

2nd UNIT.

SYNTAX ANALYSIS

- Introduction
- Top-Down parsing
- Bottom-Up parsing
- LR, SLR, LALR, CLR
- Using Ambiguous Grammar
- parser Generators - YACC

Brute-Force [eg at end of unit]

→ top-down parsing with full backup is a BRUTE-FORCE method of parsing.

Steps

1. Given a particular non-terminal that is to be expanded, the first prodⁿ for this V is applied.
2. then, within this newly expanded string, the next (leftmost) nonterminal is selected for expansion & its first prodⁿ is applied.
3. (step 2) repeats until process cannot be continued. This termination (if it ever occurs) may be due to two causes.
 - First, no more V may be present, in which case the string has been successfully parsed.
 - Second, Incorrect Expansion, then the process is "BACKED UP" by undoing the most recently applied prodⁿ. & then next prodⁿ of this V is used as next expansion, & then process continues as before.

Introduction

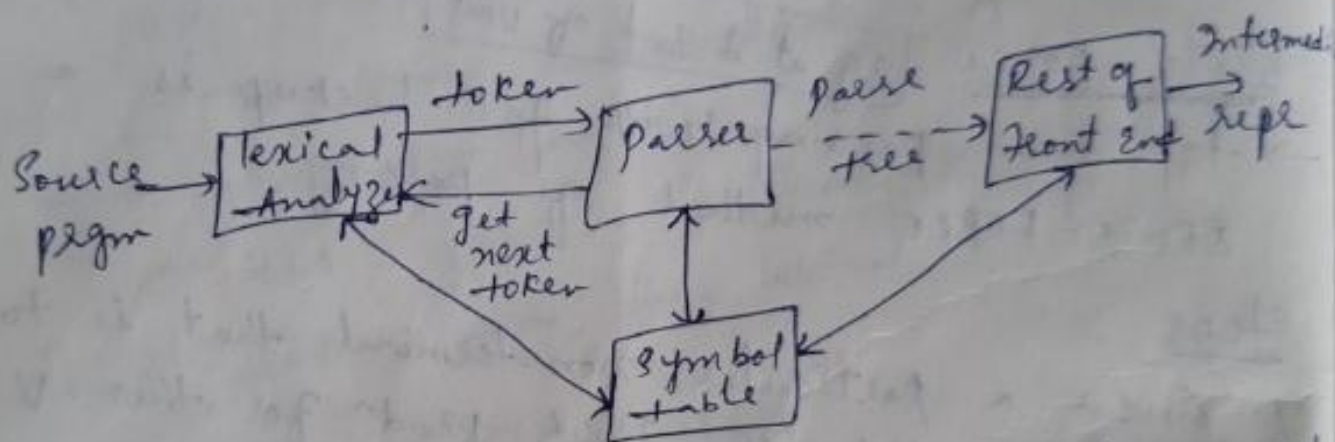
the role of the parser

Representative Grammars

Syntax Error Handling

Error Recovery Strategies

The Role of the parser



Position of parser in Compiler model

→ parser obtains a string of tokens from the lexical analyzer and verifies that the string of token names can be generated by the grammar for the source lang.

- we can expect the parser to report any syntax errors to continue in an intelligible fashion and to recover from commonly occurring

Errors to continue processing the remainder of the program.

- the parser constructs a parse tree and passes it to the rest of the compiler for further processing
- the parse tree need not be constructed explicitly, since checking and translation actions can be ~~interposed~~ interspersed with parsing.
- the parser and the rest of the front end could well be implemented by a single module.

There are 3 general types of parser for grammars

Universal parser

top-down "

Bottom-up "

Universal parser methods are too inefficient to use in production compilers.

Top-down parser builds parse trees from top (root) to the bottom (leaves)

Bottom-up parser builds parse tree from bottom (leaves) to the top (root).

In either case, the input to the parser is scanned from left to right, one symbol at a time.

Representative Grammars

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

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

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

- Expression Grammar belongs to the class of LR grammars that are suitable for Bottom-up parsing.

- It can't be used for top-down parsing because it is left recursive.

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

- Non-left-recursive variant of the Expression Grammar will be used for top-down parsing.

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

- The following grammar treats $+$ & $*$ alike, so it is useful for illustrating techniques for handling ambiguities during parsing.

Syntax Error Handling

Common programming errors can occur at many different levels.

1. Lexical Errors include misspellings of id's, keywords or operators
eg:- `ellipseSize` instead of `EllipseSize` & misquotes around text.
2. Syntactic Errors include misplaced semicolons or extra or missing braces;
3. Semantic Errors include type mismatches b/w operators & operands
4. Logical Errors assignment operator `=` instead of the comparison operator `==`.

The Error Handler in a parser has goals that are simple to state but challenging to realize:

- Report the presence of errors clearly & accurately.
- Recover from each error quickly enough to detect subsequent errors.
- Add minimal overhead to the processing of correct pgms.

Error-Recovery Strategies

Panic-mode

phrase-level

Error-productions

Global-correction

Panic-Mode Recovery

- on discovering an error, the parser discards all symbols one at a time until one of a designated set of synchronizing tokens is found.
- the synchronizing tokens are usually delimiters, such as semicolon or $\}$ whose role in the source pgm is clear & unambiguous.

- panic mode correction often skips a considerable amount of η/p w/o checking it for additional errors,
- it has the advantage of simplicity & is guaranteed not to go into an infinite loop.

phrase-level Recovery

- on discovering an error, a parser may perform local correction on the remaining η/p .
- it may replace a prefix of the remaining η/p by some string that allows the parser to continue.
- A typical local correction is to replace a comma by a semicolon, delete an extraneous semicolon, or insert a missing semicolon.
- phrase-level replacement has been used in several error-repairing compilers as it can correct any η/p string.
- its major drawback is the difficulty it has to coping with situations in which the actual error has occurred before the point of detection.

Error productions

- Common errors that might be encountered, we can augment the grammar for the lang at hand with productions that generate the erroneous construct.
- A parser constructed from a grammar augmented by these error productions detects the anticipated errors when an error production is used during parsing.
- The parser can then generate appropriate error diagnostics about the erroneous construct that has been recognized in the i/p.

Global Correction

- we would like a compiler to make as few changes as possible in processing an incorrect i/p string.
- there are algorithms for choosing a minimal sequence of changes to obtain a globally least-cost correction.
- unfortunately, these methods are in general too costly to implement in terms of time & space, so these techniques are currently only of theoretical interest.

Derivation & Parse Trees

- Derivation from S means generation of string w from S.

For constructing derivation two things are imp.

- (i) choice of non-terminal from several others.
- (ii) choice of rule from production rules for corresponding non-terminal.

Instead of choosing the arbitrary non-terminal one can choose.

- (i) either leftmost non-terminal in a sentential form then it is called leftmost derivation.
- (ii) or rightmost non-terminal in a sentential form, then it is called rightmost derivation.

$$\begin{aligned} S &\rightarrow aB \mid bA \\ A &\rightarrow a \mid aS \mid bAA \\ B &\rightarrow b \mid bS \mid aBB \end{aligned}$$

derive the string $aaabbbabbbba$ using above grammar

Sol leftmost derivation

$$\begin{aligned} S &\rightarrow aB \\ &\rightarrow a \underline{aB}B \\ &\rightarrow aa \underline{aBB}B \\ &\rightarrow aaa \underline{bS}BB \\ &\rightarrow aaa \underline{bbA}BB \\ &\rightarrow aaa \underline{bbab}B \\ &\rightarrow aaa \underline{bbab}B \end{aligned}$$

$$\begin{aligned} &aaabbbabbb \underline{S} \\ &aaabbbabbb \underline{A} \\ &aaabbbabbb a \end{aligned}$$

Derivation & Parse Trees

— Derivation from S means generation of string w from S .

For constructing derivation two things are imp.

- (i) choice of non-terminal from several others.
- (ii) choice of rule from production rules for corresponding non-terminal.

Instead of choosing the arbitrary non-terminal one can choose.

- (i) Either leftmost non-terminal in a sentential form then it is called leftmost derivation.
- (ii) or rightmost non-terminal in a sentential form, then it is called rightmost derivation.

