

TR-627

YAGLR法 : Yet Another Generalized  
LR Parser

田中 穂積、G. スレッシュユ  
(東京工業大学)

March, 1991

© 1991, ICOT

**ICOT**

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

(03)3456-3191 ~ 5  
Telex ICOT J32964

---

**Institute for New Generation Computer Technology**

1990. 2. 6

## YAGLR法: Yet Another Generalized LR Parser

田中穂積, K. G. スレッシュ  
(東京工業大学・工学部)

### 概要

プログラム言語などの人工言語のコンパイラでは, Knuthの考案したLR(k)法とよぶ統語解析アルゴリズム[Knuth 65]が使われることがある. これは先読み情報を用いて決定的に統語解析を進め, 無駄なく効率的に統語解析を行うことができる. ただしそこで使われる文法は, LR文法に限られており, 自然言語の解析に用いる一般の文脈自由文法(CFG)を扱うことができない.

富田はLR(k)法の持つ利点をそのまま保持し, しかも一般のCFGが扱えるようにLR(k)法を拡張している[Tomita 86, 87]. これは一般化LR法(generalized LR parsing algorithm; 拡張LR法)の1つであり, 富田法とよばれている. 経験的に富田法はアーリー法と比べて高速な統語解析が可能であるが[Tomita 86], 富田法に対して, アーリー法以上のオーダーの解析時間を要す特殊なCFGの存在が知られている.

我々は, アーリー法の利点と富田法の利点をそのまま引き継ぐ新しい一般化LR法(これをYAGLR法とよぶ)を提案する. それによれば, 先に述べた特殊なCFGに対する解析時間のオーダーを $n^3$ に抑えることができる. また, 富田法ではグラフ構造化スタック(GSS)を用いて, 使用記憶空間の節約を図っているが, GSSの構造は必ずしも単純でなく, 実装が容易ではない. これに対してYAGLR法では木構造化スタック(TRSS)を用いており, 経験的に, 富田法のGSSと比較して同程度の使用記憶空間に抑えることが可能であり, 構造が単純であるため扱い易い. これは, YAGLR法では富田法以上にスタックのマージが可能になるためである. 本論文では, YAGLR法による統語解析アルゴリズムの定義と動作例を説明する. また, 富田法と比較して, YAGLR法はどのような利点があるかを具体例を用いて説明する.

田中穂積, K. G. スレッシュ  
東京工業大学・工学部

## 概要

プログラム言語などの人工言語のコンパイラでは, Knuthの考案したLR(k)法とよぶ統語解析アルゴリズム [Knuth 65] が使われることがある. これは先読み情報を用いて決定的に統語解析を進め, 無駄なく効率的に統語解析を行うことができる. ただしそこで使われる文法は, LR文法に限られており, 自然言語の解析に用いる一般の文脈自由文法(CFG)を扱うことができない.

富田はLR(k)法の持つ利点をそのまま保持し, しかも一般のCFGが扱えるようにLR(k)法を拡張している [Tomita 86, 87]. これは一般化LR法 (generalized LR parsing algorithm: 拡張LR法) の1つであり, 富田法とよばれている. 経験的に富田法はアーリー法と比べて高速度な統語解析が可能であるが [Tomita 86], 富田法に対して, アーリー法以上のオーダの解析時間を要す特殊なCFGの存在が知られている.

我々は, アーリー法の利点と富田法の利点をそのまま引き継ぐ新しい一般化LR法 (これをYAGLR法とよぶ) を提案する. それによれば, 先に述べた特殊なCFGに対する解析時間のオーダを $n^3$ に抑えることができる. また, 富田法ではグラフ構造化スタック(GSS)を用いて, 使用記憶空間の節約を図っているが, GSSの構造は必ずしも単純でなく, 実装が容易ではない. これに対してYAGLR法では木構造化スタック(TRSS)を用いており, 経験的に, 富田法のGSSと比較して同程度の使用記憶空間であり, 構造が単純であるため扱い易い. これは, YAGLR法では富田法以上にスタックのマージが可能になるためである. 本論文では, YAGLR法による統語解析アルゴリズムを与え, 富田法との比較を具体例により説明する.

## [Abstract]

We have developed a new generalized LR parsing algorithm called YAGLR in which it combines some of the advantages of both Earley's and Tomita's algorithms. We show that for some special CFGs the time consumed by YAGLR parsing algorithm is less than Tomita's parsing algorithm and is in the order of  $O(n^3)$ . While YAGLR uses a tree structured stack, the structure is not so complex compared to the one of Tomita's graph-structured stack. The reason is that YAGLR parsing algorithm enables us to make much more merge operations on the tree-structured stack.

