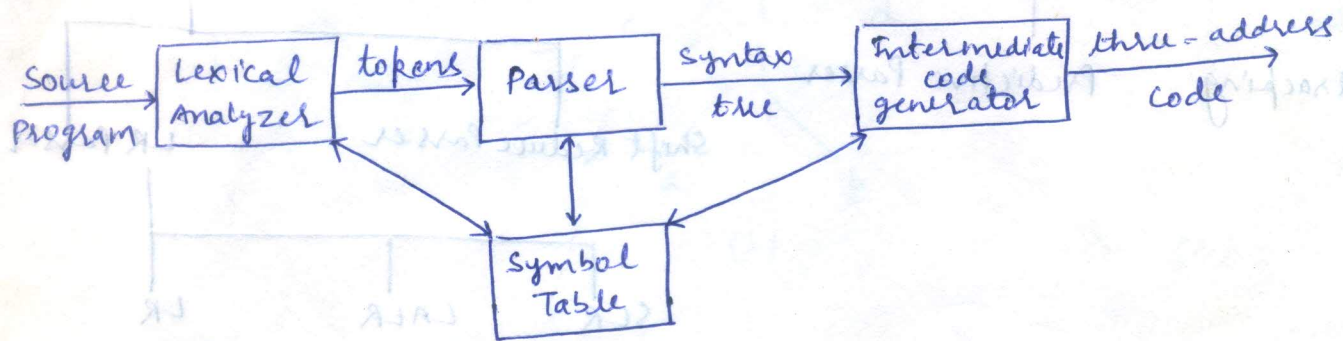


Parser : →

- The component that performs the syntax analysis is called syntax analyzer or parser.

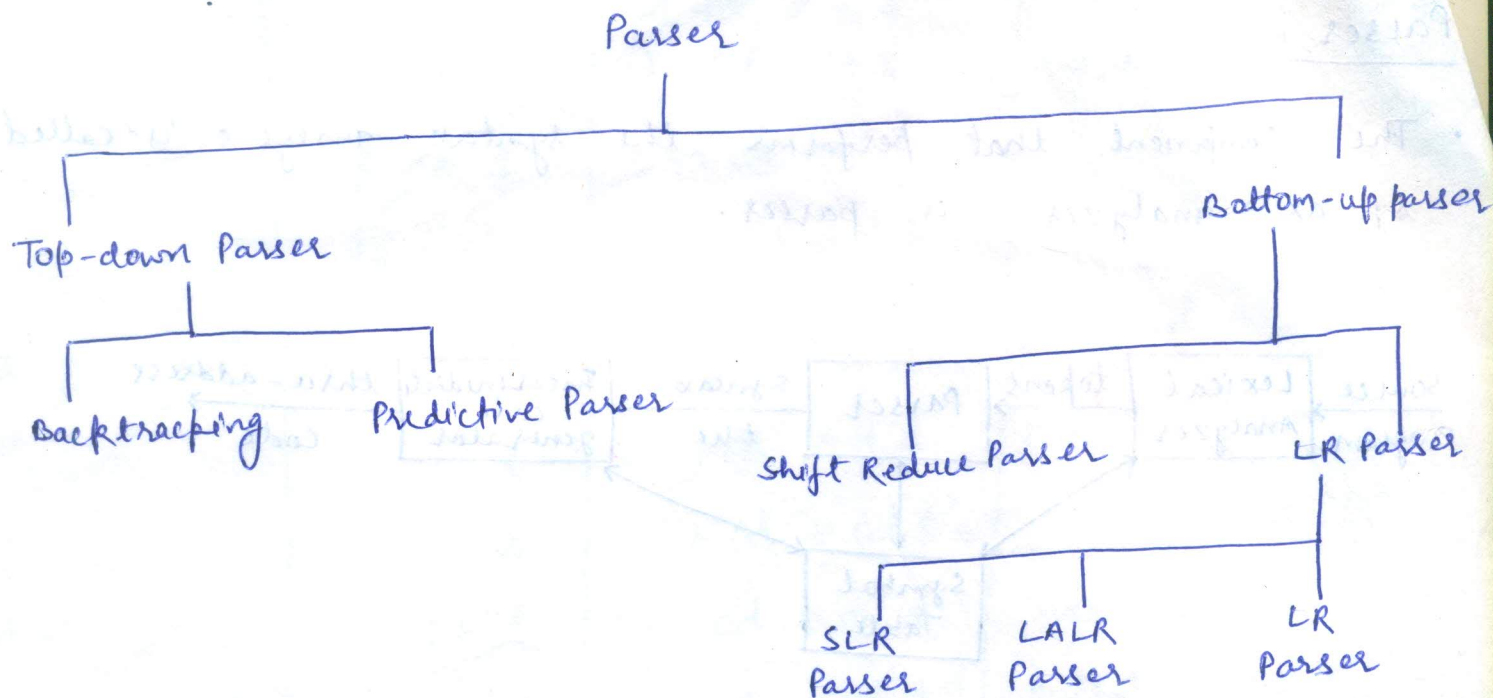


A model of a compiler front end

The main functions and tasks performed by a parser are -

- (i) A parser obtains a sequence of tokens from the lexical analyzer.
- (ii) checks that the tokens appearing in its input occur in patterns permitted by the specification for the source language.
- (iii) A parser verifies that the string can be generated by the grammar for the source language.
- (iv) A parser imposes a hierarchical tree like structure (Parse tree) on the incoming token stream from which the intermediate code can be generated.
- (v) A parser detects and reports any syntax error.
- (vi) collects all the information in the symbol table.
- (vii) Passes its output to the intermediate code generator for further processing.

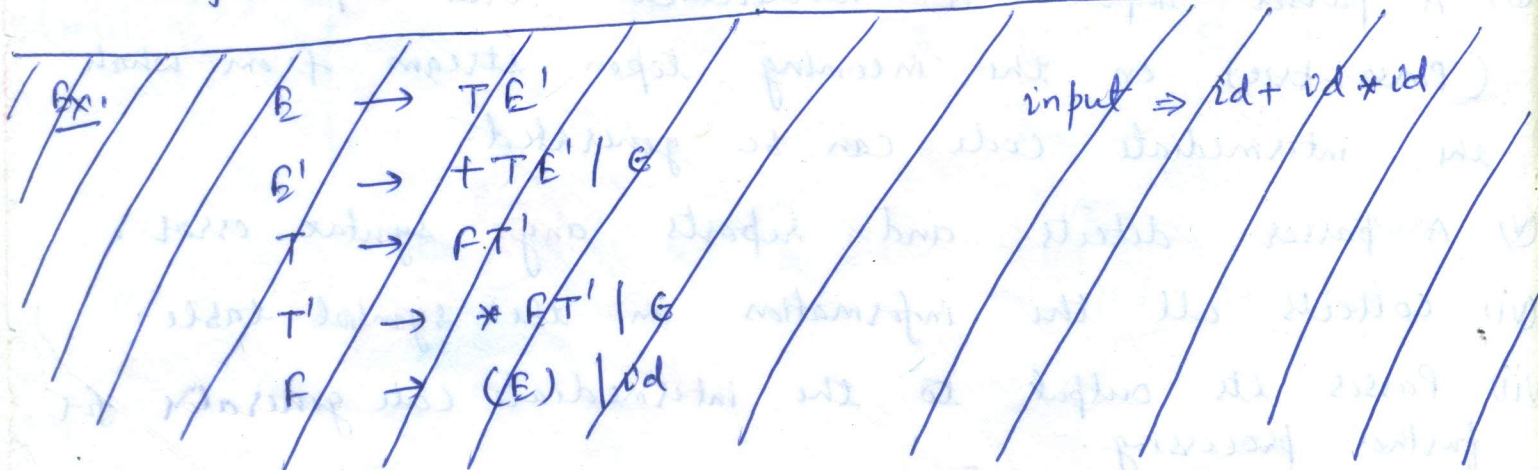
Types of parser



Parsing Techniques

Top - Down Parsing

- Top down parsing can be viewed as the problem of constructing a parse tree for the input string, starting from the root and creating the nodes of the parse tree in preorder.
- Equivalently, top down parsing can be viewed as finding a leftmost derivation for an input string.

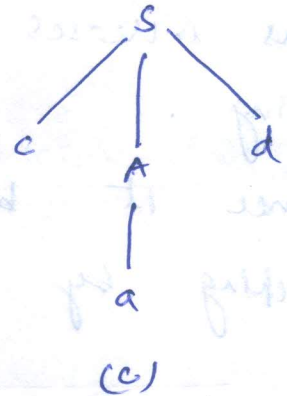
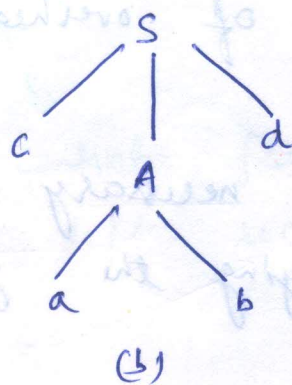
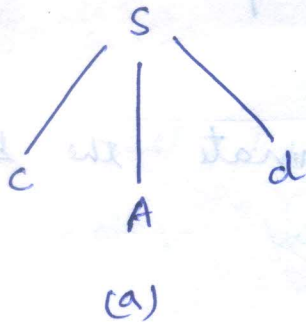


Consider the grammar

$$S \rightarrow cAd$$

$$A \rightarrow ab|a$$

and the input string $w = cad$.