$$S \rightarrow aB \mid bA$$

$$A \rightarrow a \mid aS \mid bAA$$

$$B \rightarrow b \mid bS \mid aBB$$

derive the string $aaabbabbbba$ using above grammar

Sol leftmost derivation

$$S \rightarrow a\underline{B}$$

$$a \underline{aB} B$$

$$aa \underline{aBB} B$$

$$aaa \underline{bS} BB$$

$$aaa b \underline{aBB}$$

$$aaa b b \underline{aBB}$$

$$aaa b b a \underline{b} B$$

$$aaabbabbb\underline{S}$$

$$aaabbabbb\underline{A}$$

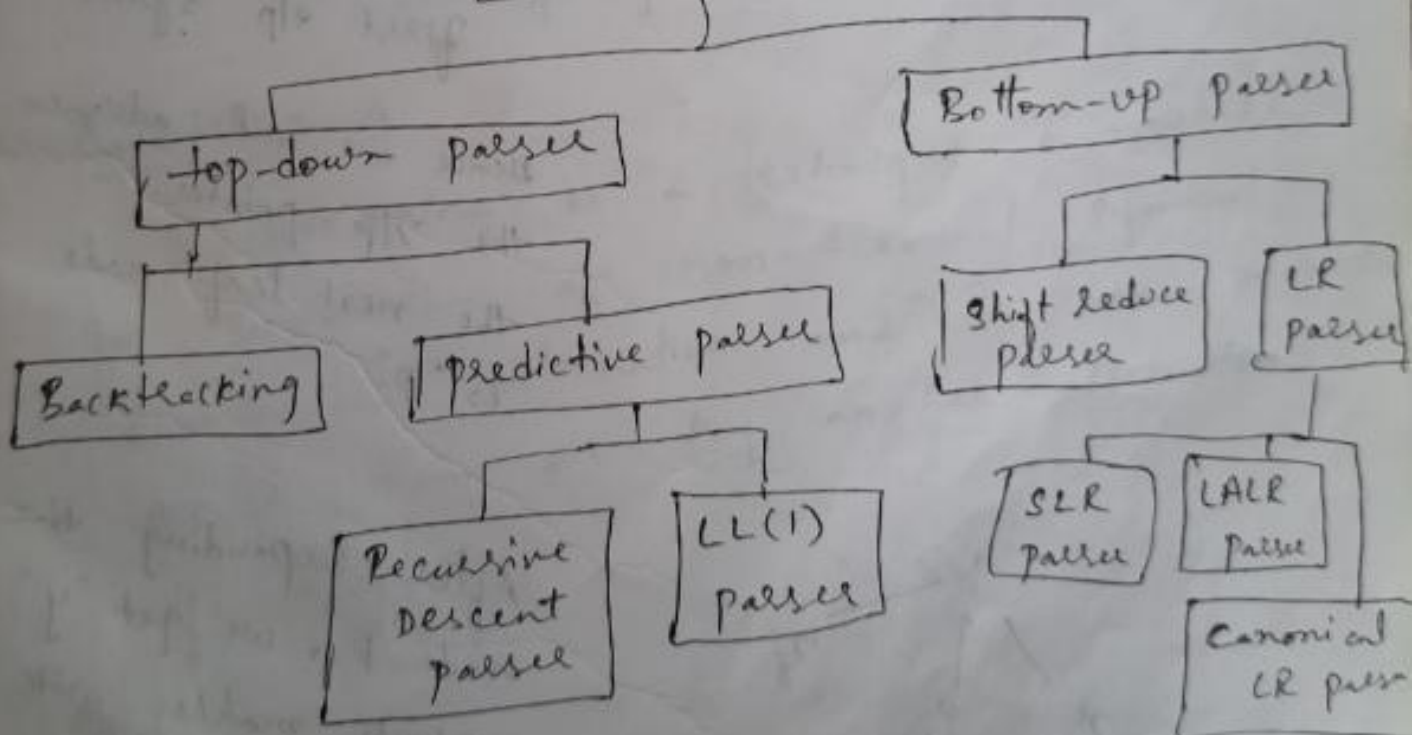
$$aaabbabbbba$$

Parsing Techniques

passing techniques work on the following principle.

1. The parser scans the input string from left to right and identifies that the derivation is leftmost or rightmost.
2. the parser make use of production rules for choosing the appropriate derivation.

Types of parser



Top Down Parser

parse tree is generated from top to bottom

Consider a Grammar

$$S \rightarrow x p z$$

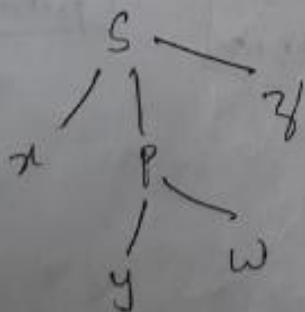
$$P \rightarrow y w | y$$

9/p :- x y z



the leftmost leaf of the parse tree matches with the first 9/p symbol "x"

Hence we will advance the 9/p pointer. the next leaf node is "P".



After expanding the node "P", we get "y" which matches with the 9/p symbol.

Now the next node is "w" which is not matching with the 9/p symbol. Hence we go back to see

- whether there is another alternative of "P".
- Now the next 9/p symbol is "z" which matches the 9/p symbol.
 - we halt and declare that the parsing is completed successfully.

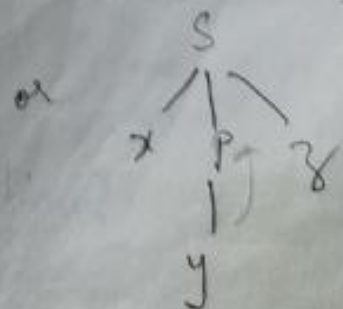
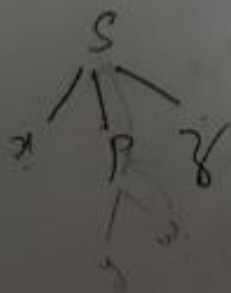
Problems with Top-down parsing

Backtracking
 left Recursion
 left factoring
 Ambiguity

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 alternatives if any,

Eg:- $S \rightarrow xPz$
 $P \rightarrow yw/z$



Left Recursion

Let A be a cfa having a production rule with left Recursion

$$A \rightarrow A\alpha$$

$$A \rightarrow \beta$$

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

$$A \rightarrow \beta A'$$

$$A' \rightarrow \alpha A'$$

$$A' \rightarrow \epsilon$$

Eg:-

Consider the Grammar

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

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

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

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

$$A \rightarrow A\alpha \mid \beta$$

$$A = E$$

$$\alpha = +T$$

$$\beta = T$$

$$A \rightarrow \beta A'$$

$$A' \rightarrow \alpha A' \mid \epsilon$$

then 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

$$T \rightarrow FT'$$

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

So the equivalent grammar will be

$$\begin{aligned} E &\rightarrow TE' \\ E' &\rightarrow +TE' \mid \epsilon \\ T &\rightarrow FT' \\ T' &\rightarrow +FT' \mid \epsilon \\ F &\rightarrow (E) \mid id \end{aligned}$$

Left Factoring

In general if

$$A \rightarrow \alpha B_1 \mid \alpha B_2$$

is a production then it is not possible for us to take a decision whether to choose first rule or second.

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

$$A \rightarrow \alpha A'$$

$$A' \rightarrow B_1 \mid B_2$$

For eg:- $S \rightarrow iEtS \mid iEtSes \mid a$

$$E \rightarrow b$$

The left factored grammar becomes

$$S \rightarrow iEtSS' \mid a$$

$$S' \rightarrow eS \mid \epsilon$$

$$E \rightarrow b$$

(iv) Ambiguity
 to remove Ambiguity we will apply
 one rule: if the grammar has left
Associative operator (+, -, *, /) then
 include the left Recursion and if
 the grammar has right Associative operator
 (Exponential operator) then include the
Right Recursion.

the Unambiguous Grammar is

$$E \rightarrow E + T$$

$$E \rightarrow T$$

$$T \rightarrow T * F$$

$$T \rightarrow F$$

$$F \rightarrow id$$

PREDICTIVE PARSER

As the name suggests. the predictive
 parser tries to predict the next construction
 using one or more lookahead symbols
 from a/p string.

there are two types of predictive parsers:

1. Recursive Descent
2. LL(1) parser.

Recursive Descent parser

- A parser that uses collection of Recursive procedures for parsing the given η/p string is called Recursive Descent (RD) parser.
- the R.H.S of the production rule is directly converted to a program.
- 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.
- if the η/p symbol is non-terminal then a call to the procedure corresponding to the non-terminal is made.
- if the η/p symbol is terminal then it is matched with the lookahead from η/p . The lookahead pointer has to be advanced on matching of the η/p symbol.

- if the production rule has many alternates ~~then~~ has to be combined into single body of procedure.

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

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

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

procedure E

{ if lookahead = num then

{ match(num);
T;

else error;

if lookahead = '\$'

{ declare success;

else error;

procedure T

{ if lookahead = '*'

{ match('*');

```

if lookahead == num
{
    match(num);
    T;
}
else error;
}
else NULL;
}

```

```

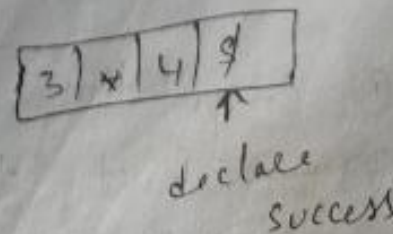
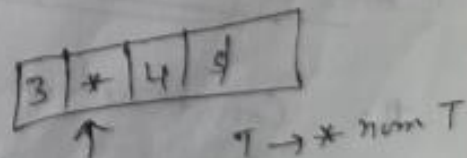
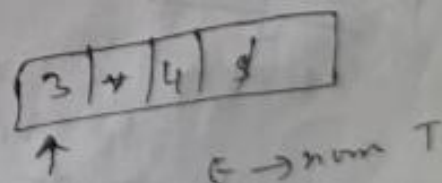
procedure match (token t)
{
    if lookahead == t
    lookahead = next-token;
    else error
}

```

```

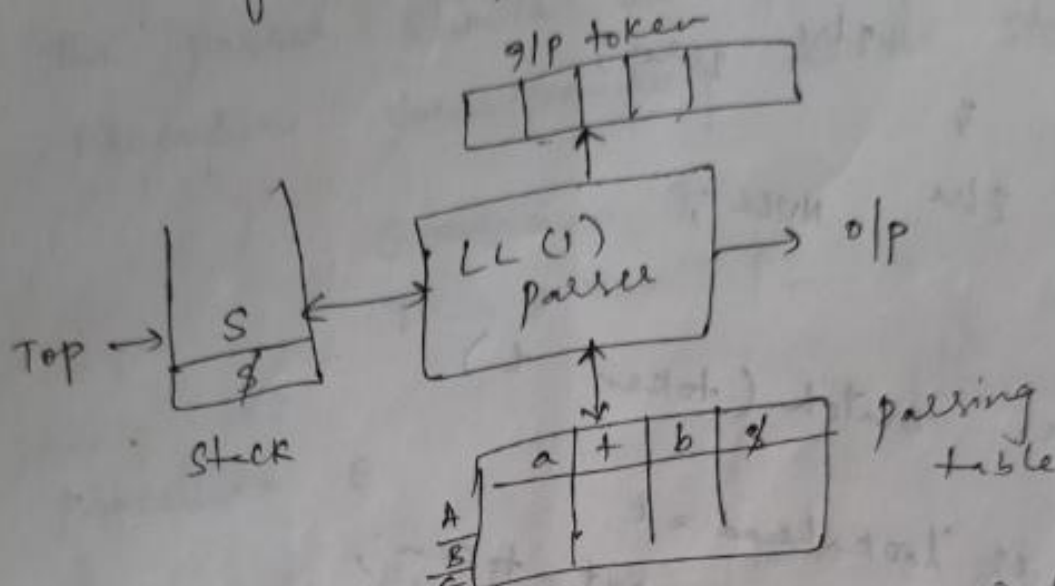
procedure error
{
    print("error!!");
}

```



Predictive LL(1) parser

LL(1) $\rightarrow$ 1 lookahead
 $\Delta, \Rightarrow$ leftmost derivation
 left to right g/p scan



Stack :- the symbols in R.H.S of rule are pushed into the stack in Reverse order.