## 1. はじめに

プログラム言語などの人工言語のコンパイラでは、Knuthの考案したLR(k)法とよぶ統語解析アルゴリズム[Knuth 65]が使われることがある。これは先読み情報を用いて決定的に統語解析を進め、無駄なく効率的に統語解析を行うことができる。ただしそこで使われる文法は、LR文法に限られており、自然言語の解析に用いる一般の文脈自由文法(CFG)を扱うことができない。

富田はLR(k)法の持つ利点をそのまま保持し、しかも一般のCFGが扱えるようにLR(k)法を拡張している[Tomita 86, 87]。これは一般化LR法(generalized LR parsing algorithm; 拡張LR法)の1つであり、富田法とよばれている。経験的に富田法はアークリー法と比べて高速な統語解析が可能であるが[Tomita 86]、富田法に対して、解析すべき文の長さを $n$ としたとき、 $n^m$  ( $m > 3$ )のオーダーの解析時間を要する特殊なCFGの存在が知られている[Johnson 89][Kipps 89]。一方、アークリー法は、いかなるCFGに対しても、解析時間のオーダーが $n^3$ であるから、この特殊なCFGに対して、富田法はアークリー法以上の解析時間を要することになる。

本稿で提案する新しい一般化LR法では、アークリー法の利点と富田法の利点をそのまま引き継ぐ。この新しい一般化LR法をYAGLR法(Yet Another Generalized LR parsing)とよんでいる。YAGLR法によれば、先に述べた特殊なCFGに対する解析時間のオーダーを $n^3$ に抑えることができる。また、富田法ではグラフ構造化スタック(GSS)を用いて、使用記憶空間の節約を図っているが、GSSの構造は必ずしも単純でなく、実装が容易ではない。これに対してYAGLR法では本構造化スタック(TRSS)を用いるが、経験的に、富田法のGSSと比較して使用記憶空間は同程度であり、構造が単純であるため扱い易い。これは、YAGLR法では富田法以上にスタックのマージが可能になるためである。

YAGLR法では、解析中に統語解析木は作らず、統語解析木を構成する部品のみを作成する。この部品を逆ドット項と呼ぶが、これはアークリー法で作成するドット項(アークリー項)と対称な構造をしている。2章では、逆ドット項を説明する。3章では、YAGLR法の鍵となるTRSSのマージ操作(統合化操作)、シフト操作、レデュース操作を説明し、最後にYAGLR法の統語解析アルゴリズムを説明する。レデュース操作の過程で逆ドット項を作成するが、逆ドット項を作成することの利点も3章で説明する。また富田法とYAGLR法による解析例を示し、富田法とYAGLR法の比較を行う。4章では、特殊なCFGに対しても、YAGLR法の統語解析時間のオーダーが $n^3$ であることを証明する。5章では、まとめと今後の研究課題について述べる。

## 2. 逆ドット項(Dot Reverse Item)

本章では、YAGLR法による統語解析で重要な役割を果たす逆ドット項を説明する。富田法がレデュース操作時に部分統語解析結果を圧縮統語森として作成するのに対して、YAGLR法ではアークリー項と対称な逆ドット項の集合を作成する。逆ドット項の定義を以下に与える。

- (1) 解析対象文 $w$ を $w = w_1 w_2 \dots w_n \in T^*$ であるとすると、ここに $T$ は、終端記号の集合 $T$ の要素を0個以上並べたものを表す。終端記号(語と考えるとよい) $w_i$ と $w_{i+1}$ の間に番号 $i$ を振り、これを位置番号とよぶ( $i$ を、語 $w_i$ の位置番号とよぶ)。ただし $w_1$ の左には位置番号0を、また $w_n$ の右には位置番号 $n$ を振る。
- (2) CFG規則 $A \rightarrow \alpha$ があり、 $\alpha$ が $\beta \gamma$ と書けるとき、 $0 \leq j \leq n$ なる $j$ に対して、 $[A \rightarrow \beta \cdot \gamma, j]$ は、 $w$ に対する逆ドット項である。特に、 $[A \rightarrow \alpha \cdot, j]$ または $[A \rightarrow \alpha \cdot, j]$ も $w$ に対する逆ドット項である。
- (3) 逆ドット項の集合 $I$ を次のように定義する。 $0 \leq i \leq j \leq n$ なる $i$ と $j$ に対して、 $[A \rightarrow \alpha \cdot \beta, j] \in I$  iff  $S \Rightarrow \gamma A \delta$ ,  $\beta \Rightarrow w_{i+1} w_{i+2} \dots w_j$ , and  $\delta \Rightarrow w_{j+1} w_{j+2} \dots w_n$

アークリー項も逆ドット項も、ドットの位置は項の集合 $I$ の添字 $i$ として表す。アークリー項と逆ドット項の違いは、項中の位置番号 $j$ の解釈にある。上記した(3)の定義から明らかに、逆ドット項では、ドット記号の右の部分 $\beta$ として解析済みであるとみなす。項内のCFG規則の右辺の直後の位置番号が $j$ であり、そこから先頭方向に $i$ まで逆戻りした部分 $w_{i+1} w_{i+2} \dots w_j$ が $\beta$ として解析済みであると解釈する。ちなみにアークリー項では解析対象文の $\alpha$ が解析済みであるとされていた。逆ドット項を作成する利点は3.6で説明する。

