がわかる。しかし全体として見れば、TWSに関する研究が コンパイラ作成にまつわる多くの重要な問題を解決しているという事が言えそうである。現在では 可能性のあるTWSのさまざまな要求を満たすのに十分な方法(technique)があり、未解決の問題は特殊な項目にすぎないのである。

TWSのシンタクティックな面は かなり注目を集めており 未解決の疑問は少ない。ここでセマンティックスに深い関係を持ち、また互いに関連する三つの問題が 浮かび上がってくる。この問題のひとつは、「シンタックス」が人間の直観に非常に近いものとなることができるように extraなシンタックスの性質を当てはめるのに満足 of いくような方法をどのようにして見つけるかということである。[Gil 66, Don 67] また これに関連して、そこでつかわれている recognizer-constructing プログラムに関して詳細な知識を持ち合せていないと 総合的な文法の規則にセマンティックスを適用する適当な方法がないということも述べられている。最後に自動的な recognizer-constructing プログラムが 優雅に退化していつているという問題もある。というのは 可能な限り 効果のある方法を使用するシステムが好まれるのであり、その時使っている方法がその任に耐えられなくなってしまってから これを放棄するよりは、速度の遅い より一般的な計画へと自動的に動いて行くからである。これに加えてシンタックス・エラーの自動復旧という問題は さらに注目を集めるであろう[EvA 63, Ir 65]。

しかし、TWSの postsyntactic な面については 余り研究が行なわれてきていない。この「セマンティックス」という問題に対するアプローチとしては、根本的に異なる次の三つがあげられる。第一の方法は、汎用のリスト処理 もしくはシンボル操作の能力(第三章Aを参照)を与えることである。第二の方法は、沢山のデータ構造と組み込みのルーチン ことに コンパイラ作成用に設計されたルーチンを供給することである(第三章Bを参照)。第三の方法は、これらの機能に code bracket および機械独立な出力の指定方法を結びつける事である(第三章B・

1を参照)。マクロやサブルーチンを用いることによって、最初のふたつの方法は 普通のユーザーにとっては 自動化されたシステムのように見えるであろう。この点から見れば、セマンティックスにおける一番重要な問題は コンパイラ作成の大事な局面を扱う汎用のルーチンをいかにして見出すかという事になるのである。我々は、TWSのアプローチは可能であることが証明されており 一般的な問題は 現在では それを展開する段階で考えるべきであると感じているのである。もちろん 現在の範囲では解明できないプログラミング言語(例えば、シミュレーション[Te 66])もいくつかあるが、しかし それらの各々は先ず研究しなければならないような基本的な概念はわずかしき含んでいないのである。つまり TWSの研究の目指すところは プログラミング言語の未解決な問題を理解する(そして最後には自動化する)事に向かっているわけである。

このように TWSの研究について系統的に述べてきた結果、勿論 我々は 各人に対する年間計画は提示してきたことになるであろう。この見解が正しいという事は metaproblemを考慮することによってもたらされるプログラミング・システムに多大の貢献を与えていることによっても明らかである。そして その章の残りの部分では TWSのアプローチに附随する多くの興味深い問題について論じ、これらの問題に関する文献の概要を述べておきたい。この節の最後に記した参考文献には ごく最新のもの、包括的なもの、及びこの論文ですでに参照したものが含まれている。

機械語の形式的記述についての問題は かなり以前から提示されていながら、いまだに解決されていない。この問題が解決されれば、それは target machine を表わすTWSへの第三の入力として使う事もできるのである。ところが この問題は 理論的にも直接的にも研究されていながら、TWSで使用可能であるような成果は 全くといってよい程上がっていない。parallel processorを利用できるということは 新しい研究分野に新しいレベルの複雑さを加えるものであり、また これをよりよいものにするものである。parallel processor に対するソフトウェアの仕事の多くは 特定の機械に関連づけられているものなので、この論文の範囲外のものであるように思われる。parallelism の研究の基礎となるような いくつかの抽象的[Kar 66]あるいは具体的[Shed 67, Sto 67]理論もある。そして高レベル言語における parallelism [Dij 65]も注目を集めるようになってきている。

また コードの選択と増大の理論(optimizationの問題)に関係した問題もかなり古くからである。この理論は弱いばかりでなく コンパイラの作成者によって一般に使用されるコードの増大(enhancement)の型は6通り程しかないのである。プログラムの能率における飛躍的な進歩の多くは 通常問題に対するアプローチ全体を組み立て直すことによるものである。このことは

広義の optimization と呼ばれるが、我々は これを nonprocedural なプログラミングの一局面として論じることにする。“nonprocedural”は“セマンティックス”と同様に一般に認められた定義がまだ確立されていない。そこで より高い程度の problem ステートメントに応じて procedural なステップの選択と組み立て直しを行なうようなプログラミング・システムを nonprocedural と呼ぶことにする。

nonprocedural なプログラミング言語は 宣言的言語 (declarative language)、問題向き言語、質問システムなど多くの項目のもとで論じられてきた。しかし この研究の多くは理論的に余り面白いものとは言えない。というのは、ある人が大きなルーチンを作り ユーザーは それにパラメータを与えるだけであるからである。( [ You 65] を参照)。また ある限られた問題の分野に対するかなり優れた nonprocedural システムが コンピュータ・グラフィックス、関係的言語 (relational languages) [ Rov 67], 配列の処理 [ Gal 67], 数値解析 [ Ri 66] において開発されてきている。言うまでもなく アナログ・コンピュータは常にこの方法でプログラムされ、ハイブリッド・コンピューティングにおいて使われている言語を拡張することによって将来発展の可能性の強いシステム [ Schl 67] が開発中である。また Cheatham は 第Ⅲ章 C・5 で述べた拡張可能なコンパイラに一般的な種類の nonprocedural な性質をつけ加えるという構想を抱いているのである。一方、自然言語の処理を通じて開発されたより複雑なシンタックスの形や変換を用いるという 別の方面からのアプローチも可能である。TWS の研究は 自然言語システムと多くの点において同じ興味を持っていると思える程なのである。これには質問言語と呼ばれるもの [ Com 66] や、言うまでもなく COBOL [ Samm 61] があるのであるが これらは ここで取り上げた問題とは表面的なつながりを持っているにすぎない。しかし 会話的あるいは nonprocedural なプログラミング言語についての最近の興味ある事柄は、syntax-directed な自然言語システム (第Ⅳ章 A を参照) とあいまって、いろいろな考え方の間の情報交換をさらに有意義なものとしていくであろう。