Parsing table :- $M[A, a]$
 $\downarrow$ terminal
 $\downarrow$ non-terminal

Construction of Predictive LL(1) parser

The construction of predictive LL(1) parser is based on two very imp fns and those are FIRST & FOLLOW

- ① Computation of FIRST & FOLLOW fn.
- ② Construct the predictive parsing table using FIRST & FOLLOW fns.

- ③. parse the g/p string with the help of predictive parsing table.

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 $FIRST$ fns.

①. if the terminal symbol "a" then $FIRST(a) = \{a\}$.

②. if there is a rule $X \rightarrow \epsilon$ then $FIRST(X) = \{\epsilon\}$

③. for the rule $A \rightarrow X_1 X_2 X_3 \dots X_K$
 $FIRST(A) = FIRST(X_1) \cup FIRST(X_2) \dots \cup FIRST(X_K)$
where $K < X_j \leq n$ such that $1 \leq j \leq K-1$

FOLLOW fn.

$\text{FOLLOW}(A)$ is defined as the set of terminal symbols that appear immediately to the right of A .

In other words

$$\text{FOLLOW}(A) = \{a \mid S \Rightarrow \alpha A a \beta\}$$

where α & β are some grammar

symbols may be terminal or non-terminal.

the rules for computing FOLLOW

1. For the start symbol S place $\$$ in $\text{FOLLOW}(S)$.
2. if there is a production $A \rightarrow \alpha B \beta$ then everything in $\text{FIRST}(\beta)$ w/o ϵ is to be placed in $\text{FOLLOW}(B)$.
3. if there is a production $A \rightarrow \alpha B \beta$ or $A \rightarrow \alpha B$ and

$\text{FIRST}(\beta) = \{\epsilon\}$ then

$\text{FOLLOW}(A) = \text{FOLLOW}(B)$ or

$\text{FOLLOW}(B) = \text{FOLLOW}(A)$

Problem for LL(1) parser

Consider the Grammar

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

Step 1:- Computing FIRST & FOLLOW

→ $E \rightarrow TE'$ is a rule in which the first symbol at R.H.S is **T**.

New $T \rightarrow FT'$ in which the first symbol at R.H.S is **F** &

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

$$\therefore FIRST(E) = FIRST(T) = FIRST(F)$$

As $F \rightarrow (E)$

$$F \rightarrow id$$

$$\text{Hence } FIRST(E) = FIRST(T) = FIRST(F) = \{ (, id \}$$

→ $E' \rightarrow +TE' \mid \epsilon$
 $FIRST(E') = \{ +, \epsilon \}$

→ $T' \rightarrow *FT' \mid \epsilon$
 $FIRST(T') = \{ *, \epsilon \}$

Symbols	FIRST
E	{ (, id }
E'	{ +, ε }
T	{ (, id }
T'	{ *, ε }
F	{ (, id }

FOLLOW(E)

(i) As there is a rule $F \rightarrow (E)$ the symbol '(' appears immediately to the right of E.

Hence '(' will be in FOLLOW(E)

Since E is a start symbol, add \$ to FOLLOW of E.

Hence $\text{FOLLOW}(E) = \{ (, \$ \}$

FOLLOW(E')

(i) $E \rightarrow TE'$

(ii) $E' \rightarrow +TE'$

$E \rightarrow TE'$ the computational rule is $A \rightarrow \alpha B \beta$

$A = E, \alpha = T, B = E', \beta = \epsilon$ then by rule 3

Everything in $\text{FOLLOW}(A) = \text{FOLLOW}(B)$

i.e. $\text{FOLLOW}(E) = \text{FOLLOW}(E')$

$\therefore \text{FOLLOW}(E') = \{ (, \$ \}$

$E' \rightarrow +TE'$ the computational rule is $A \rightarrow \alpha B \beta$

$A = E', \alpha = +T, B = E', \beta = \epsilon$ then by rule 3

Everything in $\text{FOLLOW}(A) = \text{FOLLOW}(B)$

i.e. $\text{FOLLOW}(E') = \text{FOLLOW}(E')$

$\therefore \text{FOLLOW}(E') = \{ (, \$ \}$

FOLLOW(T)

- (i) $E \rightarrow TE'$
- (ii) $E' \rightarrow +TE'$

$E \rightarrow TE'$ the computational rule is $A \rightarrow \alpha B \beta$

$A = E$, $\alpha = \epsilon$, $B = T$, $\beta = E'$ then by **rule 2**

$$\text{FOLLOW}(B) = \{\text{FIRST}(\beta) - \epsilon\}$$

$$\begin{aligned}\therefore \text{FOLLOW}(T) &= \{\text{FIRST}(E') - \epsilon\} \\ &= \{+, \epsilon\} - \epsilon \\ &= \{+\end{aligned}$$

$E' \rightarrow +TE'$ we will map it $A \rightarrow \alpha B \beta$

$A = E'$, $\alpha = +T$, $B = E'$, $\beta = \epsilon$ then by **rule 3**

$$\text{FOLLOW}(A) = \text{FOLLOW}(B)$$

i.e. $\text{FOLLOW}(E') = \text{FOLLOW}(T)$

$$\text{FOLLOW}(E') = \{+, \epsilon\}$$

finally $\text{FOLLOW}(T) = \{+\} \cup \{), \epsilon\}$
 $= \{+,), \epsilon\}$

FOLLOW(T')

- (i) $T \rightarrow FT'$
- (ii) $T' \rightarrow *FT'$

$T \rightarrow FT'$ map it with $A \rightarrow \alpha B \beta$ then **rule 3**

$A = T$, $\alpha = F$, $B = T'$, $\beta = \epsilon$ then by **rule 3**

$$\text{FOLLOW}(A) = \text{FOLLOW}(B)$$

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

$T' \rightarrow *FT'$ we will map $A \rightarrow \alpha B\beta$
 $A = T'$, $\alpha = *F$, $B = T'$, $\beta = \epsilon$ then rule 3

$$\text{FOLLOW}(A) = \text{FOLLOW}(B)$$

$$\therefore \text{FOLLOW}(T') = \text{FOLLOW}(T)$$

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

FOLLOW(F)

$$\textcircled{1} \quad T \rightarrow FT' \\ T' \rightarrow *FT'$$

$T \rightarrow FT'$ we will map $A \rightarrow \alpha B\beta$
 $A = T$, $\alpha = \epsilon$, $B = F$, $\beta = T'$ then by rule 2

$$\text{FOLLOW}(B) = \{\text{FIRST}(T') - \epsilon\}$$

$$\therefore \text{FOLLOW}(F) = \text{FIRST}(T') - \epsilon \\ = \{*\}$$

$T' \rightarrow *FT'$ we will map $A \rightarrow \alpha B\beta$
 $A = T'$, $\alpha = *F$, $B = T'$, $\beta = \epsilon$ then by rule 3

$$\text{FOLLOW}(A) = \text{FOLLOW}(B)$$

$$\therefore \text{FOLLOW}(T') = \text{FOLLOW}(T)$$

$$= \{+,), \$\}$$

Finally $\text{FOLLOW}(F) = \{ \epsilon \} \cup \{ +,), \$ \}$
 $= \{ +, *,), \$ \}$

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

STEP 2:- Construction of predictive parsing table.

For the rule $A \rightarrow \alpha$ of grammar G.

① for each 'a' in $\text{FIRST}(\alpha)$ create entry
 $M[A, a] = A \rightarrow \alpha$
 where 'a' is terminal symbol.

② for ϵ in $\text{FIRST}(\alpha)$ create entry
 $M[A, \epsilon] = A \rightarrow \alpha$
 where 'b' is symbols from $\text{FOLLOW}(A)$

③ if ϵ is in $\text{FIRST}(\alpha)$ and '\$' is in $\text{FOLLOW}(A)$
 then create entry in the table
 $M[A, \$] = A \rightarrow \alpha$