## 3. YAGLR法

本章ではTRSSの構造と節点のマージ操作を説明し、次にシフト操作、レデュース操作を説明する。最後にYAGLR法による統語解析アルゴリズムを説明する。

### 3.1 識別子付き位置番号(Position Number with an Identifier)

YAGLR法で用いるTRSSの各節点には、レデュース操作時に逆ドット項、 $[A \rightarrow \alpha \cdot \beta, j] \in I$ を作成するために、状態の他に、 $j$ と $i$ を得る情報が含まれていなければならない。さらに、マージを効率よく行うために、識別子付き位置番号を導入する。識別子付き位置番号は全解析過程でユニークな番号であり、後述のマージ操作、シフト操作、レデュース操作により、スタックトップの節点が新たに作られる度に作り出され、スタックトップの節点の位置番号にこれを付加する。識別子が付加された位置番号のことを識別子付き位置番号と呼ぶ。位置番号が $i$ で識別子が $x$ なら、この識別子付き位置番号を $x \cdot i$ と書くことにする。識別子はマージの歴史を保存しておくためのもので、それにより重複した統合操作を避けることができる(3.2)。レデュース操作時の逆ドット項の生成には、識別子付き位置番号のなかの、位置番号のみが使われることに注意したい(3.3)。

YAGSSで用いるTRSSの節点の構造を以下に示す：

[識別子付き位置番号の集合, 状態]

たとえば、 $\{^3 2, ^4 3, ^5 4\}$ ,  $\{^1 2, ^3 2, ^6 6\}$ は識別子付き位置番号の集合であり、 $\{(^3 2, ^4 3, ^5 4), 2\}$ ,  $\{(^1 2, ^3 2, ^6 6), 4\}$ はTRSSの節点である。

### 3.2 節点のマージ操作



### 〔レデュース操作〕

レデュース操作の指定するCFG規則  $A \rightarrow X_1 X_2 \cdots X_n$  により、スタックは(a)から(b)に変わると共に、(c)に示す逆ドット項を作成する。

$$(a) \cdots [N_k, s_k] - [N_{k+1}, s_{k+1}] - \cdots - [N_{k+n}, s_{k+n}] (\dagger \gamma \gamma')$$

↓

$$(b) \cdots [N_k, s_k] - [N, t] (\dagger \gamma \gamma')$$

ここで状態  $t$  は、レデュース操作後にCFG規則の左辺の非終端記号  $A$  と状態  $s_k$  とからGOTO表を参照して決めた新しい状態である。また、 $N_k = \{^0a, ^0b, \dots\}$ ,  $N_{k+1} = \{^0c, ^0d, \dots\}$ ,  $\dots$ ,  $N_{k+n-1} = \{^0e, ^0f, \dots, ^0g\}$ ,  $N_{k+n} = \{^0j\}$ ,  $N = \{^0j\}$  であるとする。スタックトップの節点の位置番号の集合  $N_{k+n}$  は、最後にシフトした語の識別子付き位置番号  $^0j$  を含み、レデュース後のスタックトップの節点の位置番号の集合  $N$  は、新しいユニークな識別子付き位置番号  $^0j$  を含むことに注意したい。これは、レデュース後のスタックのトップに新しい節点が作られるためである。レデュース操作の前後で位置番号  $j$  は変わらないことに注意したい。

(c) 逆ドット項の作成:

$$l_a \ni [A \rightarrow X_1 X_2 \cdots X_n, j]$$

$$l_b \ni [A \rightarrow X_1 X_2 \cdots X_n, j]$$

.....

$$l_c \ni [A \rightarrow X_1 X_2 \cdots X_n, j]$$

$$l_d \ni [A \rightarrow X_1 X_2 \cdots X_n, j]$$

.....

$$l_e \ni [A \rightarrow X_1 X_2 \cdots X_n, j]$$

$$l_f \ni [A \rightarrow X_1 X_2 \cdots X_n, j]$$

.....

$$l_g \ni [A \rightarrow X_1 X_2 \cdots X_n, j]$$

レデュース操作では、操作の指定するCFG規則の右辺の非終端記号の数だけ、スタックトップから節点をポップする。ポップする各節点には、CFG規則の右辺の非終端記号が対応するので、上記した(1)の場合、 $X_1, X_2, \dots, X_n$  のそれぞれには、スタックの節点、 $[N_{k+1}, s_{k+1}]$ ,  $[N_{k+2}, s_{k+2}]$ ,  $\dots$ ,  $[N_{k+n}, s_{k+n}]$  が順に対応する。TRSSに含まれる節点の位置番号は、どこまでがシフト済み(処理済み)かを示している。これを用いて逆ドット項のドット位置(逆ドット項の集合  $l$  の添字)を知ることができる。逆ドット項内に書かれる位置番号は全て、スタックトップの節点の持つ位置番号  $j$  であり、位置番号に付加された識別子は、逆ドット項の作成には寄与していないことに注意。

### 3.4 節点のマージ操作(M2)の妥当性

二つの節点  $[M, s]$  と  $[N, s]$  に対して(M2)が成立することは、次の(1), (2), (3)から証明できる。

- (1) YAGLR法では、全てのレデュース操作が終了してからシフト操作を行うため、TRSSの全てのスタックトップの節点の持つ位置番号は、レデュース操作中は最後にシフトした語の位置番号に等しく同一である。従って、レデュース操作で作成する逆ドット項中の位置番号はマージ前もマージ後と変わらず、最後にシフトした語の位置番号である。
- (2) 逆ドット項の集合  $l$  の添字は、逆ドット項中のドットの位置番号を表す。この添字は、TRSS中の節点の持つ位置番号の集合から1つずつ取り出されるもので、マージ前の位置番号の集合から個々に取り出しても、マージ後の1つの和集合から取り出しても変わらない。
- (3)  $M$  が  $N$  の部分集合でない場合に、統合前の二つの節点の持つ全ての子節点を、統合節点の子節点として良いことは明らかである。一方、 $M$  が  $N$  の部分集合である場合に、節点  $[M, s]$  の子節点をそのまま統合節点の子節点として良い理由は次の通りである。 $M$  が  $N$  の部分集合であるということは、節点  $[M, s]$  を得るために行ったのと全く同じマージの履歴が、節点  $[N, s]$  に残されていることを意味している。従って、節点  $[M, s]$  と節点  $[N, s]$  以降のマージの結果は、既に節点  $[N, s]$  の子節点以降に含まれているので、両節点の統合節点を  $[N, s]$  とし、その子節点を、 $[N, s]$  の子節点に一致させて良い。