TWS とエクゼクティブ・システム (executive system) との間の関係には まだ未解決の問題がいくつか残されている。TWS の主な利点のひとつは それぞれのコンパイラがエクゼクティブに対してインターフェイスをとる労力 (しばしば全体の半分以上にも及ぶものである) を削減する事である。しかるに、このパラグラフ以前には「エクゼクティブ」という言葉が出て来なかったという事が この分野における過去の研究の実態を示しているものであると言えよう。コンパイラについての“環境の (environmental)”問題 [ Leo 66] に興味を持っている小グループは前々からあったのであるが、それらはタイム・シェアリングという革命が起きるまではほとんど影響を及ぼすには到らなかったのである。つまり、大規模なマルチアクセス・タイム・

シェアリング・システムが使えるようになって（あるいはその要求が高まってきてから）新たなTWSに関する研究問題がもち上ってきたのである。大規模なタイム・シェアリング・システムのエクゼクティブの主な仕事は リソースの割り当てである。割り当てられるリソースには いろいろな記憶装置や処理ユニットばかりでなく、コンパイラのようなプログラムも含まれるのである。そこで研究問題は いかにしてシステムの効果をよりよくするように、トランスレータとエクゼクティブとが情報を交換することができるかという方法を案出することなのである。現存のシステムで早急に必要とされているのは、主記憶に関係したものであり、特定の言語の swapping time を減少させようとする いくつかの計画〔Bob 67, Coh 67, Roy 67〕も試みられている。これに関連する問題は TWS〔Feld 67〕とその結果のオブジェクト・コードにおいて純粋な procedure を optimal に（すなわち maximal でなく）使用するという事である。コンパイラとエクゼクティブの優れたインターフェイスは達成するのが大変にむずかしい反面、理論的には余り興味深くない解決方法が実際面では大いに役立つのである。

次にエクゼクティブ・システムに関連してふたつの問題について簡単にふれておくことにする。制御言語（control language）は、これにシンタックス処理をつけ加える事によって、また理論的に言えばTWSですでに用いられているのと同じのコードを使用する事によって改善することができる。さらに意欲的なプロジェクトは syntax-directed な方法をエクゼクティブ・プログラム自身を作るのに適用しようとしている。もうひとつのこれに関連する問題は debugging aids の問題である。on-line debugging システム〔Evl 66〕については 今までに多大な研究がなされてきたのであるが、その多くはMACプロジェクトを除いては アセンブラ言語のレベルにおけるものであった。すなわち、特定のバッチ用のコンパイラに 記号によるダンプ機能（symbolic dump facilities）の優れたものを備えたものもあるが プリントやTWSに到る道はまだ開拓されていないのである。また 自動的にエラーを検出すること、および syntax-directed プロセッサにおける復旧機能に関する研究〔Ir 65〕もわずかではあるがなされているのである。ここでもう一度述べておくが 優れているとはとても言えないようなシステムであってもユーザーにしてみれば大きな価値のあるものなのである。

最後にデータ構造に関する研究分野について論じることにする。この分野は 行列のとり扱いからファイルの処理にいたるまでのすべてを含んでおり、かつコンピュータ・サイエンスの種々の分野に深いつながりを持っているようである。ある意味では このデータ構造こそがコンピュータ・サイエンスで 現在問題となっている最も重要なものであり、このうちの未解決の問題を検討することが かなり意味の深い事であると言えよう。ここではTWSに関連したそのいくつかの面について述べ、ここで扱った他の研究問題において どのようなデータ構造に関係する考



察がなされてきたかを示すことにする。

いかなる TWS においても、翻訳時および実行時の双方で組み込まれているデータ構造の選択はひとつの主要な問題となっているのである。第 II 章では翻訳時の構造について述べ、実行時に対しては 組み込みの structure operator を与えるような事は 基本的には何もなされていないことを述べた。つまり 複雑なデータ構造の言語の多くは TWS (たとえば [Ab 66, It 66, Rov 67]) を使って書かれるのに対して、その structure operator はすべて手書きによっているのである。しかし 最近容易に TWS に組み入れることのできる単一の一般的なデータ構造 (universal data structure) を開発しようとする試み (たとえば [Ross 66, IBM 66, Wir 66b]) もなされている。ここで問題なのは データ構造が現在適用されている分野においては、既存の提案はすべて大変に効率のよくないものになってしまうということである。ある与えられたアルゴリズムに合った構造を選ぶとすれば 必らず nonprocedural なプログラミングが問題となってくる。同様に、データ構造を改善することによって グローバルな optimization や自然言語の処理を大きく進歩させることができるのである。実際、ここで扱ったすべての研究問題、およびその他の多くのものの中には密接な関係が存在しているのである。したがって TWS の問題は その本来の性質から言っても、プログラミングの研究の最新の分野と常に関連をもっているのである。

最近の TWS の研究を簡単に検討してみようとしたところが このような長い論文になってしまった。我々はここで、ほとんど独自で研究を続けているかなり多くの優れた人々が どのようにしてシステム・プログラミングの多くの面について合理的に理解し得たかということを示そうと努めてきたのであった。そして 情報交換が盛んに行なわれ、科学的な規格化がさらに進めば意義深い進展や 研究において開発されたアイディアの迅速な応用も可能となってくる。このような目的があるからこそ、我々はこの論文を書いたのであるし また読んで下さったものであると確信する次第である。

〔参考文献〕

REFERENCES FOR IV.C

Theory of machine instructions: Brat 61, Bur 64, Car 62, Don 67, Elg 64, Gil 67, Maur 65, Nar 67, Ste 61.  
Parallelism: Dij 65, Kar 66, Kuc 67, Mark 67, Shed 67, Sto 67.  
Code selection and enhancement, general references: Ar 61, Grie 65, Hill 62, Hor 66, Lucc 67.  
Nonprocedural languages: Che 66, Gall 67, Ri 66, Rov 67, Schl 67, Sib 61, Wil 64b, You 65.  
Natural language processing: Bar 64, Chom 65, Col 67, Cra 66, Hal 66, Int 63, Kun 62, Nap 67, Nar 67, Th 66.  
Executive interface: Feld 67, Le 66, Orch 66, Nob 63.  
Paging: Bob 67, Coh 67, Den 65, Rov 67.  
Debugging: EvA 63, EvT 66, Ir 65.  
Data structures: Ab 66, Brook 67c, Gall 67, IBM 66, It 66, Pra 65, Pra 66, Rov 67, Sta 67, Wir 66b.