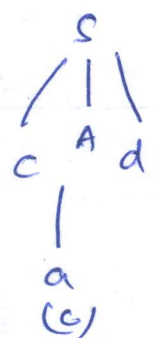
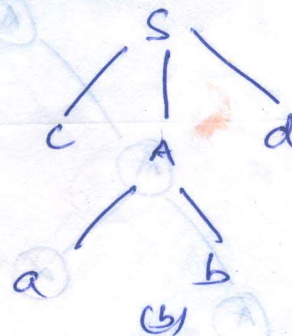
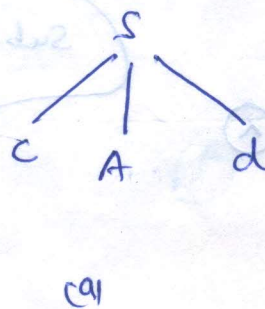
Problems with Top-down Parsing -

- In order to implement the parsing we need to eliminate following problems.

(a) Backtracking :-

- Backtracking is a technique in which for expansion of non terminal symbol we choose one alternative and if some mismatch occurs then we try another alternative if any -

ex. $S \rightarrow cAd$
 $A \rightarrow ab|a$



- If for a non-terminal there are multiple production rules beginning with the same input symbol then

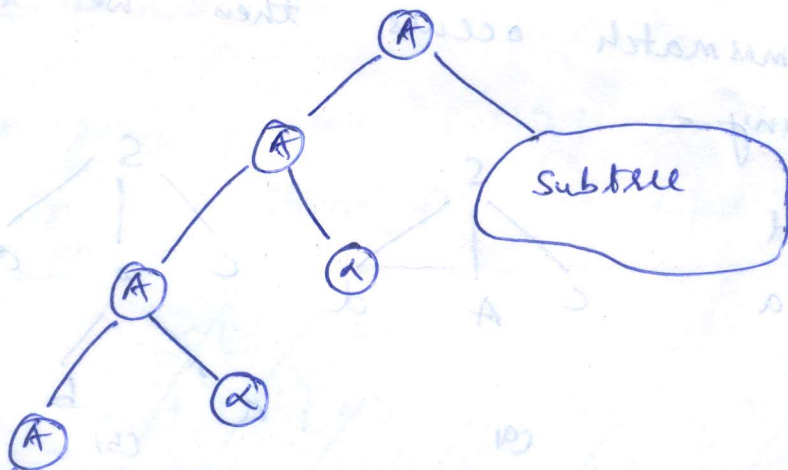
to get the correct derivation we need to try all these alternatives.

- secondly, in backtracking we need to move some levels upward in order to check the possibilities.
- This increases lot of overhead in implementation of parsing.
- Hence it becomes necessary to eliminate the backtracking by modifying the grammar.

(b) Left Recursion :-

$$A \xRightarrow{+} A \propto$$

- Here $\xRightarrow{+}$ means deriving the input in one or more steps.
- The A can be non terminal and a denotes some input string.
- Because of left recursion, the top down parser can enter in infinite loop.



left recursion

Thus expansion of A causes further expansion of A only and due to generation of $A, Ax, Axx, \dots$ the input pointer will not be advanced.

- To eliminate left recursion we need to modify the grammar.

→ Let G be a context free grammar having a production rule with left recursion.

$$A \rightarrow Ax$$

①

$$A \rightarrow B$$

Then we eliminate left recursion by re-writing the production rule as -

$$A \rightarrow BA'$$

$$A' \rightarrow xA'$$

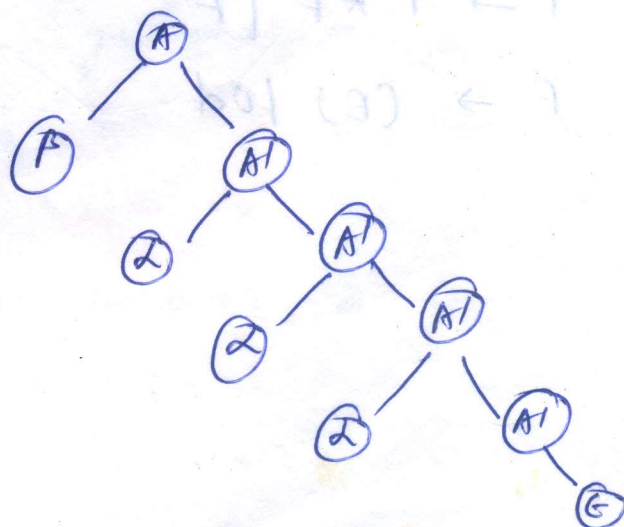
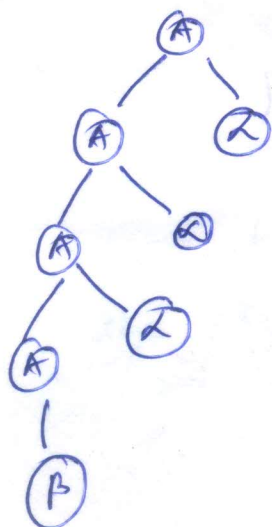
$$A' \rightarrow \epsilon$$

②

Thus a new symbol A' is introduced. We can also verify, whether the modified grammar is equivalent to original or not.

B	x	x	x
---	---	---	---

input string



for ex: Consider the grammar

$$E \rightarrow E + T \mid T$$

we can map this grammar with the rule $A \rightarrow A \prec / \beta$.

Then using equation (2) we can say,

$$A = E$$

$$\lambda = +T$$

$$\beta = T$$

then the rule becomes,

$$E \rightarrow TE'$$

$$E' \rightarrow +TE' \mid \epsilon$$

similarly for the rule,

$$T \rightarrow T * F \mid F$$

we can eliminate left recursion as

$$T \rightarrow FT'$$

$$T' \rightarrow *FT' \mid \epsilon$$

The grammar for arithmetic expression can be equivalently written as -

$$E \rightarrow E + T \mid T$$

$$E \rightarrow TE'$$

$$E' \rightarrow +TE' \mid \epsilon$$

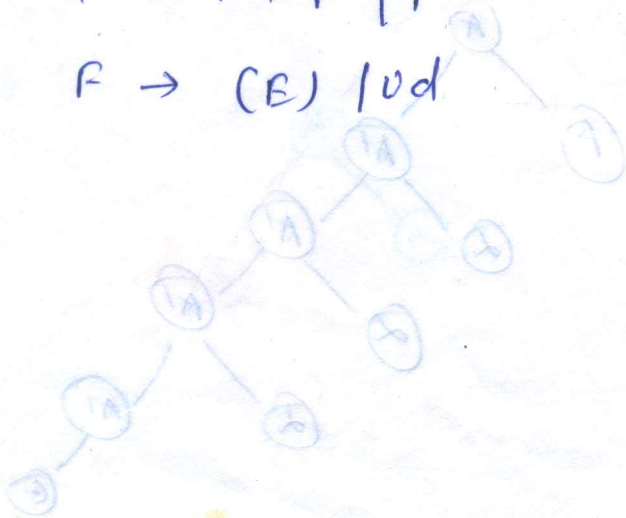
$$T \rightarrow FT'$$

$$T \rightarrow T * F \mid F$$

$$\Rightarrow T' \rightarrow *FT' \mid \epsilon$$

$$F \rightarrow (E) \mid \text{id}$$

$$\Rightarrow F \rightarrow (E) \mid \text{id}$$



1) Left factoring : $\rightarrow$

- Left factoring 'is used when it is not clear that which of the two alternatives is used to expand the non terminal.
- By left factoring we may be able to re-write the production in which the decision can be deferred until enough of the input is seen to make the right choice.

In general if

$A \rightarrow \alpha\beta_1 \mid \alpha\beta_2$ is a production, then it is not possible for us to take a decision whether to choose first rule or second.

- In such situation the above grammar can be left factored as

$$A \rightarrow \alpha A'$$

$$A' \rightarrow \beta_1 \mid \beta_2$$

For ex. consider the following grammar

$$S \rightarrow cEs \mid cEsEs \mid a$$

$$E \rightarrow b$$

the left factored grammar becomes,

$$S \rightarrow cEsS' \mid a$$

$$S' \rightarrow Es \mid \epsilon$$

$$E \rightarrow b$$

Ex - Do left factoring in the following grammar -

$$A \rightarrow aAB | aA | a$$

$$B \rightarrow bB | b$$

If the rule is $A \rightarrow \alpha\beta_1 | \alpha\beta_2 | \dots$ is a production then the grammar needs to be left factored.

$$A \rightarrow aAB | aA | a$$

$\downarrow \quad \downarrow \quad \downarrow \quad \downarrow \quad \downarrow \quad \downarrow$
 $\alpha \quad \beta_1 \quad \alpha \quad \beta_2 \quad \alpha \quad \beta_3$

We have to convert it to

$$A \rightarrow \alpha A'$$

$$A \rightarrow aA'$$

$$A' \rightarrow \beta_1 | \beta_2 \Rightarrow$$

$$A' \rightarrow AB | A | \epsilon$$

Similarly

$$B \rightarrow bB | b$$

$$\downarrow \quad \downarrow \quad \downarrow \quad \downarrow$$

$\alpha \quad \beta_1 \quad \alpha \quad \beta_2$

$$\Rightarrow B \rightarrow bB' \\ B' \rightarrow B | \epsilon$$

To summarize, the grammar with left factor operation will be -

$$A \rightarrow aA'$$

$$A' \rightarrow AB | A | \epsilon$$

$$B \rightarrow ~~B~~ bB'$$

$$B' \rightarrow B | \epsilon$$

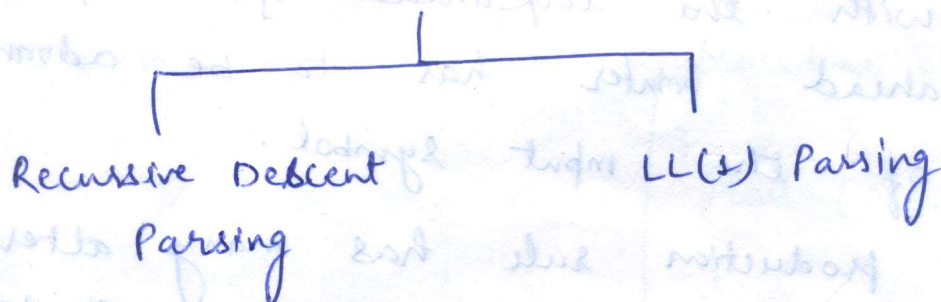
Ambiguity -

Same as in unit - 1.