4. All the remaining entries in the table M are marked as SYNTAX ERROR.

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

$$E \rightarrow TE'$$

$$A \rightarrow \alpha$$

$$A = E, \alpha = TE'$$

$$\begin{aligned} \text{FIRST}(TE') &= \{(, id), (\epsilon)\} \\ &= \{(, id) \} \end{aligned}$$

$$M[E, (] = E \rightarrow TE'$$

$$M[E, id] = E \rightarrow TE'$$

$$E' \rightarrow +TE'$$

$$A \rightarrow \alpha$$

$$A = E'$$

$$\alpha = +TE'$$

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

Hence $M[E', +] = E' \rightarrow +TE'$

$$E' \rightarrow \epsilon$$

$$A \rightarrow \alpha$$

$$A = E', \alpha = \epsilon$$

$$\text{FOLLOW}(E') = \{), \$\}$$

$$M[E',)] = E' \rightarrow \epsilon$$

$$M[E', \$] = E' \rightarrow \epsilon$$

$$T \rightarrow FT'$$

$$A \rightarrow \alpha$$

$$A = T, \alpha = FT'$$

$$\therefore T' = \epsilon$$

$$\text{FIRST}(FT') = \text{FIRST}(F) = \{ (, id \}$$

$$M[F, (] = T \rightarrow FT'$$

$$M[F, id] = T \rightarrow FT'$$

$$T' \rightarrow *FT'$$

$$A \rightarrow \alpha$$

$$A = T', \alpha = *FT'$$

$$\text{FIRST}(*FT') = \{ * \}$$

$$M[T', *] = T' \rightarrow *FT'$$

$$M[T', \epsilon] = T' \rightarrow *FT'$$

$$T' \rightarrow \epsilon$$

$$A \rightarrow \alpha$$

$$A = T', \alpha = \epsilon$$

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

$$M[T', +] = T' \rightarrow \epsilon$$

$$M[T',)] = T' \rightarrow \epsilon$$

$$M[T', \$] = T' \rightarrow \epsilon$$

$$F \rightarrow (E)$$

$$A \rightarrow \alpha$$

$$A = F, \alpha = (E)$$

$$\text{FIRST}((E)) = \{ (\}$$

$$M[F, (] = F \rightarrow (E)$$

$$F \rightarrow id$$

$$A \rightarrow \alpha$$

$$A = F, \alpha = id$$

$$\text{FIRST}(id) = \{ id \}$$

$$M[F, id] = F \rightarrow id$$

	id	+	*	(	)	\$
E	$E \rightarrow TE'$	error	error	$E \rightarrow TE'$	error	error
E'	error	$E' \rightarrow +TE'$	error	error	$E' \rightarrow \epsilon$	$E' \rightarrow \epsilon$
T	$T \rightarrow FT'$	error	error	$T \rightarrow FT'$	error	error
T'	error	$T' \rightarrow \epsilon$	$T' \rightarrow *FT'$	error	$T' \rightarrow \epsilon$	$T' \rightarrow \epsilon$
F	$F \rightarrow id$	error	error	$F \rightarrow (E)$	error	error

STEP 3:- parsing the g/p string

g/p:- id + id * id

$E \rightarrow TE'$
 $E' \rightarrow +TE' | \epsilon$
 $T \rightarrow FT'$
 $T' \rightarrow *FT' | \epsilon$
 $F \rightarrow (E) | id$

Stack	g/p	Action
\$ E	id + id * id \$	E
\$ E' T	id + id * id \$	$E \rightarrow TE'$
\$ E' T' F	id + id * id \$	$T \rightarrow FT'$
\$ E' T' id	id + id * id \$	$F \rightarrow id$
\$ E' T' +	id + id * id \$	
\$ E' T'	+ id * id \$	$T' \rightarrow \epsilon$
\$ E'	+ id * id \$	$E' \rightarrow +TE'$
\$ E' T +	+ id * id \$	
\$ E' T	id * id \$	
\$ E' T' F	id * id \$	$T \rightarrow FT'$
\$ E' T' id	id * id \$	$F \rightarrow id$
\$ E' T' +	* id \$	
\$ E' T' F +	* id \$	$T' \rightarrow *FT'$
\$ E' T' F	id \$	
\$ E' T' id	id \$	$F \rightarrow id$
\$ E' T'	\$	
\$ E'	\$	$T' \rightarrow \epsilon$
\$	\$	$E' \rightarrow \epsilon$

BOTTOM-UP PARSING

→ the parse tree is constructed from bottom to up that is from leaves to root.

→ In this process, the g/p symbols are placed at the leaf nodes after successful parsing.

→ In this process, basically parser tries to identify R.H.S of production rule & replace it by corresponding L.H.S.
this activity is called **REDUCTION**

eg:-

Consider Grammar

$S \rightarrow TL;$

$T \rightarrow \text{int} \mid \text{float}$

$L \rightarrow L, \text{id} \mid \text{id}$

the g/p string is float, id, id;

parse tree

starting from leaf node

step 1:-

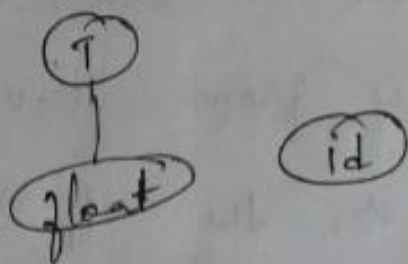
float

Step 2:- reduce float to T
 $T \rightarrow \text{float}$



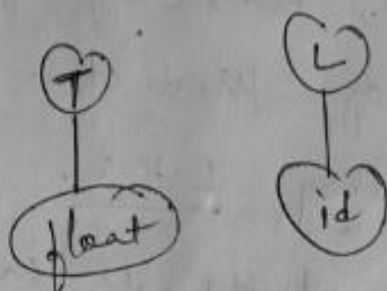
Step 3

read next string from i/p



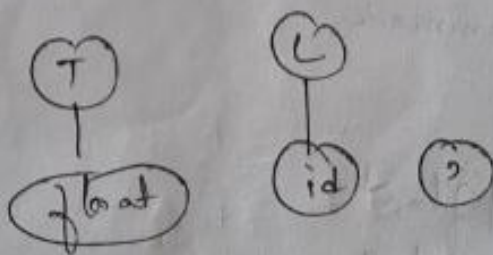
Step 4:-

Reduce id to L
 $L \rightarrow id$



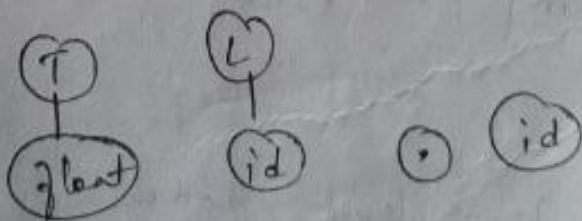
Step 5:-

read next i/p



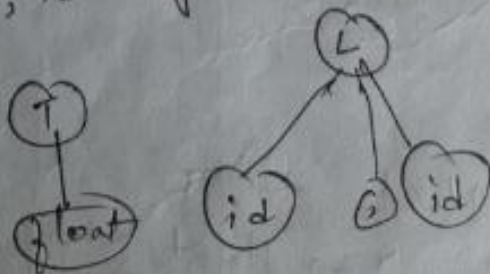
Step 6:-

read next string from i/p

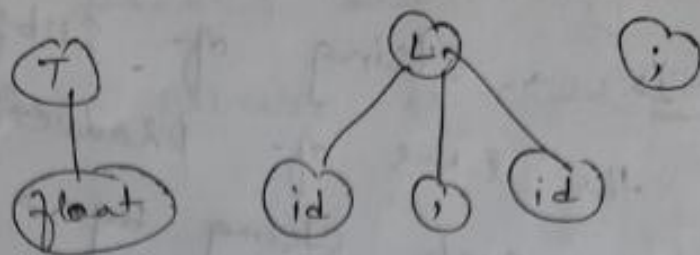


Step 7:-

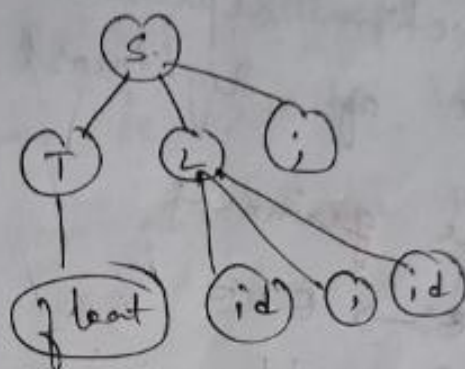
id, id gets reduced to L



step 8:- read next string ;



step 9:- TL; reduced to S



Step 10:- the sentential form produced while constructing this parse tree is

float id, id;

T id, id;

T L id;

TL;

S

HANDLE PRUNING

- Handle is a string of substring that matches the R.H.S of production and we can reduce such string by a non-terminal on L.H.S of production.
- Such reduction represents one step along the reverse of rightmost derivation.

Consider the grammar

$$E \rightarrow E + E$$

$$E \rightarrow id$$

Now the string $id + id + id$ is rightmost derivation is

$$E \rightarrow E + E$$

$$E \rightarrow E + E + E$$

$$E \rightarrow E + E + id$$

$$E \rightarrow E + id + id$$

$$E \rightarrow id + id + id$$

The Bold strings are called HANDLES

Right Sentential form	Handle	production
$id + id + id$	id	$E \rightarrow id$
$E + id + id$	id	$E \rightarrow id$
$E + E + id$	id	$E \rightarrow id$
$E + E + E$	$E + E$	$E \rightarrow E + E$
$E + E$	$E + E$	$E \rightarrow E + E$
E		

Thus Bottom parser is essentially a process of detecting handles and using them in Reduction. This process is called HANDLE PRUNING.

SHIFT REDUCE PARSER

→ Shift reduce parser attempts to construct parse tree from leaves to root.

→ Thus it works on the same principle of bottom-up parser.

→ A shift reduce parser requires following Data Structures.

- ① the I/P buffer storing the I/P string.
- ② A stack for storing & Accessing the L.H.S & R.H.S of rules.

→ the parser performs following basic operations.

① SHIFT: moving of the symbols from I/P buffer onto the stack.

② REDUCE: If the handle appears on the top of the stack then reduction of it by appropriate rule is done.

that means R.H.S of rule is popped of and L.H.S is pushed in.

③ ACCEPT :- If the stack contains START symbol only and glp buffer is EMPTY at the same time then that action is called ACCEPT.

④ ERROR :- A situation in which parser cannot either shift or reduce the symbols, it cannot even perform the accept action is called as ERROR.