謝 辞

さいごに John Reynolds氏ほか、示唆や御批判を寄せて下さった多くの  
方々に感謝の意を表したい。

## BIBLIOGRAPHY.

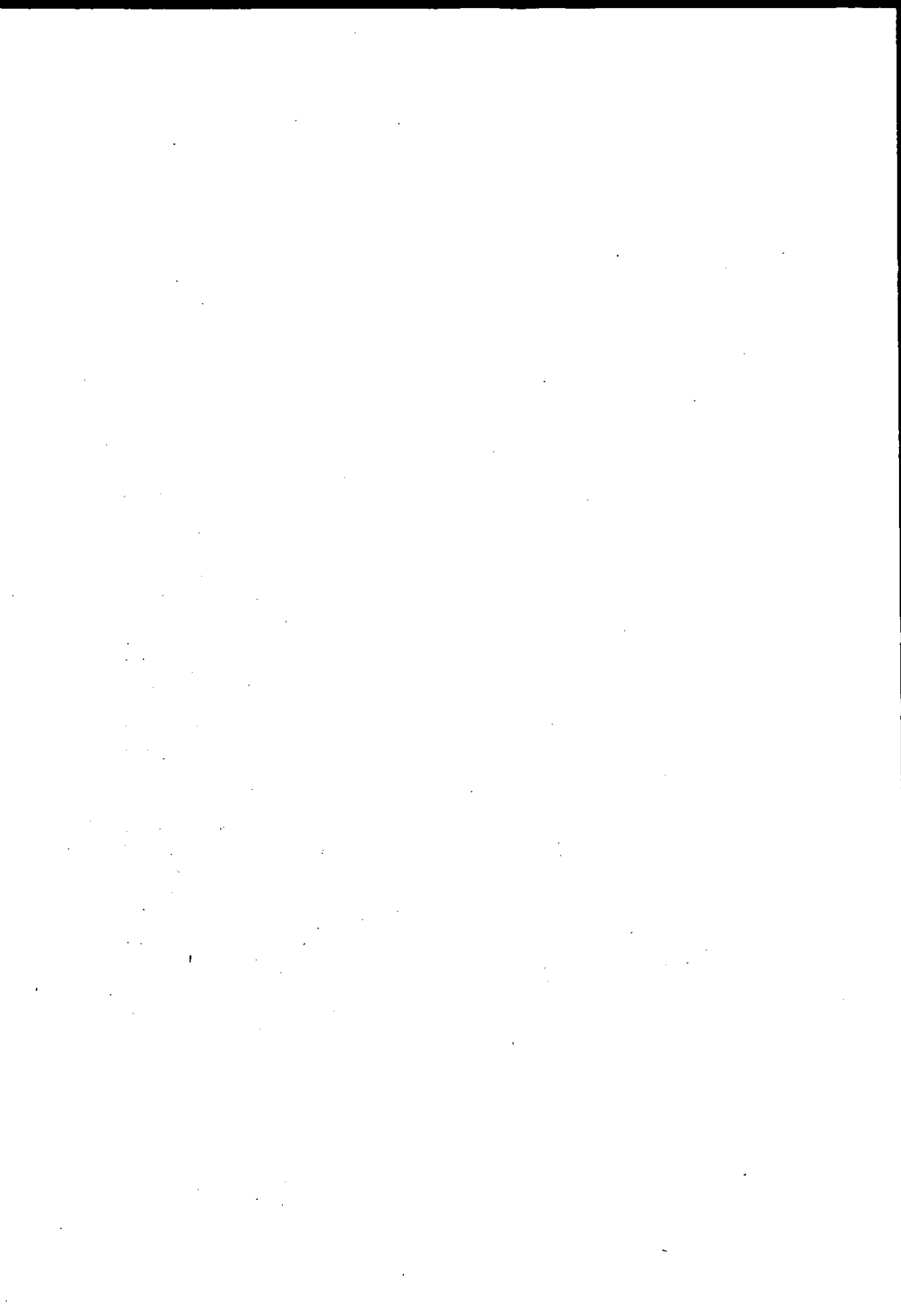
- Ab 66 ABRAHAM, P. W. The LISP 2 programming language and system. *Proc. AFIPS 1966 FJCC*, Vol. 29, pp. 661-676.
- Ac 66 ANDERSON, R. H. A two dimensional syntax for mathematical notation. Unpublished report. Harvard U., Cambridge, Mass., 1966.
- Ar 61 ANDER, B. W., GALLER, B. A., AND GRAHAM, R. M. An algorithm for equivalence declarations. *Comm. ACM* 4 (July 1961), 310-314.
- Ar 66 ———, AND ———. *Michigan Algorithmic Decoder*. U. of Michigan Press, Ann Arbor, Mich., 1966.
- Bac 57 BACKUS, J. W. ET AL. The FORTRAN automatic coding system. *Proc. Western Joint Comput. Conf.* 1957, pp. 188-198.
- Bac 59 BACKUS, J. W. The syntax and semantics of the proposed international algebraic language of the Zurich ACM-GAMM Conference. *Proc. Internat. Conf. on Information Processing*, UNESCO, 1959, pp. 125-132.
- Bak 65 DE BAKKER, J. Formal definition of algorithmic languages. M174, Mathematisch Centrum, Amsterdam, May 1965.
- Bak 67 ———. Formal definition of programming languages. Mathematisch Centrum Tract No. 16, Amsterdam, 1967.
- Bar 64 BAR HILLEL, Y. *Language and Information*. Addison-Wesley Publishing Company, Inc., Reading, Mass., 1961.
- Barn 62 BARNETT, M. P., AND FUTHELLE, R. P. Syntactic analysis by digital computer. *Comm. ACM* 5 (Oct. 1962), 515-526.
- Bart 63 BARTON, R. S. A critical review of the state of the programming art. *Proc. AFIPS 1963 SJCC*, Vol. 23, pp. 169-177.
- Ben 61a BENNETT, R. K., AND KVILEKVAL, A. SET, self extending translator. Data Processing, Inc., March 1964.
- Ben 64b ———, AND NEUMANN, D. H. Extension of existing compilers by sophisticated use of macros. *Comm. ACM* 7 (Sept. 1964), 541 (actually about assemblers).
- Ber 62 BERMAN, R., SHARP, J., AND STURGES, L. Syntactical charts of COBOL 61. *Comm. ACM* 5 (May 1962), 260.
- Bob 67 BODROW, D. G., AND MURPHY, D. Structure of a LISP system using two level storage. *Comm. ACM* 10 (March 1967), 155-160.
- Brat 63 BRAFFORT, P., AND HIRSCHBERG, D. (Eds.) *Computer Programming and Formal Systems*. North-Holland Publishing Co., Amsterdam, 1963.
- Brat 61 BRATMAN, H. An alternate form of the UNCOL diagram. *Comm. ACM* 4 (March 1961), 142.
- Brook 60a BROOKER, R. A., AND MORRIS, D. An assembly program for a phrase structure language. *Comput. J.* 3 (1960), 168-174.
- Brook 60b ———, AND ———. Some proposals for the realization of a certain assembly program. *Comput. J.* 3 (1960), 220-231.
- Brook 61 ———, AND ———. A description of Mercury Autocode in terms of a phrase structure language. *Annual Review in Automatic Programming*, Vol. 2, 1961, pp. 29-66.
- Brook 62a ———, AND ———. A general translation program for phrase structure languages. *J. ACM* 9 (Jan. 1962), 1-10.
- Brook 62b ———, ET AL. Trees and routines. *Comput. J.* 5 (1962), 33-47.