### 3.5 YAGLR法による統語解析アルゴリズム

YAGLR法は以下のアルゴリズムにより入力文を解析する。

- (1) TRSSの初期状態を、 $(\dagger \dagger A) \mid \{^0 0\}, 0 (\dagger \gamma \gamma')$
- (2) スタックトップの節点の状態と先読み語の品詞により、LR表から実行すべき操作を求めて実行する。操作にはシフト操作、レデュース操作、受理、エラーがある。受理なら解析成功として、またエラーなら解析失敗として終了。それ以外は(3)へ。
- (3) TRSSのスタックに対する操作が全てシフト操作なら、全てのシフト操作実行前と実行後にTRSSにマージ操作を施し(2)へ。さもなければ(4)へ。
- (4) レデュース操作を必要とする一つのスタック  $S$  に対して全てのレデュース操作を施し、その結果えられる新しいスタックと、 $S$  以外の既存のスタックとの間でマージを行い(3)へ。

### 3.6 なぜ逆ドット項か

レデュース操作を説明したので、YAGLR法で逆ドット項を作る意義が理解できる。実はTRSSを利用してレデュース操作時にアーリー項を作ることができるのであるが、ここで逆ドット項を作る理由を述べてみたい。今、以下の(r1)から(r3)に示すCFGを用いて、 $e \text{ You } 1 \text{ walk } 2 \text{ in } 3 \text{ the } 4 \text{ park } 5$  の解析を試みて見よう。

(r1)  $S \rightarrow NP VP$       (r2)  $VP \rightarrow V PP$       (r3)  $VP \rightarrow V$

解析が進み、“You walk”までをシフトした段階では、規則(r1)を用いたレデュース操作を行い、図3. 1に示すアーリー項を作ることができる。解析がさらに右に進み、“in the park”がPP(前置詞句)として受理され、規則(r2)により“walk in the park”がVPとして受理されたとしよう。すると再び規則(r1)を用いたレデュース操作を行い図3. 2に示すアーリー項を作ることができる。

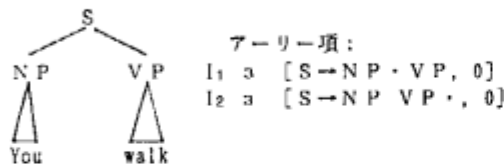


図3. 1 部分統語解析木とアーリー項(1)

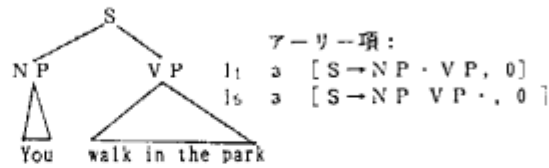


図3. 2 部分統語解析木とアーリー項(2)

ここで、項の集合 $i_1$ に所属させようとするアーリー項は、 $i_1$ の節点として既に登録済みのアーリー項と同じで重複することには注意したい。このことは、項の集合への節点の重複登録を避けるために、過去に作成したアーリー項の集合を調べる必要があることを意味している。従って過去に作成したアーリー項の全集合が、任意の時点で検索可能でなければならない。

一方逆ドット項を作成する場合には、逆ドット項中の位置番号は、レデュース操作時のスタックトップの節点のもつ位置番号である。この位置番号は、新たなシフト操作がなされ限り変わらないだけでなく、過去のシフト操作時の位置番号とも異なる。そのため逆ドット項の重複検査は、現在実行中の全てのレデュース操作により作られる項の集合を調べるだけでよい。この様にして、重複する逆ドット項を検査する範囲を局所化することが可能になり、容易にしかも高速な重複処理を行うことができる。新たなシフト操作を行う度に、重複検査のために、それまでに作られた逆ドット項を参照する必要はないのである。上記例の場合、逆ドット項なら以下が生成され、いずれも互いに重複しない。

図3. 1に対応した逆ドット項:  $I_1: S \rightarrow NP \cdot VP, 2$   
 $I_2: S \rightarrow \cdot NP VP, 2$

図3. 2に対応した逆ドット項:  $I_1: S \rightarrow NP \cdot VP, 5$   
 $I_5: S \rightarrow \cdot NP VP, 5$

これまでの説明から明らかなように、YAGLR法では、PP付加規則に対して、アーリー項なら重複した項生成が行われるが、逆ドット項なら重複した生成が行われない。一般に、重複する項の生成は、文法規則に依存するので、逆ドット項の生成が有利かアーリー項の生成が有利かは一概には決められない。しかし、逆ドット項は重複した項の検査を局所化することができることを説明した。これは、LR法が最右導出をベースとしていることと関連している。