Ex:-

$$E \rightarrow E - E$$

$$E \rightarrow E * E$$

$$E \rightarrow id$$

perform shift-reduce parsing of glp string

$$id \leftarrow id * id$$

Stack	glp buffer	Passing Action
\$	id - id * id	shift id
\$id	- id * id	Reduce $E \rightarrow id$
\$E	- id * id	Shift -
\$E -	id * id	Shift id
\$E - id	* id	Reduce $E \rightarrow id$
\$E - E	* id	Shift *
\$E - E *	id	Shift id
\$E - E * id	\$	Reduce $E \rightarrow id$
\$E - E * E	\$	Reduce $E \rightarrow E * E$

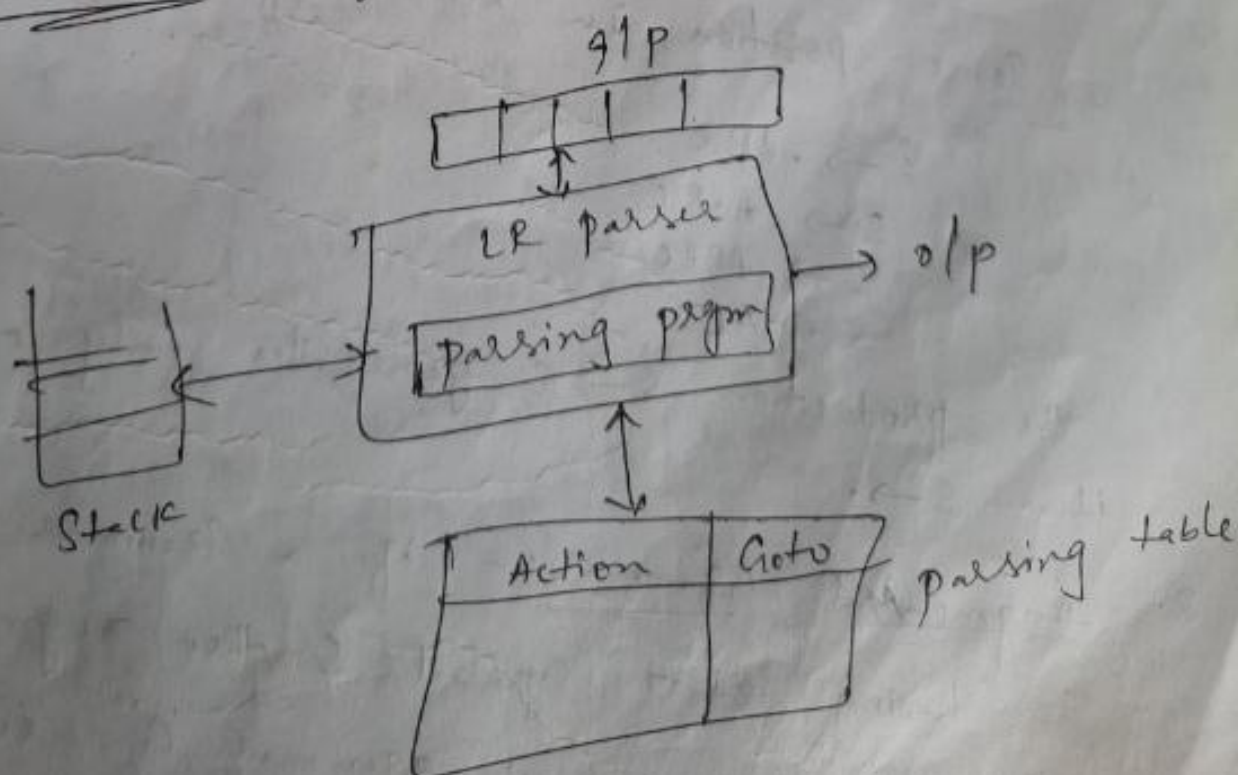
\$ E - E	\$	Reduce $E \rightarrow E - E$
\$ E	\$	Accept

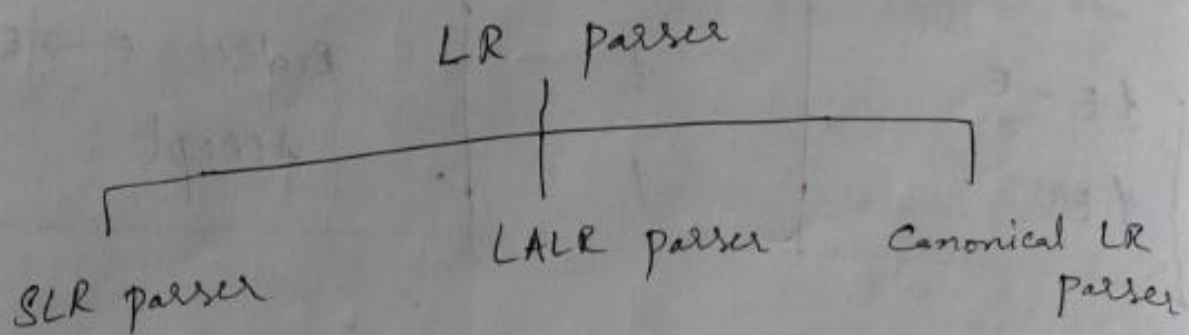
LR PARSERS

- This is the most efficient method of bottom-up parsing which can be used to parse the large class of CFG.
- This method is also called as LR(K) parsing

LR(K) $\rightarrow$ k omitted, 1 placed.
 $\hookrightarrow$ o/p scan left to right
 $\hookrightarrow$ right most derivation

Structures of LR parser





SIMPLE LR Parsing SLR

Steps.

1. Constuction of canonical set of items.
2. Constuction of SLR parsing table.
3. parsing the i/p string.

Defination of LR(0) items & related terms:-

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

$$\begin{aligned}
 S &\rightarrow \cdot ABC \\
 S &\rightarrow A \cdot BC \\
 S &\rightarrow AB \cdot C \\
 S &\rightarrow ABC \cdot
 \end{aligned}$$

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

2. Augmented Grammar:- if a Grammar G is having start symbol S then augmented grammar is a new grammar G' in which S' is a new start symbol such

that $S' \rightarrow S$.

The purpose of this grammar is to indicate the acceptance of γ/p .

that is when parser is about to reduce $S' \rightarrow S$ it reaches to Acceptance state.

3. Kernel items:- It is a collection of items $S' \rightarrow S$ and all the items whose dots γ not at the leftmost end of R.H.S of rule.

Non-Kernel items:- the collection of all the items in which γ at the left end of R.H.S of the rule.

$S \rightarrow A \cdot BC$

4. Closure & goto fn:- these γ two imp fns required to create collection of Canonical set of items.

5. Viable prefix:- 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.

Closure operation :-

For a CFG G , if I is the set of items then the $\text{closure}(I)$ can be constructed using following rules.

① Consider I is a set of Canonical items & initially every item I is added to $\text{closure}(I)$.

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

$\text{closure}(I)$:

$A \rightarrow \alpha \cdot B\beta$
 $B \rightarrow \gamma$

- This rule has to be applied until no more new items can be added to $\text{closure}(I)$.

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

- A non-terminal immediately to the right of $\cdot$ indicates that it has to be expanded shortly.

$\text{goto}(\bar{I}, b) = \text{closure}$
 $(S \rightarrow bAk, S \rightarrow bBk)$
 (A)

GOTO operation :-

the function goto can be define as follows.

If there is a production $A \rightarrow \alpha \cdot B \beta$ then goto $A \rightarrow \alpha B \cdot \beta$.

that means simply shifting of $\cdot$ one position ahead over the grammar symbol

(may be terminal / non-terminal).

goto can be written as $\text{goto}(I, B)$.

Ex: Consider the following Grammar;

$$S \rightarrow Aa \mid bAc \mid Bc \mid bBc$$

$$A \rightarrow d$$

$$B \rightarrow d$$

compute closure(I) & goto(I).

Sol

closure(I)

$$\begin{aligned} I_0: \quad & S \rightarrow \cdot Aa \\ & S \rightarrow \cdot bAc \\ & S \rightarrow \cdot Bc \\ & S \rightarrow \cdot bBc \\ & A \rightarrow \cdot d \\ & B \rightarrow \cdot d \end{aligned}$$

Now applying goto on I_0

$$\begin{aligned} I_1: \quad & \text{goto}(I_0, A) \\ & S \rightarrow A \cdot a \end{aligned}$$

$$\text{goto}(I_0, b)$$

$I_2:$

$$\begin{aligned} & S \rightarrow b \cdot A c \\ & S \rightarrow b \cdot \underline{B} c \\ & A \rightarrow \cdot d \\ & B \rightarrow \cdot d \end{aligned}$$

after non-terminal so need to apply

$$\text{goto}(I_2, B)$$

$I_3:$

$$\begin{aligned} & S \rightarrow B \cdot c \\ & \text{goto}(I_2, d) \end{aligned}$$

$I_4:$

$$\begin{aligned} & A \rightarrow d \cdot \\ & B \rightarrow d \cdot \end{aligned}$$

Prob SLR

$$E \rightarrow E + T$$

$$E \rightarrow T$$

$$T \rightarrow T * F$$

$$T \rightarrow F$$

$$F \rightarrow (E)$$

$$F \rightarrow id$$

Step 1:-

Construction of Canonical collection of set of item:

①. For the grammar G initially add $S' \rightarrow S$ in the set of item C .

②. For each set of items I_i in C and for each grammar symbol x (may be terminal or non-terminal) add closure (I_i, x) . This process should be repeated by applying $goto(I_i, x)$ for each x in I_i such that $goto(I_i, x)$ is not empty and not in C .
The set of items has to be constructed until no more set of items can be added to C .

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

I_0 :

$E' \rightarrow \cdot E$
 $E \rightarrow \cdot E + T$
 $E \rightarrow \cdot T$
 $T \rightarrow \cdot T * F$
 $T \rightarrow \cdot F$
 $F \rightarrow \cdot (E)$
 $F \rightarrow \cdot id$

Applying goto on I_0 .

goto(I_0, E)

I_1 : $E' \rightarrow E \cdot$
 $E \rightarrow E \cdot + T$

goto(I_0, T)

I_2 : $E \rightarrow T \cdot$
 $T \rightarrow T \cdot * F$

goto(I_0, F)

I_3 : $T \rightarrow F \cdot$

goto($I_0, ($)

I_4 : $F \rightarrow (\cdot E)$

goto(I_0, id)

I_5 : $F \rightarrow id \cdot$

Applying goto on I_0 is completed.

→ Now applying goto on I_1

In I_1 , there are two productions

$$E' \rightarrow E.$$

$$E \rightarrow E + T.$$

1st production is kernel item so we can't apply goto on it.