- Brook 63 ———, ET AL. The compiler-compiler. *Annual Review in Automatic Programming*, Vol. 3, 1963, p. 229.
- Brook 67a ———, MORRIS, D., AND ROHL, J. S. Compiler-compiler facilities in Atlas Autocode. *Comput. J.* 9 (1967), 350-352.
- Brook 67b ———, ———, AND ———. Experience with the compiler-compiler. *Comput. J.* 9 (1967), 345-349.
- Brook 67c ———, AND ROHL, J. S. Simply partitioned data structures: The compiler-compiler reexamined. *Machine Intelligence I*, Collins and Michie, (Eds.), Oliver and Boyd, London, 1967.
- Bruf 67 BROWN, P. J. The ML/I macro processor. *Comm. ACM* 10 (Oct. 1967), 618-623.
- BroS 63 BROWN, S. A., DRAYTON, C. E., AND MITTMAN, B. A description of the APT language. *Comm. ACM* 6 (Nov. 1963), 649-658.
- Burg 64 BURRO, W. H. The evaluation, classification and interpretation of expressions. *Proc. ACM 19th Nat. Conf.*, 1964, p. A14.
- Burk 65 BURKHARDT, W. Universal programming languages and processors. *Proc. AFIPS 1965 FJCC*, Vol. 27, pp. 1-21.
- Burs 65 BURSTELL, R. M. Some aspects of CPL semantics. No. 3, Experimental Programming Rpts., Edinburgh U., Edinburgh, April, 1965.
- Can 62 CANTOR, D. G. On the ambiguity problem of Backus Systems. *J. ACM* 9 (Oct. 1962), 477-479.
- Car 62 CARACIOLLO DI FORINO, A. On a research project in the field of languages for processor construction. *Proc. IFIP Congress*, Munich, 1962, pp. 514-515.
- Car 63 ———. Some remarks on the syntax of symbolic programming languages. *Comm. ACM* 6 (Aug. 1963), 456.
- Cas 66 CARTLE, J. A command program compiler. General Electric MSD, King-of-Prussia, Pa., 1966.
- Che 64a CHEATHAM, T. E. The architecture of compilers. CAD-64-2-R, Computer Associates, Inc., Wakefield, Mass., 1964.
- Che 64b ———, ET AL. Preliminary description of the translator generator system, II. CA-64-1-SD, Computer Associates, Inc., Wakefield, Mass., 1964.
- Che 64c ———, AND SATTLEY, K. Syntax directed compiling. *Proc. AFIPS 1964 SJCC*, Vol. 25, pp. 31-57.
- Che 65 ———. The TGS-II translator-generator system. *Proc. IFIP Congress*, New York, 1965, pp. 592-593.
- Che 66 ———. The introduction of definitional facilities into higher level programming languages. *Proc. AFIPS 1966 FJCC*, Vol. 29, pp. 623-637.
- Chom 63 CHOMSKY, N. Formal properties of grammars. In *Handbook of Mathematical Psychology*, Vol. 2, Luce, Bush and Galanter (Eds.), John Wiley & Sons, Inc., 1963, pp. 323-418.
- Chom 65 ———. *Aspects of the Theory of Syntax*. The MIT Press, Cambridge, Mass., 1965.
- Chu 51 CHURCH, A. *The Calculi of Lambda-Conversion*. Annals of Math. Studies, No. 6, Princeton U. Press, Princeton, N. J., 1951.
- Cla 66 CLAPP, L. A syntax-directed approach to automated aids for symbolic math. Summary in *Comm. ACM* 9 (Aug. 1966), 549.
- Coh 67 COHEN, J. A use of fast and slow memories in list processing languages. *Comm. ACM* 10 (Feb. 1967), 82-86.