A backtracking parser will try different production rules to find the match for the input string by backtracking each time.

- more powerful than predictive parsing but slower than ~~but~~ predictive parsing and is not preferred for practical compilers.

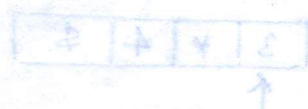
Predictive Parsers



The predictive parser tries to predict the next construction using one or more lookahead symbols from input string.

(1) Recursive Descent Parsing $\Rightarrow$

- A parser that uses collection of recursive procedures for parsing the given input string is called Recursive Descent (RD) Parser.



• The R.H.S. of the production rule is directly converted to a program. (66)

• For each non-terminal a separate procedure is written and body of the procedure (code) is R.H.S. of the corresponding non-terminal.

Basic steps for construction of RD parser :-

The R.H.S. of the rule is directly converted into program code symbol by symbol -

1. If the input symbol is non-terminal then a call to the procedure corresponding to the non-terminal is made.

2. If the input symbol is terminal then it is matched with the lookahead from input.

The lookahead pointer has to be advanced on matching of the input symbol.

3. If the production rule has many alternatives then all these alternatives have to be combined into a single body of procedure.

4. The parser should be activated by a procedure corresponding to the start symbol.

Ex - $E \rightarrow \text{num } T$

$T \rightarrow * \text{ num } T \mid \epsilon$

3	*	4	\$
---	---	---	----

↑

$E \rightarrow \text{num } T$

3	*	4	\$
---	---	---	----

↑

$T \rightarrow * \text{ num } T$

*	4	\$
---	---	----

↑

 $T \rightarrow * \text{num } T ;$

3	*	4	\$
---	---	---	----

↑

declare
success

Pseudo code for recursive descent parser of above example -

Procedure E

```

{
  if lookahead = num then
  {
    match(num);
    T; /* call to procedure T
  }
else error;
if lookahead = $
{
  declare success; /* Return on success
}
else
  error;
} /* end of procedure E

```

Procedure T

```

{
  if lookahead = '*'
  {
    match('*');
  }
  if lookahead = 'num'
  {
    match(num);
    T;
  }
  else error
}

```

else NULL /* indicates ^{alternate} error
the other is combined into same
procedure T

procedure match (token)

```

{
  if lookahead = token
  {
    lookahead = next_token;
  }
  else
    error;
}

```

procedure error

```

{
  print ("Error");
}

```

Generalised Algo.

```

void A {
  choose an A production,  $A \rightarrow x_1 x_2 \dots x_k$ 
  for ( $i=1$  to  $k$ ) {
    if ( $x_i$  is a nonterminal)
      call procedure  $X_i(i)$ ;
    elseif ( $x_i$  equals the current
      input symbol  $a$ )
      advance the i/p to the next symbol;
    else error
  }
}

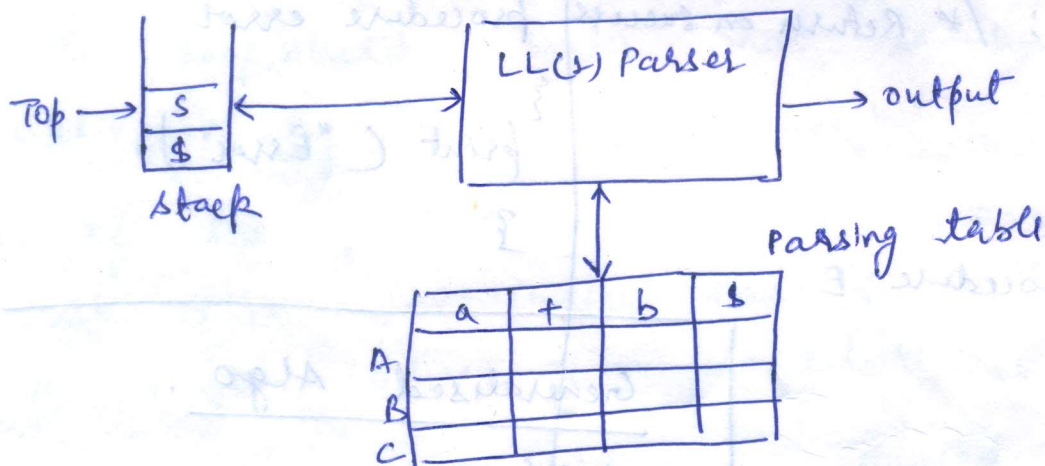
```

Predictive LL(1) Parser

- This top-down parsing algorithm is of non-recursive type and in this a table is built.
- The first L means the input is scanned from left to right.
- The second L means it uses leftmost derivation for input string.
- The number 1 in the input symbol means it uses only one input symbol (lookahead) to predict the parsing process.

input tokens

a	+	b	\$
---	---	---	----



model for LL(1) parser

- The data structures used by LL(1) are i) input buffer ii) stack iii) parsing table.
- LL(1) parser uses i/p buffer to store the input tokens.

The stack is used to hold the left sentential form. (69)

The symbols on R.H.S. of rule are pushed into the stack in reverse order i.e. from right to left.

- Thus use of stack makes this algorithm non-recursive.
- The table is basically a two dimensional array.
- The table has row for non terminal and column for terminals.
- The table can be represented as $M[A, a]$ where A is a non terminal and a is current input symbol.
- The parser consults the table $M[A, a]$ each time while taking the parsing actions.
- The configuration of LL(1) parser can be defined by top of the stack and a lookahead token.
- One by one configuration is performed and the input is successfully parsed if the parser reaches the halting configuration.
- When the stack is empty and next token is ϕ then it corresponds to successful parse.

FIRST and FOLLOW Function

The construction of both top down and bottom up parser is aided by two functions, FIRST and FOLLOW.

FIRST Function -

- $FIRST(\alpha)$ is a set of terminal symbols that are first symbols appearing at R.H.S. in derivation of α .
If $\alpha \Rightarrow \epsilon$ then ϵ is also in $FIRST(\alpha)$.

Following are the rules used to compute the FIRST functions -

1. If the terminal symbol a the $FIRST(a) = \{a\}$.
2. If there is a rule $X \rightarrow \epsilon$ then $FIRST(X) = \{\epsilon\}$.
3. For the rule $A \rightarrow X_1 X_2 \dots X_k$ $FIRST(A) = FIRST(X_1) \cup FIRST(X_2) \dots FIRST(X_k)$.

where k $X_j \in n$ such that $1 \leq j \leq k-1$.

FOLLOW function :-

FOLLOW(A) is defined as the set of terminal symbols that appear immediately to the right of A.

Alternatively $FOLLOW(A) = \{a | S \xRightarrow{*} \alpha A a \beta \text{ where } \alpha \text{ and } \beta \text{ are some grammar symbols may be terminal or non terminal.}\}$

Rules - 1. For the start symbol S place $\$$ in $FOLLOW(S)$.

If there is a production $A \rightarrow \alpha B \beta$ then everything in $FIRST(A)$ without ϵ is to be placed in $FOLLOW(B)$. (71)

3. If there is a production $A \rightarrow \alpha B \beta$ or $A \rightarrow \alpha B$ and $FIRST(B) = \{\epsilon\}$ then $FOLLOW(A) = FOLLOW(B)$ or $FOLLOW(B) = FOLLOW(A)$.

That means everything in $FOLLOW(A)$ is in $FOLLOW(B)$.

Example -

$$E \rightarrow TE'$$

$$E' \rightarrow +TE' \mid \epsilon$$

$$T \rightarrow FT'$$

$$T' \rightarrow *FT' \mid \epsilon$$

$$F \rightarrow (E) \mid \text{id}$$

$$1. FIRST(E) = FIRST(T) = FIRST(E) = \{ (, \text{id} \}$$

• Because the two productions for F have bodies that start with these two terminal symbols, id and the left parentheses.

• T has only one production, and its body starts with F .

• Since F does not derive ϵ , $FIRST(T)$ must be the same as $FIRST(F)$.