### 3. 7 富田法とYAGLR法による統語解析例

本節では図3. 3に示すCFGを用いて、富田法とYAGLR法のそれぞれに対する統語解析例を示し、両者の比較を行う。解析対象文は「文化 が きた から 伝わった から...」であるとし、解析中途まで双方のトレースを行う。図3. 3のCFGから図3. 4に示すLR表が得られる。なお、富田法のスタックの構造は、[部分木番号, 状態]である。

- (1)  $S \rightarrow PP S$       (5)  $v \rightarrow きた$   
 (2)  $S \rightarrow v$       (6)  $v \rightarrow 伝わった$   
 (3)  $PP \rightarrow n p$       (7)  $n \rightarrow きた$   
 (4)  $PP \rightarrow S p$       (8)  $n \rightarrow 文化$   
                             (9)  $p \rightarrow から$   
                             (10)  $p \rightarrow が$

図3. 3 簡単な日本語CFG

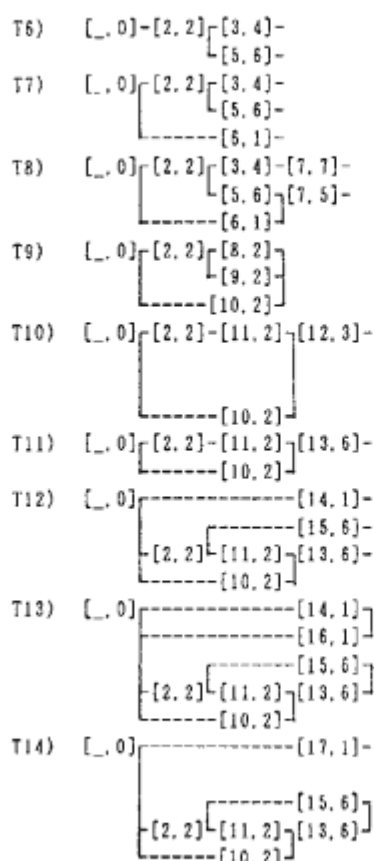
| 状態 | アクション部 |            |     |     | GOTO部 |   |
|----|--------|------------|-----|-----|-------|---|
|    | n      | p          | v   | S   | PP    | S |
| 0  | sh4    |            | sh3 |     | 2     | 1 |
| 1  |        | sh5        |     | 受理  |       |   |
| 2  | sh4    |            | sh3 |     | 2     | 6 |
| 3  |        | re2        |     | re2 |       |   |
| 4  |        | sh7        |     |     |       |   |
| 5  | re4    |            | re4 |     |       |   |
| 6  |        | sh5<br>re1 |     | re1 |       |   |
| 7  | re3    |            | re3 |     |       |   |

図3. 4 図3. 3のCFGから得られたLR表

#### [富田法による解析]

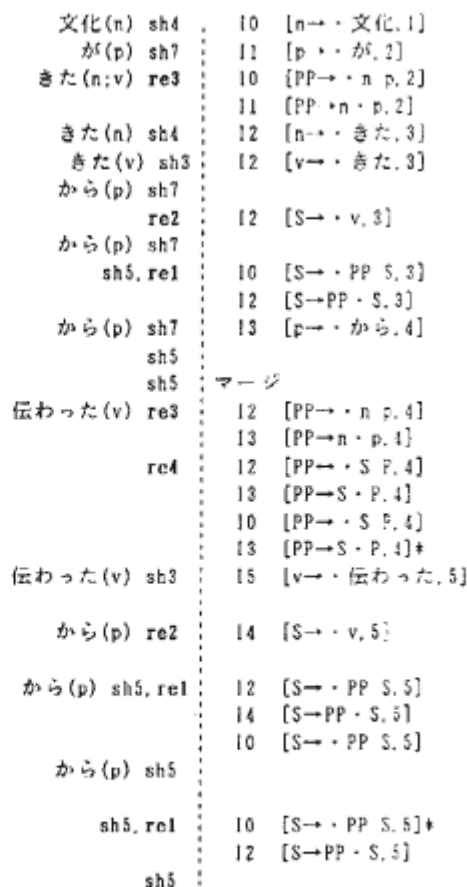
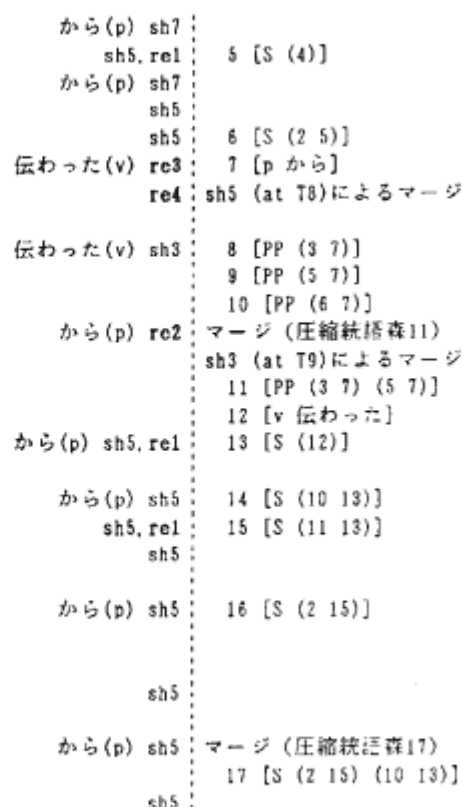
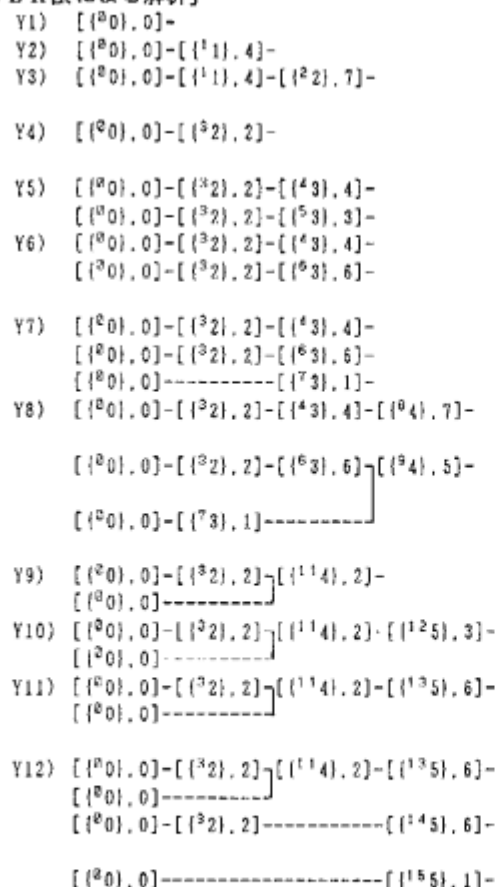
- T1)  $[_{0,0}]$ -  
 T2)  $[_{0,0}]-[0,4]$ -  
 T3)  $[_{0,0}]-[0,4]-[1,7]$ -  
 T4)  $[_{0,0}]-[2,2]$ -  
  