- Col 67 COLES, S. Syntax-directed interpretation of natural language. Doctoral dissertation, Carnegie-Mellon Inst., Pittsburgh, Pa., 1967.
- Conn 66 CONNORS, T. B. ADAM—a generalized data management system. Proc. AFIPS 1966 SJCC, Vol. 28, pp. 193-203.
- Con 63 CONWAY, M. E. Design of a separable transition-diagram compiler. *Comm. ACM* 6 (July 1963), 396.
- Con 66 COULOURIS, G. F., AND GOODEY, T. J. The CPL1 system manual. PID12/GFC, Inst. of Comput. Sci., U. of London, London.
- Con 67 ———. The compiler processor project. Internal Rep., Imperial College, London, April 1967.
- Cra 66 CRAIG, J. A., BEREZNER, S. C., GARNEY, H. C., AND LONGYEAR, C. R. DEACON: Direct English Access and CONTROL. Proc. AFIPS 1966 FJCC, Vol. 29, pp. 365-380.
- Cul 62 CULIK, K. Formal structure of ALGOL and simplification of its description. *Symbolic Languages in Data Processing*, Gordon and Breach, New York, 1962, pp. 75-82.
- Cur 58 CURRY, H. B., AND FEYS, R. *Combinatory Logic*, Vol. 1. North-Holland Publishing Co., Amsterdam, 1958.
- Den 65 DENNIS, J. B. Segmentation and the design of multiprogrammed computer systems. *J. ACM* 12 (Oct. 1965), 589-602.
- Dij 63 DIJKSTRA, E. W. On the design of machine independent programming languages. *Annual Review in Automatic Programming*, Vol. 3, 1963, pp. 27-42.
- Dij 65 ———. Solution of a problem in concurrent programming control. *Comm. ACM* 8 (Sept. 1965), 569.
- Don 67 DONOVAN, J. J., AND LEDGARD, H. F. Canonic systems and their application to programming languages. Mem. MAC-M-347, Project MAC, MIT, Cambridge, Mass., April, 1967.
- Ear 65 EARLEY, J. C. Generating a recognizer for a BNF grammar. Computation Center Rep., Carnegie Inst. of Tech., Pittsburgh, Pa., 1965.
- Ear 67 ———. An LR(K) parsing algorithm. Carnegie Inst. of Tech., Pittsburgh, Pa., 1967, (mimeo).
- Ei 63 EICKEL, J., PAHL, M., BAUER, F. L., AND SAMUELSON, K. A syntax controlled generator of formal language processors. *Comm. ACM* 6 (Aug. 1963), 451-455.
- Ei 64 ———. Generation of parsing algorithms for Chomsky 2-type languages. 6401, Mathematisches Institut der Technischen Hochschule, Munich, 1964.
- Elg 64 ELGER, C. C., AND ROBINSON, A. Random-access stored-program machines, an approach to programming languages. *J. ACM* 11 (Oct. 1964), 365-390.
- Eng 61 ENGLUND, D., AND CHAIR, E. The CLIP-translator. *Comm. ACM* 4 (Jan. 1961), 19-22.
- Eva 64 EVANS, ARTHUR. An ALGOL 60 compiler. *Annual Review in Automatic Programming*, Vol. 4, 1964, pp. 87-124.
- Evt 66 EVANS, T., AND DARLEY, D. On-line debugging techniques: a survey. Proc. AFIPS 1966 FJCC, Vol. 29, pp. 37-50.
- Feld 61 FELDMAN, J. A. A formal semantics for computer oriented languages. Carnegie Inst. of Tech., Pittsburgh, Pa., 1961.
- Feld 66 ———. A formal semantics for computer languages and its application in a compiler-compiler. *Comm. ACM* 9 (Jan. 1966), 3-9.
- Feld 67 ———, AND CORRY, J. The compiler compiler in a time sharing environment. Lecture notes on Advanced Computer Organization, U. of Michigan, Ann Arbor, Mich., 1967.
- Fer 66 FERNUSON, D. E. Evolution of the meta assembly program. *Comm. ACM* 9 (March 1966), 190-196.
- Fie 67 FIERER, J. CARAI. Memos. Computer Center Rep., Carnegie Inst. of Tech., Pittsburgh, Pa., 1967.
- Flo 61 FLOYD, R. W. A descriptive language for symbol manipulation. *J. ACM* 8 (Oct. 1961), 579-584.
- Flo 62a ———. On ambiguity in phrase structure languages. *Comm. ACM* 5 (Oct. 1962), 526, 534.
- Flo 62b ———. On the nonexistence of a phrase structure grammar for ALGOL 60. *Comm. ACM* 5 (Sept. 1962), 493-494.
- Flo 63 ———. Syntactic analysis and operator precedence. *J. ACM* 10 (July 1963), 316-333.
- Flo 64a ———. Bounded context syntactic analysis. *Comm. ACM* 7 (Feb. 1964), 62-67.
- Flo 64b ———. The syntax of programming languages: a survey. *IEEE Trans. ECTD*, 4 (Aug. 1961), 316-353.
- Flo 67 ———. Assigning meanings to programs. *AMS Symposium in Appl. Math.* 19, 1967.
- Gall 67 GALLER, R., AND PERLIS, A. J. A proposal for definitions in ALGOL. *Comm. ACM* 10 (April 1967), 204-210.
- Gar 64 GARWICK, J. V. GARGOYLE, a language for compiler writing. *Comm. ACM* 7 (Jan. 1964), 16.
- Gar 66 ———, BELL, J., AND KROGER, L. The GPL language. TER 05, Control Data, Palo Alto, Calif., 1966.
- Ger 65 GERK, C. W. High speed compilation of efficient object code. *Comm. ACM* 8 (Aug. 1965), 483-488.
- Gil 66 GILBERT, P. On the syntax of algorithmic languages. *J. ACM* 19 (Jan. 1966), 90-107.
- Gil 67 ———, AND McKEELAN, W. G. Compiler generation using formal specification of procedure oriented and machine languages. Proc. AFIPS 1967 SJCC, Vol. 30, pp. 417-455.
- Gin 66a GINSBURG, S. *The Mathematical Theory of Context Free Languages*. McGraw-Hill, Inc., New York, 1966.
- Gin 66b ———, AND GIERBACH, S. Deterministic context free languages. *Inf. Contr.* 9 (1966), 629-638.
- Gle 60 GLENNIE, A. E. On the syntax machine and the construction of a universal compiler. Tech. Rpt. No. 2, Computation Center, Carnegie Inst. of Tech., Pittsburgh, Pa., 1960.
- Gor 61 GORN, S. Specification languages for mechanical languages and their processors, a baker's dozen. *Comm. ACM* 2 (Dec. 1961), 532-542.
- Gor 63 ———. Detection of generative ambiguities in context-free mechanical languages. *J. ACM* 10 (April 1963), 196-208.
- GrnM 65 GRAHAM, M. L., AND INDERMAN, P. Z. A universal assembly mapping language. Proc. ACM 20th Nat. Conf., 1965, pp. 109-120.
- GrnR 64 GRAHAM, R. M. Bounded context translation. Proc. AFIPS 1961 SJCC, Vol. 25, pp. 17-21.
- Gran 62 GRAD, A. A. A translator oriented symbolic programming language. *J. ACM* 9 (April 1962), 480-487.

