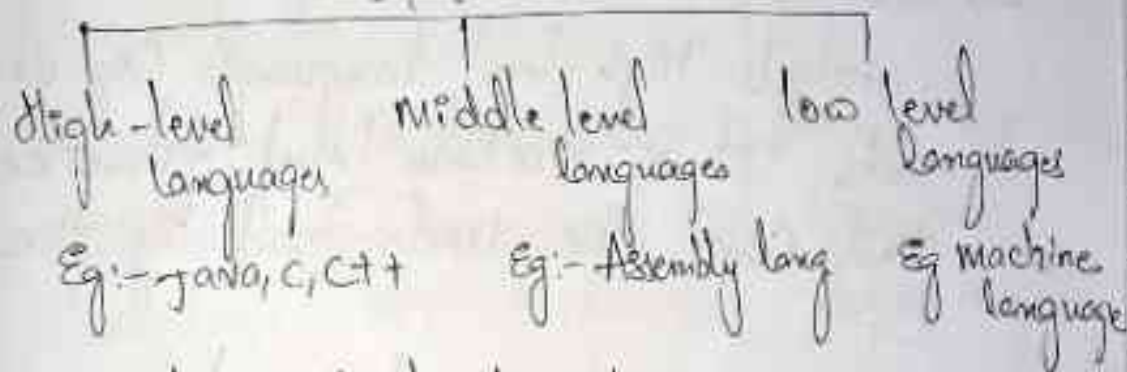


8/04/2021

Compiler Design

*Introduction:-

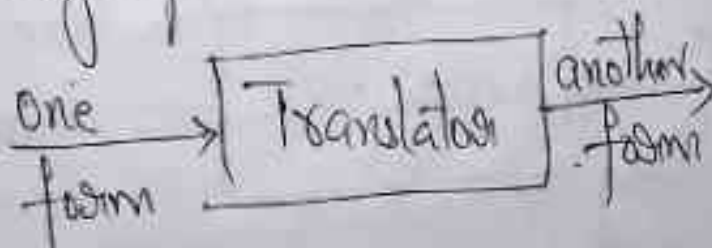
→ The programming languages are divided into three categories languages



→ Computer understands only 0's & 1's (machine language). Therefore, translators are used to convert one language to another language.

→ Translator:-

> A Translator is a programming language processor that converts a computer program from one language to another.



Examples:-

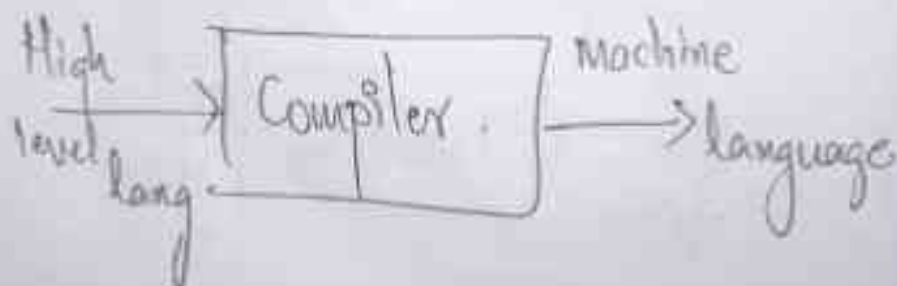
- i) Compilers
- ii) Interpreters
- iii) Assemblers

* Compilers:-

→ Compiler is a translator / Software that converts high-level languages (eg C++) into set of machine level instructions that can be understood by the CPU.

→ Compilers are very large programs, with error-checking and other abilities.

→ Some Compilers generate machine language directly while some converts high-level language to assembly language which is then translated into machine code by assembler.



* Interpreter

→ An interpreter is a translator that repeatedly reads instructions (one at a time) and translate them into machine lang.



* Assembler:-

→ Assembler is a translator which translates middle level languages (eg Assembly language) into machine language.