2nd production applying goto

goto (I_1 , +)

I_6 :

$$E \rightarrow E + \underline{T}$$
$$T \rightarrow \underline{T} * F$$
$$T \rightarrow \underline{F}$$
$$F \rightarrow \underline{(E)}$$
$$F \rightarrow \underline{id}$$

→ after non-term
So, need to expand
it

→ Now applying goto on I_2 .

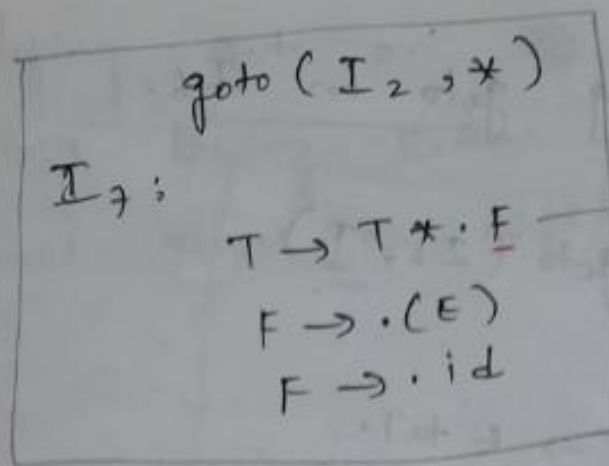
In I_2 there are two productions

$$E \rightarrow T.$$

$$T \rightarrow T * F$$

1st production is kernel item so we can't apply goto on it.

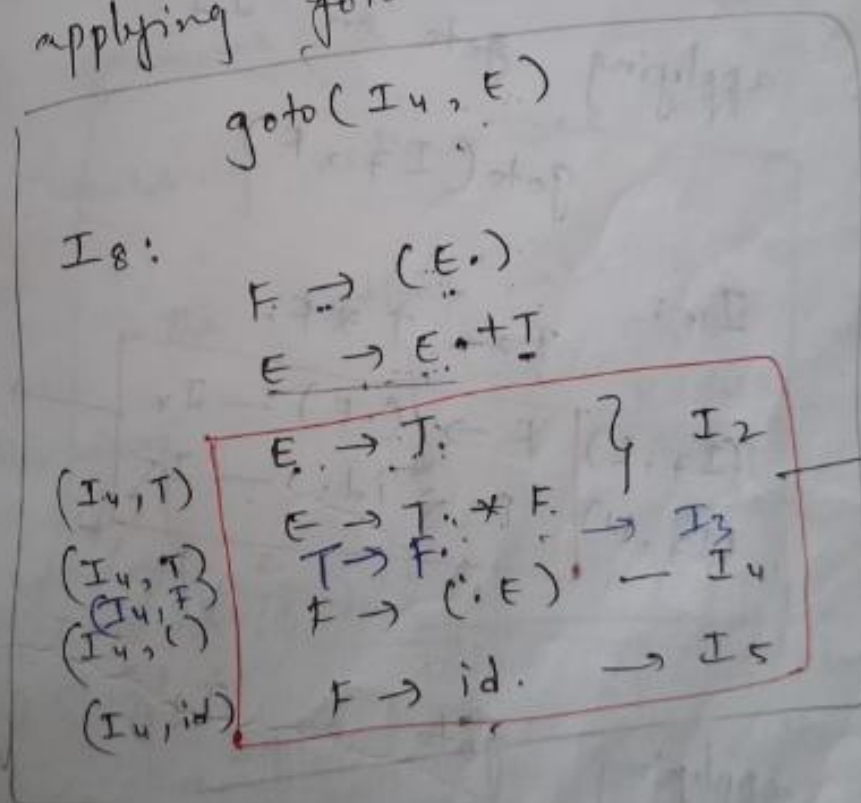
2nd production we can apply goto



After, non-termi
So, need to
Xpand it

→ Now applying goto on I_3 .
the production in I_3 is a kernel item
so, we can't apply goto.

→ Now applying goto on I_4 .



Repetition
of states

→ Now applying goto on I_5 .
the production in I_5 is a kernel
item. So, no need to apply goto.

→ Now applying goto on I_6 .

$\text{goto}(I_6, \underline{T})$

I_9 :

$E \rightarrow E + T.$

$T \rightarrow T * F$ ✓

(I_6, F)	$T \rightarrow \underline{F}.$	I_3
$(I_6, ($	$F \rightarrow (.E.)$	I_4
(I_6, id)	$F \rightarrow id.$	I_5

→ repetition of states.

→ Now applying goto on I_7

$\text{goto}(I_7, F)$

I_{10} :

$T \rightarrow T * F.$

$(I_7, ($	$F \rightarrow (.E.)$	I_4
(I_7, id)	$F \rightarrow id.$	I_5

→ repetition of states

→ Now applying goto on I_8

$\text{goto}(I_8,)$

I_{11} :

$F \rightarrow (E).,$

$(I_8, +)$	$E \rightarrow E + .T$	I_6
------------	------------------------	-------

→ repetition of state

$\text{goto on } I_9$

$\text{goto}(I_9, F)$

I_{12} $T \rightarrow T * F$ →

→ Applying goto on I_0, I_{10}, I_{11} is of no use becoz either the items are kernel items or repetition of states.

STEP 2:-

Construction of SLR parsing table

①. Initially construct set of items $C = \{I_0, I_1, I_2, \dots, I_n\}$ where C is a collection of set of LR(0) items for the grammar G .

②. The parsing actions are based on each item I_i .
The Actions are

Shift
Reduce
Accept
~~Accept~~

$T \rightarrow F. - \bar{I}_3$ some answers
 $Follow(F) = \{+, *, \$\}$

$action[\bar{I}_3 +] = action[\bar{I}_3 *] = action[\bar{I}_3 \$] = r_2$

② Shift :-

If $A \rightarrow \alpha. a\beta$ is in I_i and
set $action[i, a]$ =
 $goto(I_i, a) = I_j$ then
"Shift j".

Note that 'a' must be a terminal symbol.

(b) Reduce :-

If there is a rule $A \rightarrow \alpha$ is in I_i then set action $[i, a] = \text{reduce } A \rightarrow \alpha$ for all symbols a ,

where $a \in \text{FOLLOW}(A)$.

Note that A must not be an augmented grammar S' .

(c) Accept :-

If $S' \rightarrow S_0$ is in I_i then the entry in the action table $\text{action}[i, \$] = \text{"accept"}$

4. The Goto part of the SLR table can be filled as:

the goto transitions for state i is considered for non-terminals only.

if $\begin{cases} \text{goto}(I_i, A) = I_j \\ \text{goto}[I_i, A] = j \end{cases}$ then

5. All the entries not defined by rule 2 & 3 are considered to be as "error".

1) $E \rightarrow E + T$

2) $E \rightarrow T$

3) $T \rightarrow T * F$

4) $T \rightarrow F$

5) $F \rightarrow (E)$

6) $F \rightarrow id$

S, R, A

State	Action						Goto		
	id	+	*	(	)	\$	E	T	F
State 0	S5			S4			1	2	3
1		S6				Accept			
2		δ_2 S7			δ_2	δ_2			
3		δ_4	δ_4		δ_4	δ_4			
4	S5			S4			8	2	3
5		δ_6	δ_6		δ_6	δ_6			
6	S5			S4				9	3
7	S5			S4					10
8		S6			S11				
9		δ_1 S7			δ_1	δ_1			
10		δ_3	δ_3		δ_3	δ_3			
11		δ_5	δ_5		δ_5	δ_5			

Shift
 $I_0, id \rightarrow I_5 - \text{shift } 5$
 $I_0, (\rightarrow I_4 - S4$

$T \rightarrow F$
 $\text{Follow}(T) = \{+, *,), \$\}$

$\text{Follow}(F) = \{+, *,), \$\}$

$\text{goto}(I, id) = I_5$ $\text{goto}(I_0, E)$
 $= I_1$

Reduce

kernel item

$A \rightarrow \alpha$

$\downarrow$

$\text{Follow}(A)$

$A \rightarrow \alpha$
 (line no. of grammar)

I_2 :

$E \rightarrow T$

$E \rightarrow T$

$\text{Follow}(E) = \{+,), \$\}$

2 line of grammar

$\downarrow$
 $S_0 \delta_2$

Accept.

$S' \rightarrow S.$
 $I_1: E' \rightarrow E.$

$[I_1, \$] = \text{Accept}$

Auto.

$(I_0, E) \rightarrow I_1$
 $(I_0, T) \rightarrow I_2$
 $(I_0, F) \rightarrow I_3$
 $(I_4, E) \rightarrow I_8$
 $(I_4, T) \rightarrow I_2$
 $(I_4, F) \rightarrow I_3$
 $(I_6, T) \rightarrow 9$
 $(I_6, F) \rightarrow 3$
 $(I_7, F) \rightarrow 10$

STEP 3 parsing the 9/p string

9/p string: $id * id + id$

Stack	9/p buffer	parsing Action
\$0	$id * id + id$	shift id
\$0id5	$* id + id$	reduce $F \rightarrow id$
\$0F3	$* id + id$	reduce $T \rightarrow F$
\$0T2	$* id + id$	shift id *
\$0T2*7	$id + id$	shift id
\$0T2*7id5	$+ id$	reduce $F \rightarrow id$
\$0T2*7F10	$+ id$	reduce $T \rightarrow T * F$
\$0T2	$+ id$	reduce $E \rightarrow T$
\$0T2	$+ id$	shift +
\$0E1	id	shift id
\$0E1+6		

$\$0E1+6id5$	$\$$	reduce $F \rightarrow id$
$\$0E1+6F3$	$\$$	reduce $T \rightarrow F$
$\$0E1+6T9$	$\$$	reduce $E \rightarrow E + T$
$\$0E1$	$\$$	accept.