- Gre 62 GREEN, J. Symposium on languages for processor construction. Proc. IFIP Congress, Munich, 1962, pp. 513-517.
- Grie 65 GRIES, D., PAUL, M., AND WIEHLE, H. R. Some techniques used in the ALCOR-ILLINOIS 7090. *Comm. ACM* 8 (Aug. 1965), 495-500.
- Grie 67a —. The use of transition matrices in compiling. Tech. Rpt. CS 57, Computer Science Dept., Stanford U., Stanford, Calif., March 1967 and *Comm. ACM* 11 (Jan. 1968), 26-34.
- Grie 67b —. Internal notes on the compiler writing system. Computer Science Dept., Stanford U., Stanford, Calif., 1967.
- Grif 65 GRIFFITHS, T. V., AND PETRICK, S. R. On the relative efficiencies of context-free grammar recognizers. *Comm. ACM* 8 (May 1965), 289-299.
- Gro 66 GROSS, M. Applications geometriques des langages formels. *ICC Bull.* 5 (Sept. 1966), 141-167.
- Hal 64 HALPERN, M. XPOP: a metalanguage without metaphysics. Proc. AFIPS 1964 FJCC, Vol. 26, pp. 57-68.
- Hal 66 —. Foundations of the case for natural language programming. Proc. AFIPS 1966 FJCC, Vol. 29, pp. 639-649.
- Hal 67a —. Toward a general processor for programming languages. *Comm. ACM* 11, (Jan. 1968), 15-26.
- Hal 67b —. Foundations of the case for natural language programming. *IEEE Spectrum* (March 1967), 140-149.
- Hal 67c —. A manual of the XPOP programming system. Electronic Sciences Lab. Lockheed Missiles & Space Company, Palo Alto, Calif., March 1967.
- Hals 62 HALSTEAD, M. H. *Machine-Independent Computer Programming*. Spartan Books, Washington, D.C., 1962.
- Hill 62 HILL, V., LANGMAACK, H., SCHWARZ, H. R., AND SEEGMÜLLER, G. Efficient handling of subscripted variables in ALGOL 60 compilers. Proc. Symbolic Languages in Data Processing, Gordon and Breach, New York, 1962, 331-340.
- Hoa 65 HOARE, C. A. R. A programming language for processor construction. Notes from NATO summer school lectures, 1966.
- Hor 66 HORWITZ, L. P., KARP, R. M., MILLER, R. E., AND WINGRAD, S. Index register allocation. *J. ACM* 15 (Jan. 1966), 43-61.
- Hus 62 HUSKEY, HARRY D. Languages for aiding compiler writing (panel discussion). Proc. Symbolic Languages in Data Processing, Gordon and Breach, New York, 1962, 187-204.
- IBM 66 IBM System/360 Operating System PL/I Language Specification. Form C28-6571-4.
- Ing 62 INGERMAN, P. Z. Techniques for processor construction. Proc. IFIP Congress, Munich 1962, pp. 527-528.
- Ing 66 —. *A Syntax Oriented Translator*. Academic Press, Inc., New York, 1966.
- Int 63 International Standards Organization. Survey of programming languages and processors. *Comm. ACM* 6 (March 1963), 93.
- Ir 61 IRONS, E. T. A syntax directed compiler for ALGOL 60. *Comm. ACM* 4 (Jan. 1961), 51-55.
- Ir 63a —. The structure and use of the syntax-directed compiler. *Annual Review in Automatic Programming*, Vol. 3, 1963, pp. 207-227.
- Ir 63b —. Towards more versatile mechanical translators. *AMS Symposium in Appl. Math.* 16, 1963, pp. 41-50.
- Ir 64 —. Structural connections in formal languages. *Comm. ACM* 7 (Feb. 1964), 67-71.
- Ir 65 —. An error correcting parse algorithm. *Comm. ACM* 6 (Nov. 1965), 669-673.
- It 68 ITURRIAGA, R., STANISH, T. A., KRUTAR, R. A., AND EARLEY, J. C. Techniques and advantages of using the formal compiler writing system FSL to implement a formula ALGOL compiler. Proc. AFIPS 1966 SJCC, Vol. 28, pp. 241-252.
- Kar 66 KARP, R. M., AND MILLER, R. E. Properties of a model for parallel computations: determinacy, termination, queueing. *SIAM J.* (Nov. 1966), 1390-1411.
- Kerr 67 KERR, R. H., AND CLEGG, J. The Atlas ALGOL Compiler—an ICT implementation of ALGOL using the Brooker-Morris syntax-directed compiler. *Comput. J.* (1967).
- Kir 66 KIRKLEY, C., AND RULIFSON, J. LOTS, a syntax-directed compiler. Internal Rep., Stanford Research Inst., Menlo Park, Calif., May 1966.
- Knu 62 KNUTH, D. E. History of writing compilers. Proc. ACM 17th Natl. Conf., 1962, pp. 43, 126.
- Knu 65 —. On the translation of languages from left to right. *Inf. Contr.* 8 (Oct. 1965), 607-639.
- Knu 67 —. The remaining trouble spots in ALGOL 60. *Comm. ACM* 10 (Oct. 1967), 611-618.
- Kuc 67 KUCK, D. Programming the ILLIAC IV. Talk given at AFIPS 1967 SJCC. Paper not yet available.
- Kun 62 KUNO, S., AND OETTINGER, A. G. Multiple-path syntactic analyzer. *Information Processing 68* (IFIP Congress), Popplewell (Ed.), North-Holland Publishing Co., Amsterdam, 1962, pp. 306-311.
- Lande 62 LANDEN, W. H., AND WATTENBURG, W. H. On the efficient construction of automatic programming systems. Proc. ACM 17th Natl. Conf., 1962, p. 91.
- Landi 63 LANDIN, P. J. The mechanical evaluation of expressions. *Comp. J.* 6 (1963), 308.
- Landi 65 —. A correspondence between ALGOL 60 and Church's  $\lambda$ -notation. *Comm. ACM* 8 (Feb. and March 1965), 89-101, 153-167.
- Landi 66 —. The next 700 programming languages. *Comm. ACM* 9 (March 1966), 157-166.
- Landw 64 LANDWEBER, P. S. Denison problems of phrase structure grammars. *IEEE Trans. EC*, 13 (Aug. 1964), 354-362.
- Lang 64 LANGMAACK, H., AND EICKEL, J. Präzisierung der begriffe Phrasenstruktur und strukturelle Mehrdeutigkeit in Chomsky-Sprachen. Rep. No. 6414, Rechenzentrum der Technischen Hochschule, Munich, 1964.
- Lea 64 LEAVENWORTH, B. M. FORTRAN IV as a syntax language. *Comm. ACM* 7 (Feb. 1964), 72-80.
- Lea 66 —. Syntax macros and extended translation. *Comm. ACM* 9 (Nov. 1966), 790-793.
- Leo 66 LEONARD, G., AND GOODROE, J. More extensible machines. *Comm. ACM* 9 (March 1966), 183-188.
- Let 65 LETICHEVSKII, A. A. The representation of context-free languages in automata with a push down type store. *Cybernetics* (Kibernetika) Vol. 1, No. 2, The Faraday Press, New York (1965), 81-86.