 T5)  $[_{0,0}]-[2,2]$   $[[3,4]$ -  
                              $[4,3]$ -

- 文化(n) sh4 :  
 が(p) sh7 : 0 [n 文化]  
 きた(n;v) re3 : 1 [p が]  
 きた(n) sh4 : 2 [PP (0 1)]  
 きた(v) sh3 :  
 から(p) sh7 : 3 [n きた]  
                             re2 : 4 [v きた]



(以下省略)

#### [YAGLR法による解析]





Y13) [(<sup>0</sup>0), 0] - [(<sup>2</sup>2), 2] - [(<sup>3</sup>2, <sup>1</sup>4), 2] - [(<sup>1</sup>5), 6] -  
 [(<sup>0</sup>0), 0] -----  
 [(<sup>0</sup>0), 0] ----- [(<sup>0</sup>5), 1] -  
 (以下省略)

から(p) sh5 マージ  
 sh5

注) \*付きの項は、重複項を表す、

ここで富田法とYAGLR法の比較をして見たい。T12) からT14) までとY11) からY13) までを比較してみよう。YAGLR法は、TRSSの構造が単純な分だけ、実装が容易である。また、YAGLR法は富田法以上にスタックのマージを進めることができるので、GSSを用いずとも計算効率と計算時に使用する記憶空間の縮小をはかることが可能になり、解析結果に含まれる曖昧性が増大するにつれて、YAGLR法と富田法の違いが際だってくる。4章ではその一端を説明する。

#### 4 解析対象文の長さに対する計算の複雑性

Johnsonらによれば、次の一連の文法Lnに対して、解析対象文の長さをn+2としたとき、解析に要する時間は $\Omega(n^2)$ である[Johnson 89][Kipps 89]。

- (1)  $S \rightarrow a$
- (2)  $S \rightarrow SS$
- (3)  $S \rightarrow S^{n+2}$

ただし $S^{n+2}$ は記号Sがn+2個連結したものの省略形である。解析対象文を $a^{n+2}S$ とする(ただし $n \geq m$ )。上記した規則に対して右のLR表が得られる。今、 $i (n+3 < i < n+2)$  番目のaをシフトしrelを施した後のTRSSの状態を図4. 1に示す。ここで $N_{p,q}$ は節点を表し、各節点の持つ識別子付き位置番号の集合間には、矢印で示す包含関係がある。節点XとYの間に $X \rightarrow Y$ があれば、Xの識別子付き位置番号の集合がYの部分集合であることを示す。

容易に示すことができるように、節点 $N_{p,q}$ の持つ状態 $T_{p,q}$ は、

$$T_{p,q} = \begin{cases} q & \text{if } 1 \leq q \leq m+3, \\ m+3 & \text{if } m+3 < q, \\ 0 & \text{if } q = 1. \end{cases}$$

図4. 1に対して、次のシフト操作を施すまで、レデュース操作を次々に施す必要があるが、そのための計算コストを考えてみよう。まず、規則(2)を適用してレデュース操作を進める場合の方が、規則(3)を適用する場合に比べてコストがかかることは明らかである。したがって、計算コストの

| 状態   | アクション部      |         | GOTO部 |
|------|-------------|---------|-------|
|      | a           | S       |       |
| 0    | sh1         |         | 2     |
| 1    | rel         | rel     |       |
| 2    | sh1         | acc     | 3     |
| 3    | sh1/re2     | re2     | 4     |
| 4    | sh1/re2     | re2     | 5     |
| .... | ....        | ....    | ....  |
| m+1  | sh1/re2     | re2     | m+2   |
| m+2  | sh1/re2     | re2     | m+3   |
| m+3  | sh1/re2/re3 | re2/re3 | m+3   |