2. $FIRST(E') = \{+, \epsilon\}$

Because one of the two productions for E' has a body that begins with terminal $+$, and the other body is ϵ .

- Whenever a nonterminal derives ϵ , we place ϵ in $FIRST$ for that nonterminal.

3. $FIRST(T') = \{*, \epsilon\}$ because analogous to that for $FIRST(E')$.

4. $FOLLOW(E) = FOLLOW(E') = \{), \$\}$.

- Since E is the start symbol, $FOLLOW(E)$ must contain $\$$.

- The production body (E) explains why the right parenthesis is in $FOLLOW(E)$.

- For E' , this nonterminal appears only at the ends of bodies of E -productions. Thus $FOLLOW(E')$ must be the same as $FOLLOW(E)$.

5. $FOLLOW(T) = FOLLOW(T') = \{+,), \$\}$

Because T appears in bodies only followed by E' .

- Thus everything except ϵ that is in $FIRST(E')$ must be in $FOLLOW(T)$; that explains the symbol $+$.
- However, since $FIRST(E')$ contains ϵ (i.e. $E' \xRightarrow{*} \epsilon$) and E' is the entire string following T in the bodies of the E -productions, everything in $FOLLOW(E)$ must also be in $FOLLOW(T)$.

that explains the symbols \$ and the right parentheses. (73)

As for T' , since it appears only at the ends of the T -productions, it must be that $\text{FOLLOW}(T')$ = $\text{FOLLOW}(T)$.

6. $\text{FOLLOW}(F) = \{+, *,), \$\}$ because ~~is~~ analogous to that for T in point (5).

Symbols	FIRST	FOLLOW
E	$\{C, id\}$	$\{), \$\}$
E'	$\{+, \epsilon\}$	$\{), \$\}$
T	$\{C, id\}$	$\{+,), \$\}$
T'	$\{*, \epsilon\}$	$\{+,), \$\}$
F	$\{C, id\}$	$\{+, *,), \$\}$

Construction of a predictive parsing table

For each production $A \rightarrow \alpha$ of the grammar, do the following -

- For each terminal a on $\text{FIRST}(\alpha)$, add $A \rightarrow \alpha$ to $M[A, a]$.
- If ϵ is on $\text{FIRST}(\alpha)$, then for each terminal b in $\text{FOLLOW}(A)$, add $A \rightarrow \alpha$ to $M[A, b]$.
- If ϵ is on $\text{FIRST}(\alpha)$ and $\$$ is in $\text{FOLLOW}(A)$ add $A \rightarrow \alpha$ to $M[A, \$]$ as well.

If after performing the above, there is no production at all in $M[A, a]$, then set $M[A, a]$ to error (empty entry in the table).

Ex - $E \rightarrow TE'$

$E' \rightarrow +TE' \mid \epsilon$

$T \rightarrow FT'$

$T' \rightarrow *FT' \mid \epsilon$

$F \Rightarrow (E) \mid \text{id}$

Non Terminal	Input Symbol					
	id	+	*	(	)	\$
E	$E \rightarrow TE'$			$E \rightarrow TE'$		
E'		$E' \rightarrow +TE'$			$E' \rightarrow \epsilon$	$E' \rightarrow \epsilon$
T	$T \rightarrow FT'$			$T \rightarrow FT'$		
T'		$T' \rightarrow \epsilon$	$T' \rightarrow *FT'$		$T' \rightarrow \epsilon$	$T' \rightarrow \epsilon$
F	$F \rightarrow \text{id}$			$F \rightarrow (E)$		

consider production $E \rightarrow TE'$

since $\text{FIRST}(TE') = \text{FIRST}(T) = \{\epsilon, \text{id}\}$

this production is added to $M[E, (]$ and $M[E, \text{id}]$

Production $E' \rightarrow +TE'$ is added to $M[E', +]$ since

$\text{FIRST}(+TE') = \{+\}$

since $\text{FOLLOW}(E') = \{), \$\}$, production $E' \rightarrow \epsilon$ is added

to $M[E',)]$ and $M[E', \$]$.

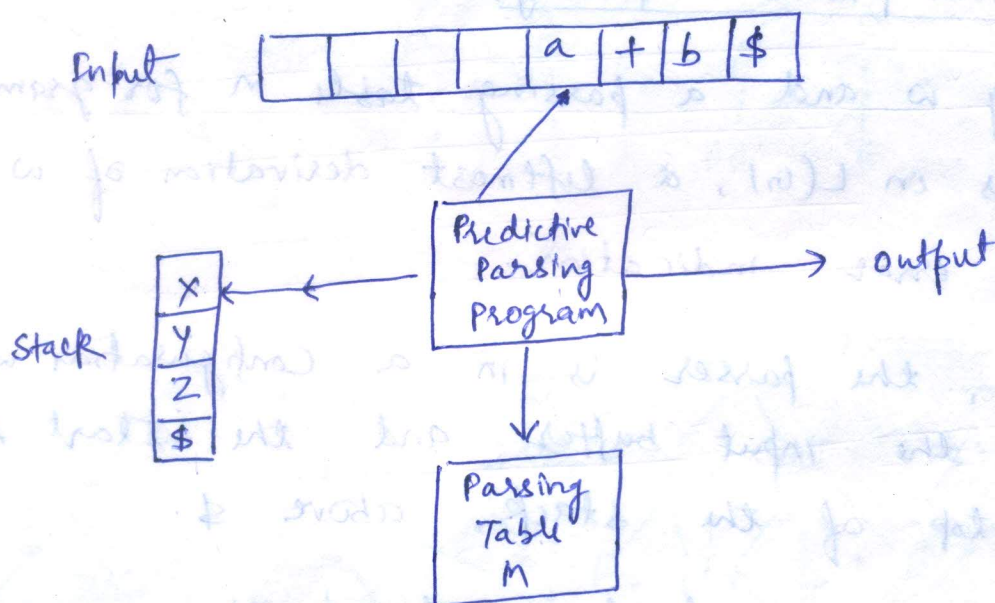
Non recursive Predictive Parsing

(75)

A non recursive predictive parser can be built by maintaining a stack explicitly, rather than implicitly via recursive calls.

- If w is the input that has been matched so far, then the stack holds a sequence of grammar symbols α such that

$$S \xrightarrow[\text{lm}]{*} w\alpha$$



Model of a table-driven predictive parser

- Parsing table is constructed by previous algo. and according to the above diagram, the input buffer contains the string to be parsed, followed by the endmarker \$.
- We reuse the symbol \$ to mark the bottom of the stack, which initially contains the start symbol of the grammar on top of \$.