- Luc 65 LUCAS, P. Definition of a subset of PL/I by finite local state vectors. Working paper to IFIP WG2.1, July, 1965.
- Luc 67 LUCIO, F. A comment on index register allocation. *Comm. ACM* 10 (Sept. 1967), 572-574.
- Mark 61 MARKOV, A. A. *Theory of Algorithms*. US Bureau of Standards, OITS 60-51085, available from Clearing House, Springfield, Va.
- Mar 67 MARTIN, D., AND ESTRIN, G. Models of computations and systems. *J. ACM* 14 (April 1967), 281-294.
- Mis 60 MARTERSON, K. S. Compilation for two computers with NELIAC. *Comm. ACM* 3 (Nov. 1960), 607-611.
- Maui 65 MAURER, W. A theory of computer instructions. Mem. MAC-M-282, Project MAC, MIT Cambridge, Mass., Sept. 1965.
- May 61 MAYOH, B. H. Letter to the editor correcting Ir 61. *Comm. ACM* 4 (June 1961), 284.
- McCar 62a MCCARTHY, J. Toward a science of computation. *Information Processing 68* (IFIP Congress), Paplewelt (Ed.). North-Holland Publishing Co., Amsterdam, 1962, pp. 21-28.
- McCar 62b — ET AL. LISP 1.5 programmers manual. Computation Lab Report, MIT (1962).
- McCar 67 —, AND PAINTER, J. Correctness of a compiler for arithmetic expressions. *AMS Symposium in Appl. Math.* 19, 1967.
- McCl 65 McCLELL, R. M. TMG: a syntax-directed compiler. *Proc. ACM 20th Natl. Conf.*, 1965, pp. 262-274.
- McIl 60 McILROY, M. D. Macro instruction extension of compiler language. *Comm. ACM* 3 (April 1960), 214-220.
- McKee 64 McKEEMAN, W. M. An approach to computer language design. Tech. Rpt. CS 48, Computer Science Dept., Stanford U., Stanford, Calif., Aug. 1965.
- Men 63 MEALY, G. A generalized assembly system. RM-3646-PR Rand Corp., Santa Monica, Calif., Aug. 1963.
- Met 64 METCALFE, H. H. A parametrized compiler based on mechanical linguistics. *Annual Review in Automatic Programming*, Vol. 4, 1964, pp. 125-165.
- Min 67 MINSKY, M. L., AND PAPER, S. Perceptions and pattern recognition. Mem. MAC-M-358, Project MAC, MIT, Cambridge, Mass., Sept. 1967.
- Mond 67 MONDSCHEN, L. VITAL compiler compiler reference manual. TN 1967-1, Lincoln Lab., MIT, Lexington, Mass., Jan. 1967.
- Mon 65 MOORE, C., AND DRETTEN, L. P. TRAC, a text handling language. *Proc. ACM 20th Natl. Conf.*, 1965, pp. 229-246.
- Mor 67 MORRIS, D., AND WILSON, I. A system program generator. Computer Science Dept., U. of Manchester, Manchester, 1967.
- Nap 67 NAPIER, R. B. E. The third order compiler. A context for free man-machine communication. *Machine Intelligence 1*, Collins and Michie (Eds.), Oliver and Boyd, London, 1967.
- Nar 66 NARASIMHAN, R. Syntax-directed interpretation of classes of pictures. *Comm. ACM* 9 (March 1966), 166-173.
- Nar 67 —. Programming languages and computers: a unified meta theory. In *Advances in Computer Science*. Academic Press, Inc., New York, 1967, Chap. 5.
- Naur 60 NAUR, P. (Ed.) Report on the algorithmic language ALGOL 60. *Numer. Math.* 2 (1969), 106-136; *Comm. ACM* 3 (May 1960), 280-31.
- Naur 63a —. Documentation problems: ALGOL 60. *Comm. ACM* 6 (March 1963), 77-79.
- Naur 63b —. Revised report on the algorithmic language ALGOL 60. *Comm. ACM* 6 (Jan. 1963), 1-17; *Numer. Math.* 4 (1963), 420-452; *Comp. J.* 3 (1963), 349-367.
- Nob 63 NOBLE, A. S., AND TALMADGE, R. B. Design of an integrated programming and operating system, I and II. *IBM Syst. J.* 2 (June 1963), 152-181.
- Nor 63 NORTHCOPE, R. S. The structure and use of a compiler-compiler system. *Proc. Australian Comput. Conf.*, Dec. 1963.
- Op 62 OPLER, A. "Tool": a processor construction language. *Proc. IFIP Congress*, Munich, 1962, pp. 513-514.
- Orch 66 ORCHARD-HAYS, WILLIAM. Multilevel operating systems. *Comm. ACM* 9 (March 1966), 189-190, (abstract only).
- Org 67 ORGANS, R. J. A mathematical theory of computing machine structure and programming. RC1863, IBM, Yorktown Heights, New York, 1967.
- Pai 66 PAINTER, J. Semantic correctness of a compiler for an ALGOL-like language. AI Rep. No. 43, Stanford U., Stanford, Calif., 1966.
- Par 61 PARIKH, R. J. Language generating devices. Quarterly progress rep. no. 60, Research Lab. of Electronics, MIT, Jan. 1961, 199-212. Reprinted with minor editorial revisions under the title: On context-free languages. *J. ACM* 15 (Oct. 1966), 570-581.
- Paul 62 PAUL, M. ALGOL 60 processors and a processor generator. *Proc. IFIP Congress*, Munich, 1962, pp. 493-497.
- Plas 66 PLASKOW, J., AND SCHUMAN, S. The TRACON system on the M460 computer. AFCTR-66-516 (July 1966).
- Pos 43 POST, E. L. Formal reductions of the general combinatorial decision problem. *Am. J. Math.* 66 (1943), 197-217.
- Pra 65 PRATT, T. W. Syntax-directed translation for experimental programming languages. TNN-11, Computation Center U. of Texas, Austin, Texas, 1965.
- Pra 66 —, AND LINDSAY, R. R. A processor-building system for experimental programming language. *Proc. AFIPS 1966 FJCC*, Vol. 29, pp. 613-621.
- Rab 62 RABINOWITZ, I. N. Report on the algorithmic language FORTRAN II. *Comm. ACM* 5 (June 1962), 327-337.
- Ran 64 RANDALL, B., AND BUSSEL, D. J. *ALGOL 60 Implementation*. Academic Press, Inc., London, 1964.
- Rey 65 REYNOLDS, J. C. An introduction to the CORGENT programming system. *Proc. ACM 20th Natl. Conf.*, 1965, pp. 422-436.
- Ri 66 RICE, J., AND ROSEN, S. NAPSS, numerical analysis and problem solving system. *Proc. ACM 21st Natl. Conf.*, 1966, pp. 51-56.
- Rig 62 RIQUET, J. Programmation et theories des categories. *Proc. Rome Symposium on Symbolic Languages in Data Processing*, Gordon and Breach, New York, (1962), pp. 83-98.
- Rob 66 ROBERTS, L. G. A graphical service system with variable syntax. *Comm. ACM* 9 (March 1966), 173-176.
- Ros 64a ROSEN, S. A compiler-building system developed by Brookner and Morris. *Comm. ACM* 7 (July 1964), 403-414.
- Ros 64b —. Programming systems and languages. *Proc. AFIPS 1964 SJCC*, Vol. 25, pp. 1-15.