$[0, id] = 55 \rightarrow 5$ in parsing string

$$[0, F] = 3$$

$$[0, T] = 2$$

$$[2, *] = 57 \rightarrow 7$$

$$[1, +] = 56 - 6$$

$$[6, id] = 55 - 5$$

$$[7, id] = 55 - 5$$

$$[7, F] = 10$$

$$[0, E] = 1$$

$$[6, F] = 3$$

$$[6, T] = 9$$

LR(K) parser / SIMPLY LR / CANONICAL LR

→ The Canonical Set of items in the parsing technique in which a lookahead symbol is generated while constructing set of items.

→ Hence the collection of set of items is referred as LR(1).
↓
1 lookahead

Steps

1. Construction of canonical set of items along with the lookahead.
2. Building canonical LR parsing table.
3. parsing the g/p string.

STEP 1:-

1. For the grammar G initially add $S' \rightarrow \cdot S$ in the set of item C .
2. Same as SLR 2nd step.
3. The closure for can be computed as follows.

$$A \rightarrow \alpha \cdot X \beta, a$$

$$X \rightarrow \gamma, b \in \text{FIRST}(\beta, a)$$

$$x \rightarrow \cdot y, b$$

4. goto Same as SLR. $x \rightarrow \cdot y \leftarrow$

this process is repeated until no more set of items can be added to the collection C.

Ex:-

$$S' \rightarrow S$$

$$S \rightarrow cc$$

$$c \rightarrow ac \mid d$$

Construct LR(1) set of items for above grammar.

Sol

we initially add $S' \rightarrow \cdot S, \$$ as the first item in I_0 .

$$\text{Now match } S' \rightarrow \cdot S, \$$$

$$A \rightarrow \alpha \cdot x \beta, a$$

$$A = S', \alpha = \epsilon, x = S, \beta = \epsilon, a = \$$$

$$x \rightarrow y, b \in \text{FIRST}(\beta, a)$$

$$S \rightarrow \cdot cc, b \in \text{FIRST}(\epsilon \$)$$

$$b = \$$$

$\therefore S \rightarrow \cdot cc, \$$ will be added I_0

Now $S \rightarrow .cc$, $\$$ is in I_0 . it match

$$A \rightarrow \alpha \cdot x \beta, a$$

$$A = S, \alpha = \epsilon, x = c, \beta = c, a = \$$$

$$X \rightarrow \gamma, b \in \text{FIRST}(\beta, a)$$

$$c \rightarrow .ac, b \in \text{FIRST}(c, a)$$

$$c \rightarrow .d$$

$$(a|d|\$)$$

$\$$ not added cuz
it added only to
starting symbol

$$\therefore \left. \begin{array}{l} c \rightarrow .ac, a|d \\ c \rightarrow .d, a|d \end{array} \right\} \text{ added in } I_0$$

Hence I_0 :

$$S' \rightarrow .S, \$$$

$$S \rightarrow .cc, \$$$

$$c \rightarrow .ac, a|d$$

$$c \rightarrow .d, a|d$$

→ Applying goto on I_0

$$\text{goto}(I_0, S)$$

I_1 :

$$S' \rightarrow S, \$$$

$$\text{goto}(I_0, c)$$

I_2 :

$$S \rightarrow c.c, \$$$

$$A \rightarrow \alpha \cdot x \beta, a$$

$$A = S, \alpha = c, x = c, \beta = \epsilon, a = \$$$

$x \rightarrow \gamma$, $b \in \text{FIRST}(\beta, a)$
 $c \rightarrow \cdot ac$, $b \in \text{FIRST}(\epsilon, \$)$
 $c \rightarrow \cdot d$ = $\$$: I

$\text{goto}(I_0, c)$
 $I_2:$
 $S \rightarrow c \cdot c, \$$
 $c \rightarrow \cdot ac, \$$
 $c \rightarrow \cdot d, \$$

$\text{goto}(I_0, a)$
 $I_3:$
 $c \rightarrow a \cdot c, a|d$
 $c \rightarrow \cdot ac, a|d$
 $c \rightarrow \cdot d, a|d$

$\text{goto}(I_0, d)$
 $I_4: c \rightarrow d \cdot, a|d$

~~$\text{goto}(I_0, a)$
 $c \rightarrow \cdot ac, a|d$
 $A \rightarrow x \cdot xB, a$
 $x \rightarrow \gamma$
 $a \rightarrow ?$
 So rule is not satisfying for
 $c \rightarrow \cdot ac, a|d$
 $c \rightarrow \cdot d, a|d$~~

$\rightarrow$ Applying goto on I_1 ,
 the production is a kernel item
 so γ can't apply goto on I_1 .

$\rightarrow$ Applying goto on I_2

$\text{goto}(I_2, c)$
 $I_5: S \rightarrow cc \cdot, \$$

goto(I_2, a)

I_6 :

$c \rightarrow a.c, \$$

$c \rightarrow .ac, \$$

$c \rightarrow .d, \$$

goto(I_2, d)

I_7 :

$c \rightarrow d., \$$

→ Applying goto on I_3

goto(I_3, c)

I_8 : $c \rightarrow ac., a|d$

goto(I_3, a)

$c \rightarrow a.c, a|d$ — I_3

goto(I_3, d)

$c \rightarrow d., a|d$ — I_4

→ repetition of states

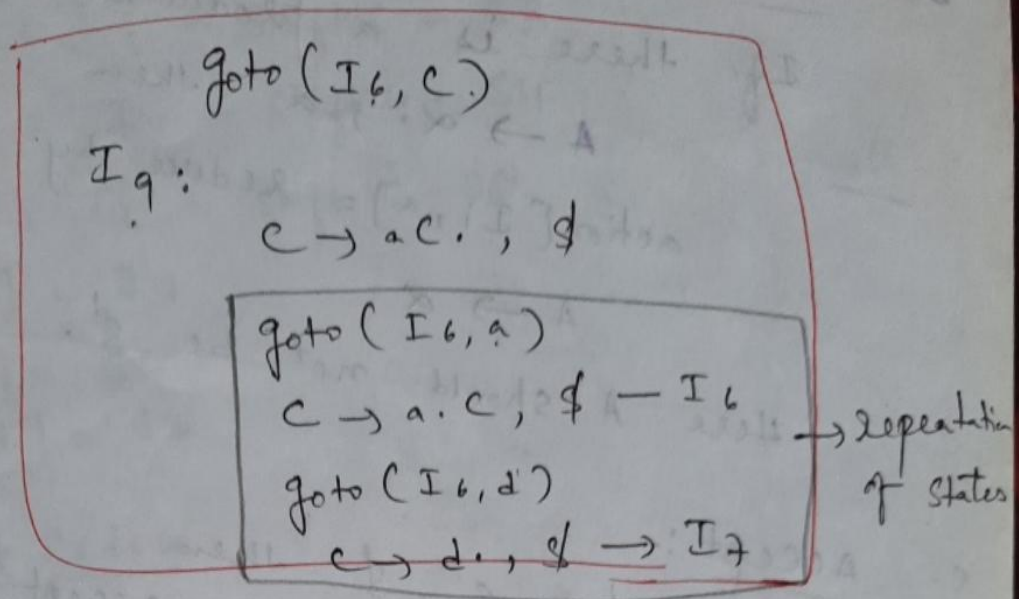
→ Applying goto on I_4

production is kernel item so goto can't be applicable.

→ Applying goto on I_5 .

It's a kernel item. so can't apply goto

→ Applying goto on I_6



→ I_7 kernel item

I_8 kernel item & repetition states

→ I_9 kernel item & repetition states

So can't apply goto.

States are from $I_0 - I_9$.

STEP 2:-

1. Shift is same as SLR step

2. The parsing actions are
Shift
Reduce
Accept

② Shift:- If $[A \rightarrow \alpha.a\beta, b]$ is in I_i and $\text{goto}(I_i, a) = I_j$ then create an entry in the action table
 $\text{action}[I_i, a] = \text{shift } j$.

b. Reduce :-

If there is a production
 $A \rightarrow \alpha$, a then
 $\text{action}[I_i, a] = \text{reduce by}$

$$A \rightarrow \alpha$$

Here A should not be S' .

c. Accept :-

$S' \rightarrow S$, \$ then
 $\text{action}[I_i, \$] = \text{Accept}$

3. goto part is same as SLR

4. Same as SLR

Grammar

- ① $S \rightarrow ccc$
- ② $c \rightarrow ac$
- ③ $c \rightarrow d$

States	ACTION			GOTO	
	a	d	\$	S	C
0	S3	S4		1	2
1			Accept		
2	S6	S7			5
3	S3	S4			8
4	r3	r3			
5			r1		
6	S6	S7			9
7			r3		
8	r2	r2			
9			r2		

shift

$$[I_0, a] = I_3 - \text{shift } 3 - S_3$$

$$[I_0, d] = I_4 - \text{shift } 4 - S_4$$

$$[I_2, a] = I_6 - \text{shift } 6 - S_6$$

$$[I_2, d] = I_7$$

$$[I_6, a] = S_6$$

$$[I_6, d] = S_7$$

Accept Reduce

$$A \rightarrow \alpha., a$$

$\downarrow$

$$A \rightarrow \alpha$$

line no. of grammar

$$I_4: c \rightarrow d., a, d - \text{line } 3$$

$$I_5: s \rightarrow cc., \$ - \text{line } 1$$

$$I_8: e \rightarrow ac., a, d - \text{line } 2$$

$$I_9: c \rightarrow ac., \$ - \text{line } 2$$

$$I_7: c \rightarrow d., \$ - \text{line } 3$$

Accept

$$I_1: s' \rightarrow s., \$ - \text{accept}$$

goto

$$[I_0, s] = I_1 - 1$$

$$[I_0, c] = I_2 - 2$$

$$[I_2, c] = I_5 - 5$$

$$[I_3, c] = I_8 - 8$$

$$[I_6, c] = I_9 - 9$$

STEP 3:- LR(0) Gram & DPDA
 parsing the slp string

slp:- "a add"
 referring the parsing table to parse
 the slp string.