Example -

$$E \rightarrow TB'$$

$$E' \rightarrow +TB' \mid \epsilon$$

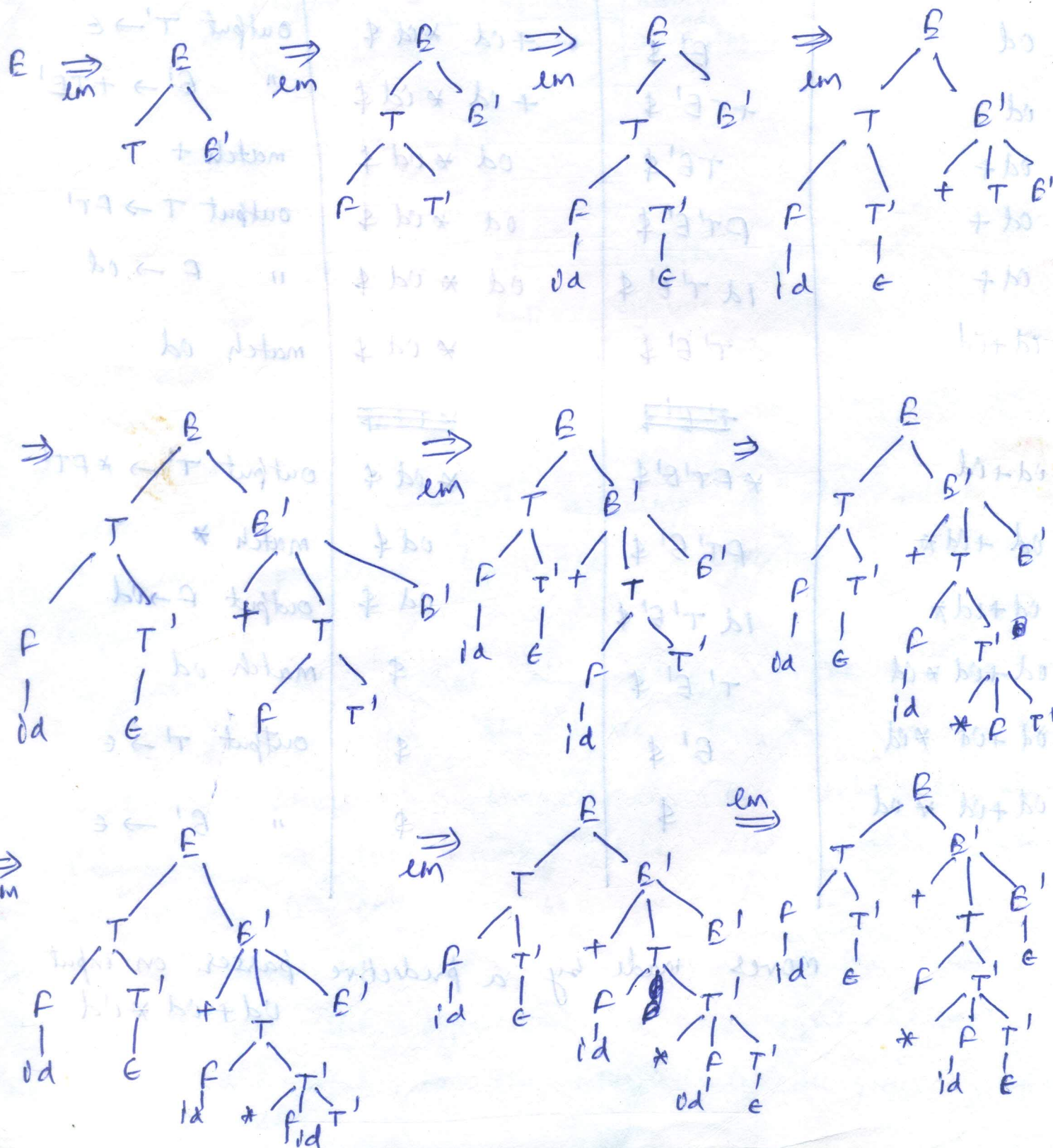
$$T \rightarrow FT' \mid \epsilon$$

$$T' \rightarrow *FT' \mid \epsilon$$

$$F \rightarrow (E) \mid id$$

Apply non recursive predictive algo. on input $id+id*id$

Derivation



matched	stack	Input	Action
	E \$	od + od * od \$	output E → TE'
	T B' \$	od + od * od \$	" T → FT'
	F T B' \$	od + od * od \$	" F → od
	od T' E' \$	id + od * od \$	matched od
od	T' E' \$	+ od * od \$	output T' → E
od	E' \$	+ od * od \$	" E' → + TE'
od	+ TE' \$	+ od * od \$	match +
od +	TE' \$	od * od \$	output T → FT'
od +	FT' E' \$	od * od \$	" F → od
od +	id T' E' \$	od * od \$	match od
od + od	T' E' \$	* od \$	output T' → * FT'
od + od	TE' \$	od \$	match *
od + od *	* FT' E' \$	od \$	output F → od
od + od *	FT' B' \$	od \$	match od
od + od *	id T' E' \$	\$	output T' → E
od + od * od	T' E' \$	\$	" E' → E
od + od * od	B' \$	\$	
od + od * od	\$	\$	

moves made by a predictive parser on input
 od + od * od

Bottom-Up Parsing

(79)

- A bottom-up parse corresponds to the construction of a parse tree for an input string beginning at the leaves (bottom) and working up towards the root (top).

Ex- $E \rightarrow E + T \mid T$

$$T \rightarrow T * F \mid F$$

$$F \rightarrow (E) \mid id$$

Derive parse tree for the token stream $id * id$.

$id * id$

$F * id$

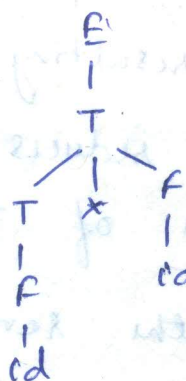
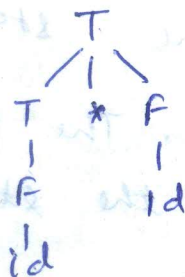
$T * id$

$T * F$

$\mid$
 id

$\mid$
 F
 $\mid$
 id

$\mid$
 F
 $\mid$
 id



A bottom-up parse for $id * id$

Reductions :-

- Bottom-up parsing is a process of reducing a string w to the start symbol of the grammar.
- At each reduction step, a specific substring matching the body of a production is replaced by the non-terminal at the head of that production.

According to the previous example, reduction is performed (80)
as below -

$vd * vd, F * vd, T * vd, T * F, T, E$
leaves (bottom) root (top)

- The sequence starts with the input string $vd * vd$.
- The first reduction produces $F * vd$ by reducing the leftmost vd to F , using the production $F \rightarrow id$.
- The second reduction produces $T * vd$ by reducing F to T .
- Now there ~~is~~ is choice between reducing the string T , which is the body of $E \rightarrow T$, and the string consisting of the second vd , which is the body of $F \rightarrow id$.
- Rather than reduce T to E , the second vd is reduced to F , resulting in the string $T * F$.
- This string then reduces to T . The parse completes with the reduction of T to the start symbol E .
- A reduction is the reverse of a step in derivation.

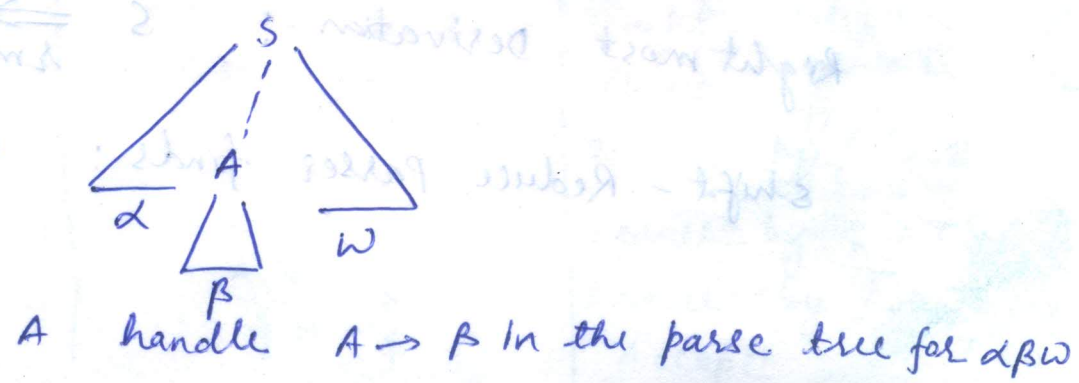
Handle Pruning →

- Bottom-up parsing during a left to right scan of the input constructs a rightmost derivation in reverse.
- Handle is a substring that matches the body of a production and whose reduction represents one step along the reverse of a rightmost derivation.

Right sentential form	Handle	Reducing Production
$vd_1 * id_2$	id_1	$A \rightarrow id$
$F * id_2$	F	$T \rightarrow F$
$T * id_2$	vd_2	$F \rightarrow vd$
$T * F$	$T * F$	$T \rightarrow T * F$
T	T	$E \rightarrow T$

Handles during a parse of $id_1 * id_2$

- Formally, if $S \xRightarrow{rm} \alpha A w \xRightarrow{rm} \alpha \beta w$, as in next figure then production $A \rightarrow \beta$ in the position following α is a handle of $\alpha \beta w$.
- Alternatively, a handle of a right sentential form γ is a production $A \rightarrow \beta$ and a position of γ where the string β may be found such that replacing β at that position ~~of~~ by A produces the previous right sentential form in a rightmost derivation of γ .
- The string w to the right of the handle must contain only terminal symbols.
- For convenience, we refer to the body β rather than $A \rightarrow \beta$ as a handle.



- A rightmost derivation in reverse can be obtained by "handle pruning".
- That is, we start with a string of terminals w to be parsed.
- If w is a sentence of the grammar at hand, then let $w = \gamma_n$, where γ_n is the n th right sentential form of some as yet unknown rightmost derivation.

$$S = \gamma_0 \xRightarrow{rm} \gamma_1 \xRightarrow{rm} \gamma_2 \xRightarrow{rm} \dots \xRightarrow{rm} \gamma_{n-1} \xRightarrow{rm} \gamma_n = w$$

Shift - Reduce Parsing

- A shift reduce parser tries to reduce the given input string into the starting symbol.
- At each reduction step, a substring of the input matching to right side of a production rule is replaced by the non terminal at the left side of that production rule.
- If the substring is chosen ~~to~~ correctly, the right most derivation of that string is created in the reverse order.

Right most Derivation : $S \xRightarrow{rm}^* w$

Shift - Reduce parser finds : $w \xleftarrow{rm} \dots \xleftarrow{rm} S$

• We use \$ to mark the bottom of the stack and also the right end of the input.

stack

Input

\$

w \$

- During a left to right scan of the input string, the parser shifts zero or more input symbols onto the stack, until it is ready to reduce a string β of grammar symbols on top of the stack.
- It then reduces β to the head of the appropriate production.
- Parser repeats this cycle until it has detected an error or until the stack contains the start symbol and the input is empty.

Ex.

$$E \rightarrow E + T \mid T$$

$$T \rightarrow T * F \mid F$$

$$F \rightarrow (E) \mid id$$

Shift Reduce Parser
on input $id_1 * id_2$

stack	Input	Action
\$	$id_1 * id_2$ \$	shift
\$ id_1	* id_2 \$	Reduce by $F \rightarrow id$
\$ F	* id_2 \$	Reduce by $T \rightarrow F$
\$ T	* id_2 \$	shift
\$ $T *$	id_2 \$	shift
\$ $T * id_2$	\$	Reduce by $F \rightarrow id$
\$ $T * F$	\$	Reduce by $T \rightarrow T * F$
\$ T	\$	Reduce by $E \rightarrow T$
\$ E	\$	Accept

There are four possible actions

(84)

1. Shift - shift the next input symbol onto the top of the stack.
2. Reduce - The right end of the string to be reduced must be at the top of the stack.
 - Locate the left end of the string within the stack and decide with what nonterminal to replace the string.
3. Accept - Announce successful completion of parsing.
4. Error - Discover a syntax error and call on error recovery routine.