- Ross 63 ROSS, D., AND RODRIGUEZ, J. Theoretical foundations of the computer aided design system. Proc. AFIPS 1963 SJCC, Vol. 23, pp. 306-322.
- Ross 61 ROSS, D. T. On context and ambiguity in parsing. *Comm. ACM* 7 (Feb. 1964), 131-133.
- Ross 66 . AED bibliography. Mem. MAC-M-278-2, Project MAC, MIT Cambridge, Mass., Sept. 1966.
- Ross 67 . The AED approach to generalized computer aided design. Proc. ACM 22nd Natl. Conf. 1967, pp. 367-385.
- Row 67 ROYKEL, P., AND FREEDMAN, J. An associative processing system for conventional digital computers. TN 1967-19, Lincoln Lab., MIT, Lexington, Mass., April 1967.
- Rut 62 RUTENHAUSER, H. Panel on techniques for processor construction. Proc. IFIP Congress, Munich, 1962, pp. 524-531.
- Sam 60 SAMELSON, K., AND BAUER, F. L. Sequential formula translation. *Comm. ACM* 3 (Feb. 1960), 76-83.
- Sam 62 . Programming languages and their processing. Proc. IFIP Congress, Munich, 1962, pp. 487-492.
- Samu 61 SAMMER, J. E. A definition of COBOL 61. Proc. ACM 16th Natl. Conf., 1961.
- Schl 67 SCHLESINGER, S., AND SASHKIN, L. POSE: a language for posing problems to a computer. *Comm. ACM* 10 (May 1967), 279-285.
- Schm 63 SCHMIDT, L. Implementation of a symbol manipulator for heuristic translation. Proc. ACM 18th Natl. Conf., 1963.
- Sch 64 SCHNEIDER, F. W., AND JOHNSON, G. D. META-3: A syntax directed compiler writing compiler to generate efficient code. Proc. ACM 19th Natl. Conf. 1964, p. 141.5-1.
- Scho 65 SCHOTT, H. Analytic differentiation using a syntax directed compiler. *Comput. J.* 7 (Jan. 1965), 290-298.
- Schor 61 SCHROEDER, D. V. META-11: A syntax-oriented compiler writing language. Proc. ACM 19th Natl. Conf., 1964, p. 141.3.
- Schü 63 SCHUTZENBERGER, M. P. Context-free languages and pushdown automata. *Inf. Comb.* 6 (Sept. 1963), 246-261.
- Sh 58 SHABBE Ad Hoc Committee on Universal Languages. The problem of programming communication with changing machines: a proposed solution. *Comm. ACM* 1 (Aug. 1958), 12-18.
- ShaA 67 SHAW, A. A formal description and parsing of pictures. Doctoral dissertation, Computer Science Dept., Stanford U., Stanford, Calif., 1968.
- Shaw 63 SHAW, C. J. A specification of JOVIAL. *Comm. ACM* 6 (Dec. 1963), 721-735.
- Shed 67 SHENKER, G. Parallel numerical methods for the solution of equations. *Comm. ACM* 10 (May 1967), 286-291.
- Sib 61 SIRLEY, R. A. The SLANG-system. *Comm. ACM* 4 (Jan. 1961), 75-81.
- Sta 67 STANDEN, T. A. A data definition facility for programming languages. Computer Science Rep., Carnegie Inst. of Tech. Pittsburgh, Pa., May 1967.
- Sto 65 STARK, R. ALTEXT: multiple purpose language. Lockheed Missiles & Space Company Tech. Rep. 6-75 65-15, March, 1965.
- Ste 61 STEIN, T. B. A first version of UNCOL. Proc. Western Joint Comput. Conf., 1961, pp. 371-378.
- Ste 66 . (Ed.) Formal language description languages for computer programming. Proc. IFIP Conf., Baden, Sept. 1964, North-Holland Publishing Co., Amsterdam, 1966.
- Stu 67 STONE, H. S. One-pass compilation of arithmetic expressions for parallel processor. *Comm. ACM* 10 (April 1967), 220-223.
- Str 65 STRACHEY, C. A general purpose macrogenerator. *Comput. J.* 8 (1965), 225-241.
- Tar 56 TARSKI, A. *Logic, Semantics, Metamathematics*. Clarendon Press, London, 1956.
- Tay 61 TAYLOR, W., TURNER, L., AND WAYHOFF, R. A syntactical chart of ALGOL 60. *Comm. ACM* 14 (Sept. 1961), 393.
- Te 66 TEICHROW, D., AND LUBIN, J. F. Computer simulation—discussion of the technique and comparison of languages. *Comm. ACM* 9 (Oct. 1966), 727-741.
- Th 66 THOMPSON, F. B. English for the computer. Proc. AFIPS 1966 FJCC, Vol. 29, pp. 349-356.
- Tix 67 TIXIER, V. Recursive functions of regular expressions in language analysis. Tech. Rpt. CS 58, Computer Science Dept., Stanford U., Stanford, Calif., March 1967.
- Tro 67 TROUT, R. G. A compiler-compiler system. Proc. ACM 22nd Natl. Conf., 1967, pp. 317-322.
- Wai 67 WAITE, W. A language independent macro processor. *Comm. ACM* 10 (July 1967), 433-441.
- War 61 WARSHALL, S. A syntax directed generator. Proc. Eastern Joint Comput. Conf., 1961, pp. 295-305.
- War 64 . AND SHAPIRO, R. M. A general-purpose table-driven compiler. Proc. AFIPS 1964 SJCC, Vol. 25, pp. 59-65.
- Weg 62 WEGNER, P. (Ed.) *Introduction to Systems Programming*. Academic Press, Inc., New York, 1962.
- Wij 66 VAN WIJNGAARDEN, A. Recursive definition of syntax and semantics in *Formal Language Description Languages for Computer Programming*, North-Holland Publishing Co., Amsterdam, 1966, pp. 13-24.
- Wil 61a WILKES, M. V. An experiment with a self-compiling compiler for a simple list-processing language. *Annual Review in Automatic Programming*, Vol. 4, 1964, pp. 1-48.
- Wil 61b . Constraint-type statements in programming languages. *Comm. ACM* 7 (Oct. 1964), 587.
- Wir 60a WIRTH, N. A programming language for the 360 computers. Tech. Rpt. CS 53, Computer Science Dept., Stanford U., Stanford, Calif., Dec. 1966.
- Wir 66b . AND HOARE, C. A. R. A contribution to the development of ALGOL. *Comm. ACM* 9 (June, 1966), 413-432.
- Wir 66c . AND WEBER, H. EULER—a generalisation of ALGOL, and its formal definition: Part I, Part II. *Comm. ACM* 9 (Jan., Feb. 1966), 13-25, 89-99.
- Yer 65 YERSHOV, A. P. ALPHA—an automatic programming system of high efficiency. Proc. IFIP Congr., New York, 1965, pp. 622-623.
- You 65 YOUNG, J. W., JR. Nonprocedural languages. 7th Ann. ACM Tech. Symp., S. Calif. Chapter, March 1965.
- Zar 67 ZARA, R. F. A semantic model for a language processor. Proc. ACM 22nd Natl. Conf., 1967, pp. 323-339.
- Zem 66 ZEMANEE, H. Semiotics and programming languages. *Comm. ACM* 9 (March 1966), 130-143.





— 禁 無 断 転 載 —

昭和48年 3 月発行

発行所 財団法人 日本情報処理開発センター

東京都港区芝公園 3-5-8

機械振興会館内

TEL (434)8211(代表)

印刷所 株式会社 三 州 社

東京都港区芝大門一丁目一番二十一号

TEL (433)1481(代)