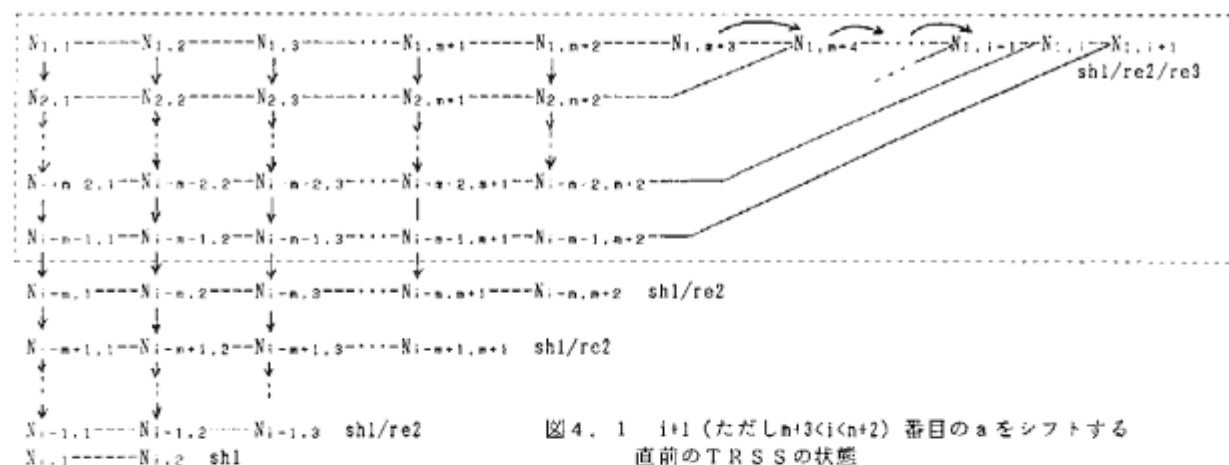
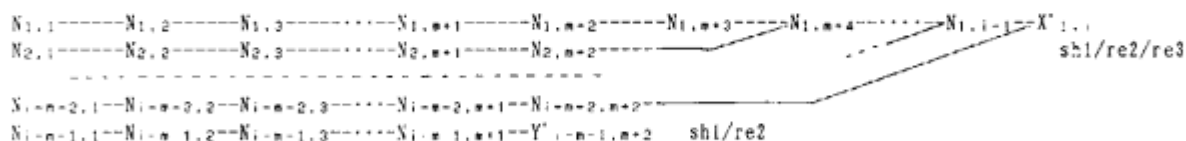


図4. 1  $i+1$  (ただし $n+3 < i < n+2$ ) 番目のaをシフトする直前のTRSSの状態

オーダを求めるためには、規則(2)を適用するレデュースを考えれば十分である。また、最上部の点線で囲まれたTRSSの部分のみが、入力文の長さに依存しているので、これに対するレデュース操作の計算コストを考えれば十分である。今、図4. 1の最上部のスタックに対して、規則(2)によるレデュース操作を一度施すと、下図の二つのスタック構造が得られる。



上図のスタックは、図4. 1の統合化可能なスタックとマージされる。たとえば、sh1を待つ最長スタック同士のスタックトップの状態はn+3で等しいのでマージされる。このとき、図4. 1に記述した節点間の識別子付き位置番号の集合間、包含関係が成り立つため、マージ操作は先頭から二つの目までの節点同士( $\langle N_{1,i-1} \text{ と } X'_{1,i} \rangle$ ,  $\langle N_{1,i} \text{ と } N_{i-1,i} \rangle$ ,  $\langle N_{i-n-1,n+2} \text{ と } N_{i-n-2,n+2} \rangle$ )で行えばよい。先頭から3番目の節点同士( $\langle N_{1,i-1} \text{ と } N_{1,i-2} \rangle$  または  $\langle N_{i-n,n+1} \text{ と } N_{i-n-2}$

$\langle n-1 \rangle$  のマージなどは、識別子付き位置番号の集合間に部分集合関係が成立しているので、それ以上マージを葉の方向に向けて進める必要はない。

以上の考察から、re2操作後のマージには先頭から3番目のマージまでを考慮すればよいことが分かる。しかも、それぞれのマージには、識別子付き位置番号の集合の要素が高々  $i+1$  含まれるだけであるから、マージに要する計算コストは  $i$  のオーダーである。一方、re2は繰り返し行われるが、その回数は高々  $i$  回である。したがって、レデュース操作とマージ操作の総和は、 $\Omega(i^2)$  である。以上より、 $a^{n^2}$  の入力文を解析するために要する総計算コストのオーダーは、 $\sum_{i=1}^n \Omega(i^2) = \Omega(n^3)$  となり、証明された。