• Initial stack just contains only the end-marker \$ and also end of the input string is marked by the end-marker \$.

$E \rightarrow E + T \mid T$
 $T \rightarrow T * F \mid F$
 $F \rightarrow (E) \mid id$

Right-most Derivation of $id + id * id$

$E \Rightarrow E + T \Rightarrow E + T * F \Rightarrow E + T * id \Rightarrow E + F * id$
 $\Rightarrow E + id * id \Rightarrow T + id * id \Rightarrow F + id * id$
 $\Rightarrow id + id * id$

Right most sentential form

$id + id * id$
 $F + id * id$
 $T + id * id$
 $E + id * id$
 $E + F * id$
 $E + T * id$
 $E + T * F$
 $E + T$
 E

Reducing Production

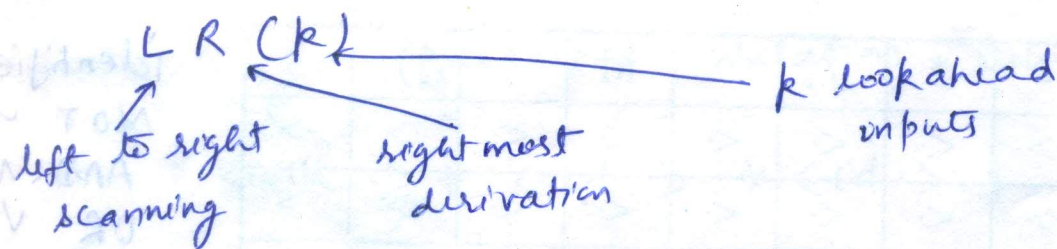
$P \rightarrow id$
 $T \rightarrow F$
 $E \rightarrow T$
 $P \rightarrow id$
 $T \rightarrow F$
 $P \rightarrow id$
 $T \rightarrow T * F$
 $E \rightarrow E + T$

* Handles are underlined

Conflicts During Shift-Reduce Parsing

(85)

- There are context free grammars for which shift-reduce parsers can not be used.
- stack contents and the next input symbol may not decide action.
- Shift/Reduce conflict - whether make a shift operation or a reduction.
- Reduce/reduce conflict - The parser cannot decide which of several reductions to make
- If a shift reduce parser cannot be used for a grammar, that grammar is called as non-LR(CR) grammar



Operator Precedence Parser

A grammar G is said to be operator precedence if it posses following properties -

1. No production on the right side is ϵ .
2. There should not be any production rule possessing two adjacent non-terminals at the right hand side.

Ex - $E \rightarrow EAE \mid (E) \mid -E \mid id$

This grammar is not an operator precedence grammar as in the production rule -

$E \rightarrow EAE$; it contains two consecutive non terminals.

Operator-Precedence Relations

	+	-	*	/	^	id	(	)	\$
+	>	>	<	<	<	<	<	>	>
-	>	>	<	<	<	<	<	>	>
*	>	>	>	>	<	<	<	>	>
/	>	>	>	>	<	<	<	>	>
^	>	>	>	>	<	<	<	>	>
id	>	>	>	>	>			>	>
(	<	<	<	<	<	<	<	=	
)	>	>	>	>	>			>	>
\$	<	<	<	<	<	<	<		

identijer.id

NOT ✓

AND ✓

OR ✓

(Open Bracket

* / ^

+ -

) Close Bracket

\$

In operator precedence parsing we will first define ⁽⁸⁷⁾ precedence relation $<, =, \text{ and } >$ between pairs of terminals.

$p < q \rightarrow p$ gives more ~~prec~~ precedence than q .

$p = q \rightarrow p$ has same precedence as q .

$p > q \rightarrow p$ takes precedence over q .

$$E \rightarrow E + E \mid E - E \mid E * E \mid E / E \mid E \wedge E$$

$$E \rightarrow (E) \mid -E \mid \text{id}$$

Now consider the string -

$$\text{id} + \text{id} * \text{id}$$

- We will insert $\$$ symbols at the start and end of the input string.

- We will also insert precedence operator by referring the precedence relation table.

$$\$ < \text{id} > + < \text{id} > * < \text{id} > \$$$

We will follow following steps to parse the given string -

- 1) scan the input from left to right until first $>$ is encountered.
- 2) scan backwards over = until $<$ is encountered
- 3) The handle is a string between $<$ and $>$.

The parsing can be done as follows -

$\$ \langle \cdot id \cdot \rangle + \langle \cdot id \cdot \rangle * \langle \cdot id \cdot \rangle \$$

Handle id is obtained between $\langle \cdot \cdot \rangle$ Reduce this by $E \rightarrow id$

$E + \langle \cdot id \cdot \rangle * \langle \cdot id \cdot \rangle \$$

Handle id is obtained b/w $\langle \cdot \cdot \rangle$ Reduce this by $E \rightarrow id$

$E + E * \langle \cdot id \cdot \rangle \$$

Handle id is obtained b/w $\langle \cdot \cdot \rangle$ Reduce this by $E \rightarrow id$

$E + E * E$

Remove all the non-terminals

$+ *$

Insert $\$$ at the beginning at the end also insert the precedence operator

$\$ \langle \cdot + \cdot \rangle \langle \cdot * \cdot \rangle \$$

The $*$ operator is surrounded by $\langle \cdot \cdot \rangle$ This indicates that

$*$ becomes handle. That means we have to reduce $E * E$ operation first

$\$ \langle \cdot + \cdot \rangle \$$

Now $+$ becomes handle. Hence we evaluate $E + E$

$\$ \$$

Parsing is done.

How to create operator Precedence Relations (89)

step 1. check whether grammar is operator grammar or not

step 2. calculate the leading and trailing for every non terminal involved in given grammar.

1. For any production like $X \rightarrow YaZ$

Leading $(X) = \{a\}$ if Y is or a single non terminal

2. For any production like $X \rightarrow Bx$

Leading $(X) = \text{Leading}(B)$

3. For any production like $X \rightarrow YaZ$

Trailing $(X) = \{a\}$ if Z is or a single non terminal

4. For any production like $X \rightarrow \alpha B$

Trailing $(X) = \text{Trailing}(B)$

5. For any production like $A \rightarrow B$

Leading $(A) = \text{Leading}(B)$, Trailing $(A) = \text{Trailing}(B)$

6. For any production like $A \rightarrow a$

Leading $(A) = \{a\}$, Trailing $(A) = \{a\}$

$$E \rightarrow E + T \Rightarrow LCB = + \text{ (rule-1)}$$

$$T(CB) = + \text{ (rule-3)}$$

$$| E \rightarrow T \Rightarrow LCB = L(T) \text{ (Rule-5)}$$

$$T(CB) = T(T)$$

$$T \rightarrow T * F \Rightarrow L(T) = * \text{ (rule-1)}$$

$$T(T) = * \text{ (rule-3)}$$

$$| T \rightarrow F \Rightarrow L(T) = L(F) \text{ (rule-5)}$$

$$T(T) = T(F)$$

$$F \rightarrow (E) \Rightarrow L(F) = (\text{ (rule-1)}$$

$$T(F) =) \text{ (rule-3)}$$

$$F \rightarrow id \Rightarrow L(F) = T(F) = id \text{ (Rule-6)}$$

Non-Terminal	E	T	F
Leading (L)	+, *, (, id	*, (, id	(, id
Trailing (T)	+, *,), id	*,), id	), id

Step 3 - Now establish the precedence relation between terminals as follows -

(i) $a = \cdot b \Rightarrow$ if on RHS of any production $x \rightarrow \alpha \beta b \delta$ where β is nothing or single non terminal

(ii) $a < \cdot b$: if b is on Lead (B) where B is immediate

right non terminal of a in any production,

(91)

$$A \rightarrow \alpha a B \beta \quad a(\text{row}) < \{ \text{all lead of } B \} \quad (\text{column})$$

3. $a \rightarrow b$: if a is in Last(B) where B is immediate left non terminal of a in any production,

$$A \rightarrow \alpha B b \beta \quad \{ \text{all trailing of } B \} \rightarrow b$$

4. $\$ < \{ \text{all lead of } S \}$

5. $\{ \text{all trailing of } S \} \rightarrow \$$

	+	*	(	)	od	\$
+	.	<	<	.	<	.
*	.	.	<	.	<	.
(	<	<	<	=	<	
)	.	.		.		.
od	.	.		.		.
\$	<	<	<		<	

$$E \rightarrow E + T \quad \text{rule-2}$$

$$E \rightarrow T \quad \text{rule-3}$$

$$T \rightarrow T * F \quad \text{rule-2}$$

$$T \rightarrow F \quad \text{rule-3}$$

$$P \rightarrow (E) \quad \text{rule-1 \& 2}$$

$$P \rightarrow od \quad \text{rule-3}$$

$$\text{rule-4}$$

$$\text{rule-5}$$

Non-Terminal	E	T	F
Leading (L)	+, *, (, id	*, (, id	(, od
Trailing (T)	+, *,), od	*,), id	), id

Disadvantages of Operator Precedence Parsing

- Disadvantages:

- It is difficult to handle operators (like - unary minus) which have different precedence (the lexical analyzer should handle the unary minus).
- Small class of grammars.
- Difficult to decide which language is recognized by the grammar.