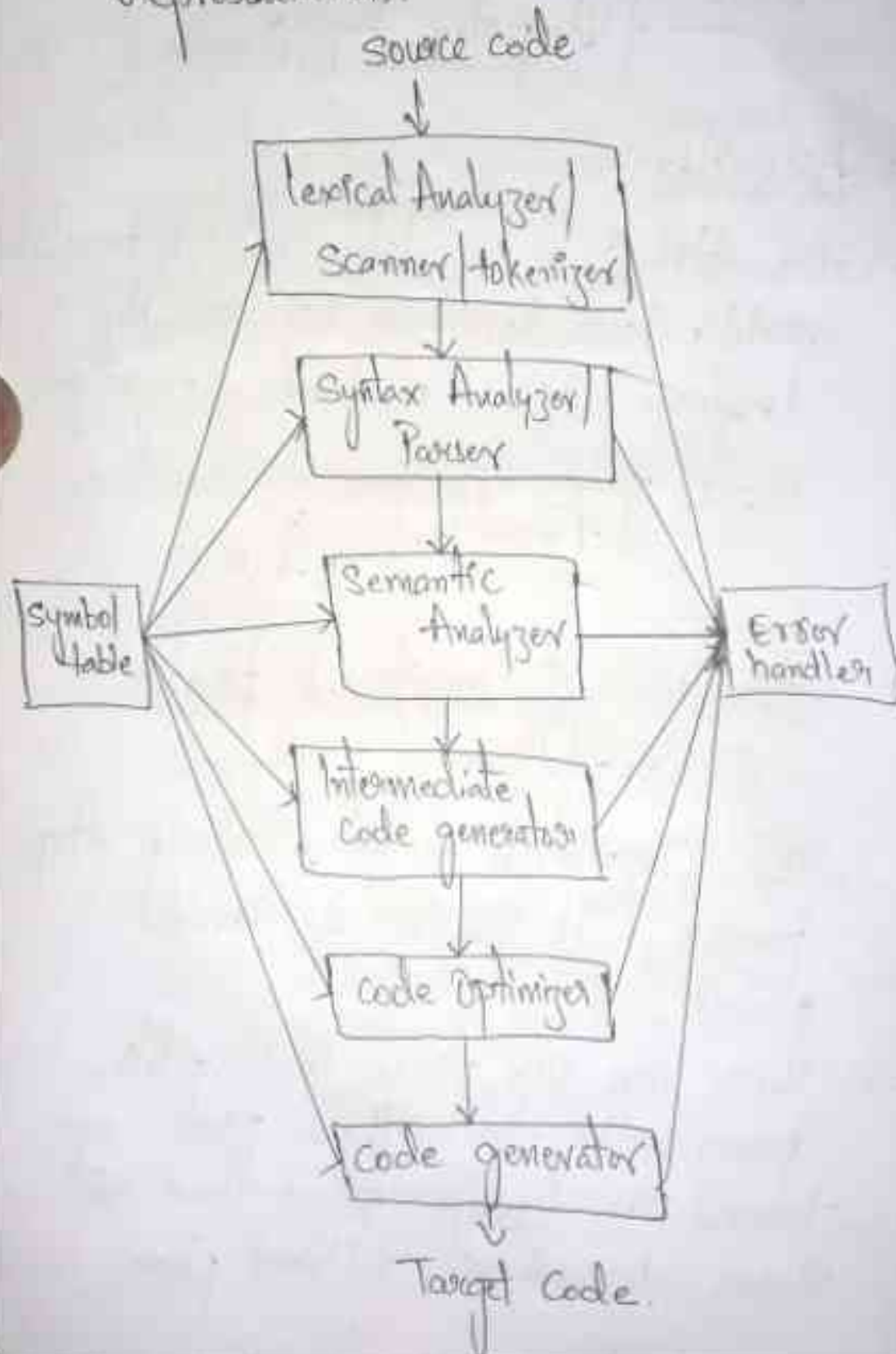


* Structure of Compiler is Phases of Compilation

→ The compilation is not a single step process, it consists of several phases

→ There are six phases in compilation process, the first three phases are termed as Analysis phase and next three termed as Synthesis phase.

- Analysis phase Creates an Intermediate Representation from given source code
- Synthesis phase Creates an Equivalent target program from Intermediate Representation.