Stack	slp buffer	parsing Action
\$0	a add \$	shift a
\$0a3	add \$	shift a
\$0a3a3	dd \$	shift d
\$0a3a3d4	d \$	shift reduce $c \rightarrow d$
\$0a3a3c8	d \$	reduce $c \rightarrow ac$
\$0a3c8	d \$	reduce $c \rightarrow ac$
\$0c2	d \$	shift d
\$0c2d7	\$	reduce $c \rightarrow d$
\$0c2c5	\$	reduce $s \rightarrow cc$
\$0st	\$	Accept

thus the given string is successfully
 parsed using LR parser (or) Canonical
 LR parser.

LALR parser

→ In this type of parser the lookahead symbol is generated for each set of item.

→ the table obtained by this method are smaller in size than LR(K) parser.

→ In fact the states of CLR and LALR parsing are always same.

STEPS:-

1. Construction of canonical set of items along with the lookahead.
2. Building LALR parsing table.
3. parsing the g/p string using canonical LR parsing table.

STEP 1:-

→ Construction of canonical set of items will be as the LR(1) or CANONICAL LR.

→ But the only difference is that:
In construction of LR(1) items for LR parser, we have differed the two states if SECOND component is different

but in this case, we will MERGE the two states by merging of 1st & 2nd components from both the states.

Ex:-

→ Consider the SAME grammar as LR(1) or Canonical

→ Consider the SAME Canonical set of items as LR(1) / Canonical

$S \rightarrow CC$

$C \rightarrow aC$

$C \rightarrow d$

$I_0:$

$S' \rightarrow .S, \$$

$S \rightarrow .CC, \$$

$C \rightarrow .aC, ald$

$C \rightarrow .d, ald$

$I_3: \text{goto}(I_0, a)$

$C \rightarrow a.C, ald$

$C \rightarrow .aC, ald$

$C \rightarrow .d, ald$

$I_1: \text{goto}(I_0, S)$

$S' \rightarrow S., \$$

$I_4: \text{goto}(I_0, d)$

$C \rightarrow d., ald$

$I_2: \text{goto}(I_0, C)$

$S' \rightarrow C.C, \$$

$C \rightarrow .aC, \$$

$C \rightarrow .d, \$$

$I_5: \text{goto}(I_2, C)$

$S \rightarrow CC., \$$

$I_6: \text{goto}(I_2, a)$

$c \rightarrow a.c, \$$

$c \rightarrow .ac, \$$

$c \rightarrow .d, \$$

$I_8: \text{goto}(I_3, c)$

$c \rightarrow ac, ald$

$I_9: \text{goto}(I_6, c)$

$c \rightarrow ac, \$$

$I_7: \text{goto}(I_2, d)$

$c \rightarrow d, \$$

Now we will merge 3, 6

4, 7

8, 9

Because they are having 1st component same.

$I_{36}: \text{goto}(I_0, a)$

$c \rightarrow a.c, ald|\$$

$c \rightarrow .ac, ald|\$$

$c \rightarrow .d, ald, \$$

$I_{47}: \text{goto}(I_0, d)$

$c \rightarrow d, ald|\$$

$I_{89}: \text{goto}(I_3, c)$

$c \rightarrow ac, ald|\$$

STEP 2:- Construction of parsing table

Step 1:- Construct the LR(1) set of items.

Step 2:- Merge the two states I_i & I_j if the 1st component are matching & create a new state replacing one of the older state such as
$$I_{ij} = I_i \cup I_j$$

Step 3:- Action table SAME as LR(1) of CANONICAL LR

Step 4:- GOTO table SAME as LR(1) / CANONICAL LR

Step 5:- If the parsing action CONFLICT then the algorithm fails to produce LALR parser & grammar is not LALR(1).

All the entries not defined by rule 3 & rule 4 are considered to be "Error".

States	Action			Goto	
	a	d	\$	S	C
0	S36	S47		1	2
1			accept		
2	S36	S47			5
36	S36	S47			89
47	r3	r3	r3		
5			r1		
89	r2	r2	r2		

Shift

$[I_0, a] - \text{shift } 36 - S36$

$[I_0, d] - S47$

$[I_2, a] - S36$

$[I_2, d] - S47$

$[I_3, a] \} - S36$

$[I_6, a] \}$

$[I_3, d] \} - S47$

$[I_6, d] \}$

REDUCE

$[I_4, a] \}$

$[I_7, a] \}$

$A \rightarrow \alpha$

$C \rightarrow d, a, d - \text{line } 3$

$[I_5, \$] - s \rightarrow cc., \$ - \text{line } 1 - 2,$

$\left. \begin{array}{l} [I_8, a|d|\$] \\ [I_9, a|d|\$] \end{array} \right\} - c \rightarrow ac., a|d|\$ - \text{line } 2 - 22$

ACCEPT:

$[I_1, \$] s' \rightarrow s., \$ - \text{Accept}$

GOTO

$[I_0, S] - I_1 - 1$

$[I_0, C] - I_2 - 2$

$[I_2, C] - I_5 - 5$

$\left. \begin{array}{l} [I_3, C] \\ [I_6, C] \end{array} \right\} - \left. \begin{array}{l} I_8 \\ I_9 \end{array} \right\} - 89$

LR(1) / LR(0) G

STEP 3:-

parsing the g/p string "aadd"

stack	g/p buffer	parsing action
\$0	aadd\$	shift a
\$0a36	add\$	shift a
\$0a36a36	dd\$	shift d
\$0a36a36a36	d\$	reduce $c \rightarrow d$
\$0a36a36d47	d\$	reduce $c \rightarrow ac$
\$0a36a36c89	d\$	reduce $c \rightarrow ac$
\$0a36c89	d\$	shift d
\$0c2	\$	reduce $c \rightarrow d$
\$0c2d47	\$	reduce $s \rightarrow cc$
\$0c2c5	\$	Accept
\$0s1		

Thus the LALR & LR parser will mimic one another on the same g/p.

HANDLING

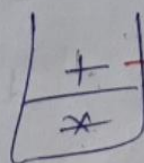
AMBIGUOUS

GRAMMAR

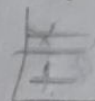
→ Ambiguous Grammar creates CONFLICTS
& we cannot parse the s/p string.

→ Conflict can be avoided by using
PRECEDENCE & ASSOCIATIVITY.

{ * , + }



error



* is having priority &
+ is having less priority. So, we
can't push + onto the stack.

→ if S/R conflict then it can be
avoided by using SHIFT only.

→ if R/R conflict then it can be
avoided by using first Reduce.

AUTOMATIC PARSER GENERATOR / YACC

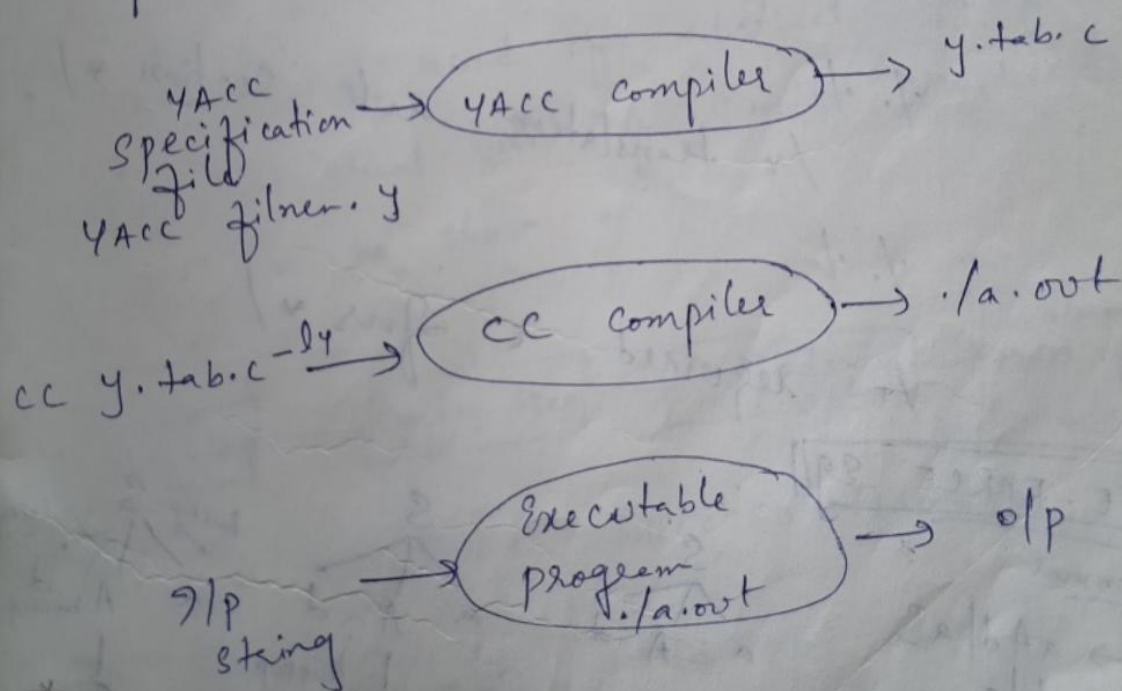
→ YACC stands for YET ANOTHER COMPILER COMPILER

which is basically the utility available from UNIX.

→ Basically YACC is LALR parser generator.

→ YACC can report Conflicts or ambiguities (if at all) in the form of error messages.

The typical YACC translator can be represented as:-



YACC: PARSER GENERATOR MODEL

YACC Specification

→ YACC specification file consists of three parts

Declarative Section
TRANSLATION Rule
c fns

%. {

/* declaration section */

%. }

%. %.

/* translation rule section */

%. %.

/* required c fns */

BRUTE-FORCE eg

Grammar

$S \rightarrow aAd | aB$

$A \rightarrow b | c$

$B \rightarrow ccd | ddc$

9/p:- accd