- Advantages:

- Simple and Easy to implement
- powerful enough for expressions in programming languages
- Can be constructed by hand after understanding the grammar.
- Simple to debug

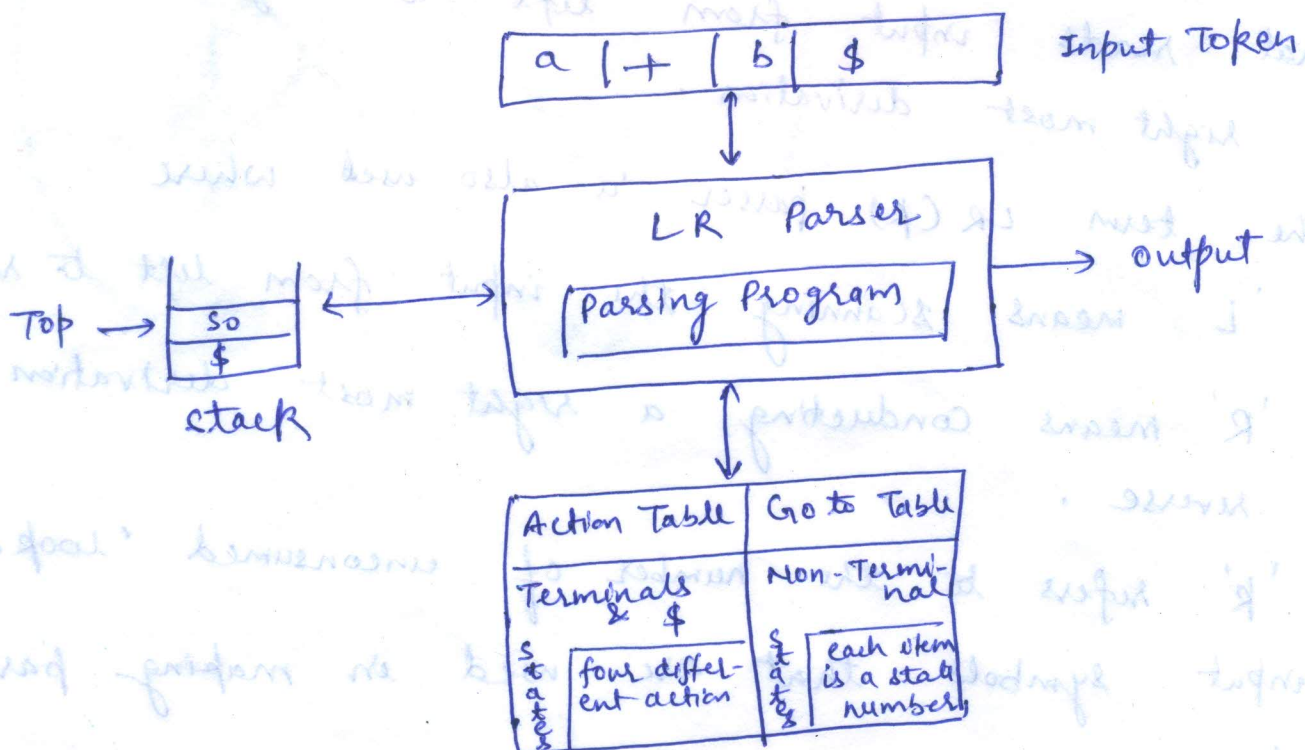
LR Parser

- An LR parser is a parser for context free grammars that reads input from left to right and produces a right most derivation.
- The term LR (k) parser is also used where
 - 'L' means scanning the input from left to right.
 - 'R' means conducting a right most derivation in reverse.
 - 'k' refers to the number of unconsumed 'lookahead' input symbols that are used in making parsing decisions.

Properties of LR parser -

- (i) LR parsers can be constructed to recognize most of the programming languages for which context free grammar can be written.
 - (ii) The class of grammar that can be passed by LR parser is a superset of class of grammars that can be passed using predictive parsers.
 - (iii) LR parser works using non backtracking shift reduce technique yet it is efficient one.
- LR parsers detect syntactical errors very efficiently.

Structure of LR parser →



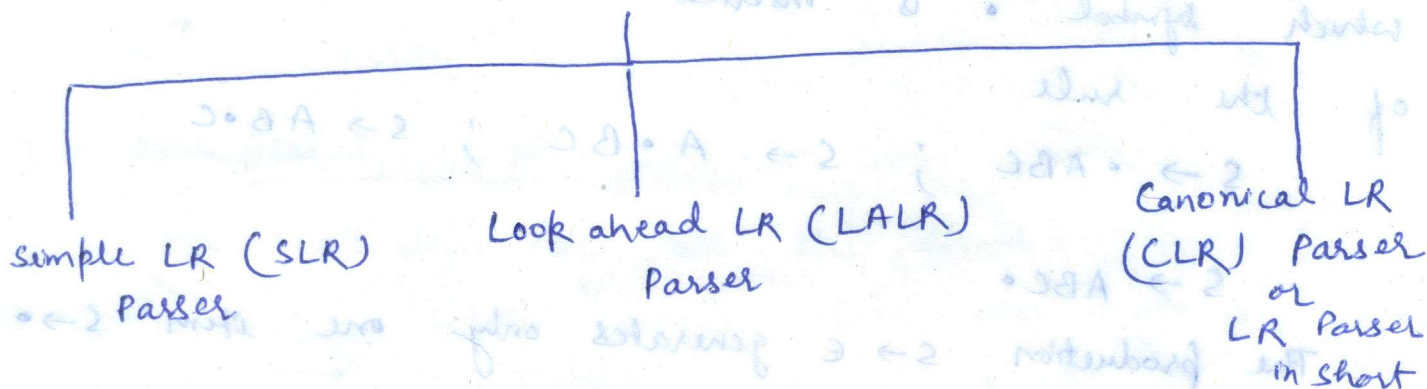
Structure of LR parser

The driving program works on following line —

(96)

1. It initializes the stack with start symbol and invokes scanner (lexical analyzer) to get next token.
2. It determines s_i the state currently on the top of the stack and a_i the current input symbol.
3. It consults the parsing table for the action $[s_i, a_i]$ which can have one of the four values.
 - (i) s_i means shift state i .
 - (ii) r_j means reduce by rule j .
 - (iii) Accept means successful parsing is done.
 - (iv) Error indicates syntactical error.

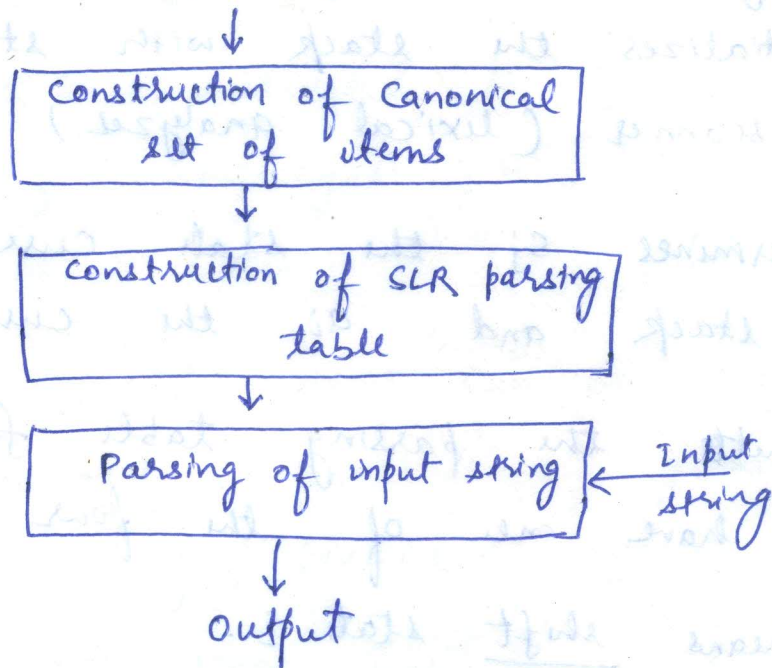
LR Parser



Simple LR (SLR) Parser

- It is the weakest of the three methods but it is easiest to implement.
- A grammar for which SLR parser can be constructed is called SLR grammar.

Context free grammar



Working of SLR (1)

Definition of LR(0) items and related terms →

(i) The LR(0) item for grammar G is production rule in which symbol $\cdot$ is inserted at some position in RHS of the rule.

$$S \rightarrow \cdot ABC \quad ; \quad S \rightarrow A \cdot BC \quad ; \quad S \rightarrow AB \cdot C$$
$$S \rightarrow ABC \cdot$$

The production $S \rightarrow \epsilon$ generates only one item $S \rightarrow \cdot$.

(ii) Augmented grammar :-

• If a grammar G is having start symbol S then augmented grammar is a new grammar G' in which S' is the new start symbol such that $S' \rightarrow S$.

(98)

The purpose of this grammar is to indicate the acceptance of input i.e. when parser is about to reduce $S' \rightarrow S$ it reaches to acceptance state.

(ii) Nonkernel items $\rightarrow$

- It is collection of items in which $\cdot$ are at the left end of R.H.S. of the rule.

(iv) Functions closure and go to $\rightarrow$

- These are two important functions required to create collection of canonical set of items.