次に図4、1の各節点中の識別子付き位置番号の数は：節点  $N_{i,j}$  に対して、 $2 \leq j \leq m+3$  の場合1、 $n+4 \leq j \leq i-1$  の場合2、 $i \leq j \leq i+1$  の場合1、節点  $N_{k,j}$  ( $2 \leq k \leq i-m-1$ ,  $2 \leq j \leq m+2$ ) に対して  $k$  である。したがって、点線で囲まれた部分の節点中の識別子付き位置番号の総和は、 $1 \cdot (n+4) + 2 \cdot (i-m-4) + (n+2) \cdot (2+3+\dots+(i-m-1))$  である。この総和のオーダーは  $i^2$  である。これは、点線で囲まれた部分以外の節点に含まれる識別子付き位置番号の総和のオーダーと等しい。以上から、入力文  $a^{n^2}$  を解析するために必要な TRSS の空間のオーダーは  $n^2$  であり、また、作成される逆ドット項の数の総和のオーダーは、 $n^3$  であることが証明された。

## 5 おわりに

本論文ではYAGLR法と呼ぶ新しい統語解析アルゴリズムを提案した。YAGLR法は、既存の方法と比べていくつかの利点を持っている。

- (A) Johnsonらの与えた特殊な文法に対しても、解析時間のオーダーは、解析対象文の長さ  $n$  に対してアーリー法と同等の  $n^3$  のオーダーに抑えることができる。
  - (B) レデュース操作時に、アーリー項と対称な逆ドット項を作成しながら統語解析を行う。逆ドット項を生成する利点について考察した。
  - (C) 最終情報を用いているので、最終的に生成される逆ドット項の数はアーリー法より少ない。
  - (D) 富田法以上にスタックのマージを行うことができるため、グラフ構造化スタックより構造が単純で操作し易い木構造化スタックを用いることができる。それによりスタック操作が単純になる。
  - (E) 解析終了後に、作成した逆ドット項を組み合わせて統語解析木を作る。
- 今後の検討課題として、以下のものを挙げるができる。
- (1) YAGLR法の統語解析時間のオーダーが、一般のCFGに対して  $n^3$  のオーダーであることの証明。
  - (2) YAGLR法による統語解析に要する空間量の評価。
  - (3) YAGLR法の並列解析アルゴリズムの研究。
  - (4) 逆ドット項からの統語構造の並列抽出アルゴリズムの研究。

(1) は重要な研究課題であるが、(2) に述べたように、解析に要する空間の大きさの評価も行う必要がある。Kippsは富田法に比較的単純な修正を施すことで、一般のCFGに対する計算時間のオーダーを  $\Omega(n^3)$  に抑えることができることを示した[Kipps 89]。しかし、このアルゴリズムは富田法以上の空間を要す。これに対してYAGLR法では、スタックのマージを徹底的に行うアルゴリズムであるから、計算空間を抑えることができる。一般のCFGに対するYAGLR法の使用記憶空間の大きさの詳細な検討を行う必要がある。

(3) に述べたように、富田法と同様、YAGLR法の並列解析アルゴリズムを検討することも意味があるだろう[Tanaka 89]。本稿で述べたアルゴリズムは橋型探索の逐次型アルゴリズムであり、できるだけ無駄な解析を省いたアルゴリズムである。この長所を生かしつつ、シフトレデュース・コンフリクトが発生したときの待ち合わせができるだけ少なくなる並列解析アルゴリズムを研究する必要がある。それには沼崎らの研究が参考になる[Numazaki 90]。

YAGLR法では、統語解析後に得られた逆ドット項の集合から統語解析木を作り出さねばならない。これはアーリー法と双対なアルゴリズムを用いれば良い。一つの統語解析木を計算するための時間は  $n^2$  のオーダーであることが知られている[Aho 72]。この統語解析木の計算にも多くの並列性が観察されるので、統語解析木抽出用の並列アルゴリズムについて今後研究する必要がある。

## 参考文献

- [Aho 72] Aho, A.V. and Ullman, J.D.: The Theory of Parsing, Translation and Compiling. Prentice-Hall, New Jersey(1972).
- [Earley 70] Earley, J.: An Efficient Augmented-Context-Free Parsing Algorithm. Comm. of ACM, 13, 1-2, 95-102(1970).
- [Johnson 89] Johnson, M.: The Computational Complexity of Tomita's Algorithm. International Parsing Workshop '89, Carnegie-Mellon University, 203-208(1989).
- [Kipps 89] Kipps, J. R.: Analysis of Tomita's Algorithm for General Context-Free Parsing. International Parsing Workshop '89, Carnegie-Mellon University, 193-202(1989).
- [Numazaki 90] Numazaki, H. and Tanaka, H.: A new Parallel Algorithm for Generalized LR Parsing. COLING'90 305-310(1990).
- [Suresh 90] Suresh, K.G.: Yet Another Generalized LR Parser-Report4. Tanaka Lab., Tokyo Inst. of Technology (1990.4.16)
- [田中 89] 田中穂積：自然言語解析の基礎、産業図書(1989)。
- [Tanaka 89] Tanaka, H. and Numazaki, H.: Parallel Generalized LR Parser Based on Logic Programming. 1st Australian-Japan Joint Symposium on Natural Language processing, 201-211(1989).
- [Tomita 86] Tomita, M.: Efficient Parsing for Natural Language. Kluwer, Boston, Mass(1986).
- [Tomita 87] Tomita, M.: An Efficient Augmented-Context-Free Parsing Algorithm. Computational Linguistics, 13, 31-46(1987).