(v) Viable prefix $\rightarrow$

- It is the set of prefixes in the right sentential form of production $A \rightarrow \alpha$.
- This set can appear on the stack during shift/reduce action.

(vi) ~~Non~~ kernel items $\rightarrow$

- The collection of ~~all~~ the items $\leftarrow S' \rightarrow \cdot S$ and all the items whose dots are not at the left most end of R.H.S. of the rule.

~~Closure~~ Closure operation $\rightarrow$

For a context free grammar G , if I is the set of items then the function closure ($\bar{I}$) can be constructed using these rules -

1. Consider I is a set of canonical items and initially every item I is added to closure(I).

2. If rule $A \rightarrow \alpha \cdot B\beta$ is a rule in closure(I) and there is another rule for B such as $B \rightarrow \gamma$ then

closure(I): $A \rightarrow \alpha \cdot B\beta$

$B \rightarrow \cdot \gamma$

This rule has to be applied until no more new items can be added to closure(I).

• The meaning of rule $A \rightarrow \alpha \cdot B\beta$ is that during derivation of the input string ~~xxx~~ at some point we may require strings ~~xxxxxx~~ derivable from $B\beta$ as input.

• A nonterminal immediately to the right of $\cdot$ indicates that it has to be expanded shortly.

Goto operation -

• If there is a production $A \rightarrow \alpha \cdot B\beta$ then $\text{goto}(A \rightarrow \alpha \cdot B\beta, B) = A \rightarrow \alpha B \cdot \beta$.

• That means simply shifting of $\cdot$ one position ahead over the grammar symbol (may be terminal or non terminal)

• The rule $A \rightarrow \alpha \cdot B\beta$ is in I then the same goto function can be written as $\text{goto}(I, B)$.

Ex. $X \rightarrow xb/a$

compute closure (I) and goto (I)

Let $I: X \rightarrow \cdot xb$

$$\begin{aligned} \text{closure (I)} &= X \rightarrow \cdot xb + X \rightarrow \cdot a \\ &= X \rightarrow \cdot a \end{aligned}$$

goto function can be computed as

$$X \rightarrow \cdot \overset{\curvearrowright}{x} b \Rightarrow X \rightarrow x \cdot b$$

$$X \rightarrow \cdot \overset{\curvearrowright}{a} \Rightarrow X \rightarrow a \cdot$$

$$\text{goto (I, x)} = X \rightarrow x \cdot b$$

$$\text{goto (I, a)} = X \rightarrow a \cdot$$

Ex - Consider the following grammar -

$$E \rightarrow E + T / T$$

$$T \rightarrow T * F / F$$

$$F \rightarrow (E) / \text{vd}$$

(a) Write the augmented grammar corresponding to the given grammar

(b) If $I = \{ [E' \rightarrow \cdot E] \}$

then find closure (I)

(c) If $I = \{ [E' \rightarrow E \cdot], [E \rightarrow E \cdot + T] \}$

then find goto (I, +).

(a) augmented grammar

$$E' \rightarrow x E$$

$$E \rightarrow E + T / T$$

$$T \rightarrow T * F / F$$

$$F \rightarrow (E) / \text{vd}$$

(b) $I = \{ [E' \rightarrow \cdot E] \}$

closure (I) :

$$E' \rightarrow \cdot E$$

$$E \rightarrow \cdot E + T$$

$$E \rightarrow \cdot T$$

$$T \rightarrow \cdot T * F$$

$$T \rightarrow \cdot F$$

$$F \rightarrow \cdot (E)$$

$$F \rightarrow \cdot \text{vd}$$

given that $I = \{ [E' \rightarrow E.] , [E \rightarrow E. + T] \}$

$$\text{goto}(I, +) = \text{closure}([E \rightarrow E + \cdot T])$$

therefore $\text{goto}(I, +)$ consists of

$$E' \rightarrow E + \cdot T$$

$$T \rightarrow \cdot T * F$$

$$T \rightarrow \cdot F$$

$$F \rightarrow \cdot (B)$$

$$F \rightarrow \cdot \text{od}$$

Construction of Canonical Collection of LR(0) items $\rightarrow$

- For the grammar G initially add $S' \rightarrow \cdot S$ in the set of item $C(\text{closure})$.
- For each set of items I_i in C and for each grammar symbol x (may be terminal or non-terminal) add $\text{closure}(I_i, x)$.
- Above process should be repeated by applying $\text{goto}(I_i, x)$ for each x in I_i such that $\text{goto}(I_i, x)$ is not empty and not in C .
- set of items has to be constructed until no more set of items can be added to C .

$$E \rightarrow E + T$$

$$E \rightarrow T$$

$$T \rightarrow T * F$$

$$T \rightarrow F$$

$$F \rightarrow (E)$$

$$F \rightarrow id$$

- In this grammar we will add the augmented grammar $E' \rightarrow \cdot E$ in the I then we have to apply closure.

I_0 :

$$E' \rightarrow \cdot E$$

$$E \rightarrow \cdot E + T$$

$$E \rightarrow \cdot T$$

$$T \rightarrow \cdot T * F$$

$$T \rightarrow \cdot F$$

$$F \rightarrow \cdot (E)$$

$$F \rightarrow \cdot id$$

→ The item I_0 is constructed starting from the grammar $E' \rightarrow \cdot E$.

→ Now immediately right to $\cdot$ is E . Hence we have applied closure (I_0) and thereby we add E -productions with $\cdot$ at the left end of the rule.

→ That means we have added $E \rightarrow \cdot E + T$ and $E \rightarrow \cdot T$ in I_0 .

→ But again the rule $E \rightarrow \cdot T$ which we have added contains non terminal T immediately right to $\cdot$.

→ So we have added T -productions in I_0 .

$$T \rightarrow \cdot T * F \quad \text{and} \quad T \rightarrow \cdot F$$

→ In T -productions after $\cdot$ comes T and F respectively.

• But since we have already added T productions so we will not add these.

• But we will add the F -productions having dots.

- The $F \rightarrow \cdot (E)$ and $F \rightarrow \cdot id$ will be added.
- None after dot 'c' and 'id' are coming in these two productions.
- The c and id are terminal symbols and are not deriving any rule. Hence our closure function terminates over here.
- Since there is no rule further we will stop creating I_0 .

Now apply goto (I_0, E)

$E' \rightarrow \cdot E$	shift dot	$E' \rightarrow E \cdot$
$E \rightarrow \cdot E + T$	to right	$E \rightarrow E \cdot + T$

Thus I_1 becomes,

goto (I_0, E)
 $I_1: E' \rightarrow E \cdot$
 $E \rightarrow E \cdot + T$

Since in I_1 , there is ~~no~~ no nonterminal after dot we cannot apply closure (I_1).

By applying goto on T of I_0

goto (I_0, T)
 $I_2: E \rightarrow T \cdot$
 $T \rightarrow T \cdot * F$

Since on I_2 there is no nonterminal after dot we cannot apply closure (I_2).

By applying goto on F of I_0

goto (I_0, F)
 $I_3: T \rightarrow F \cdot$

since after dot in I_3 there is nothing, hence (104)
we cannot apply closure (I_3).

- By applying goto on C of I_0 but after dot E comes hence we will apply closure on E, then on T, then on F.

goto (I_0, C)
 I_4 : $T \rightarrow (\cdot E)$
 $E \rightarrow \cdot E + T$
 $E \rightarrow \cdot T$
 $T \rightarrow \cdot T * F$
 $T \rightarrow \cdot F$
 $F \rightarrow \cdot (E)$
 $F \rightarrow \cdot id$

$\Rightarrow$

By applying goto on
id of I_0

goto (I_0, id)
 I_5 : $F \rightarrow id \cdot$

since in I_5 there is no non-terminal to the right of dot
we cannot apply closure here.

• Thus we have completed applying gotos on I_0 .

→ We will consider I_1 for applying goto.

• In I_2 there are two productions $E' \rightarrow E \cdot$ and

$E \rightarrow E \cdot + T$

• There is no point applying goto on $E' \rightarrow E \cdot$,
hence we will consider $E' \rightarrow E \cdot + T$ for applying
of goto

goto ($I_1, +$)
 I_6 : $E \rightarrow E + \cdot T$
 $T \rightarrow \cdot T * F$
 $T \rightarrow \cdot F$
 $F \rightarrow \cdot (E)$
 $F \rightarrow \cdot id$

There is no other rule in I_1 for applying goto. (105)

• so we will move to I_2 . In I_2 there are two productions: $E \rightarrow T \cdot$ and $T \rightarrow T \cdot * F$.

• we will apply goto on $*$.

goto ($I_2, *$)

$I_7: T \rightarrow T * \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow \cdot id$

• The goto cannot be applied on I_3 . Apply goto on E in I_4 . In I_4 there are two productions having E after dot ($F \rightarrow \cdot E$ and $E \rightarrow \cdot E + T$).

• Hence we will apply goto on both of these productions. The I_8 becomes -

goto (I_4, E)

$I_8: F \rightarrow (E \cdot)$

$E \rightarrow E \cdot + T$

• If we will apply goto on (I_4, T) but we get $E \rightarrow T \cdot$ and $T \rightarrow T \cdot * F$ which is I_2 only. Hence we will not apply goto on T .

• Similarly we will not apply goto on F , C and id as we get the states I_3 , I_4 , I_5 again.