- The Compiler has two modules namely - front-end and backend.
- Front-end Constitutes of
- i) Lexical Analyzer
 - ii) Syntax Analyzer
 - iii) Semantic Analyzer
- Back-end Constitutes
- i) Intermediate code generator
 - ii) Code Optimizer
 - iii) Code Generator.

→ The first three phases are language independent, next three phases are machine dependent or target dependent.

Lexical Analyzer:-

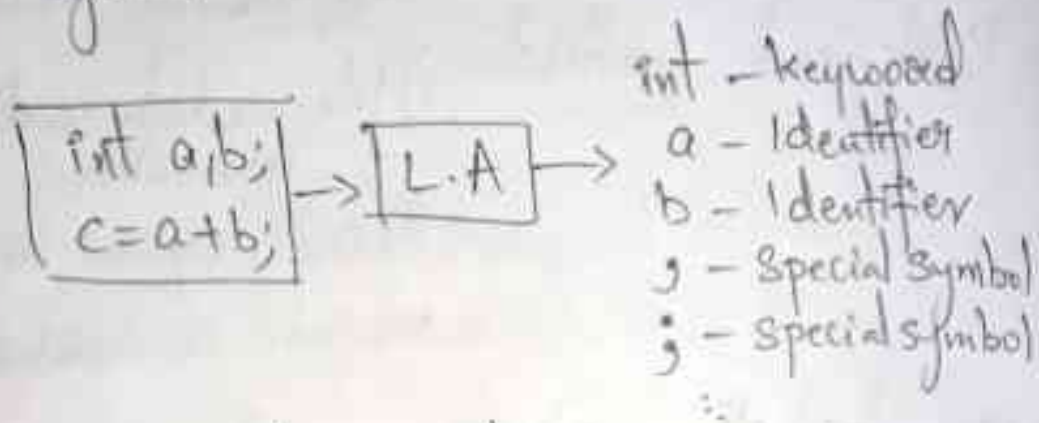
→ The Lexical Analyzer Scans the whole program and breaks down into meaningful units called tokens.

→ It is also called Scanner because it scans the source program and also called as tokenizer because it breaks down the source program into tokens.

Input:- Source program

Output:- Tokens.

- It removes Comments and Whitespaces
- The dictionary required for lexical Analyzer is Regular Expressions
- It is a phase that makes entry into Symbol table



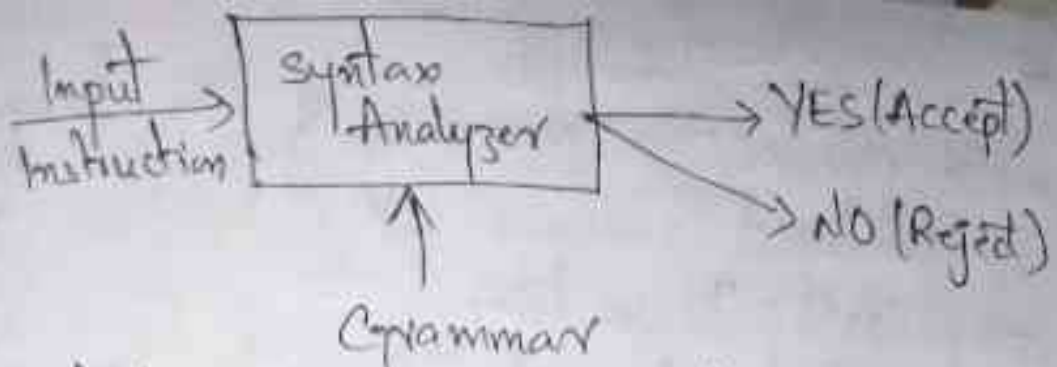
- It reports lexical errors such as
 - Missing of a character
 - Addition of a character
 - Transposition of a character
 - Writing outside of comment.

* Syntax Analyzer:-

- Every language has its own set of rules called grammar.
- The Syntax Analyzer will check whether the given instruction follows the grammar or not.

Input:- Tokens

Output:- Parse Tree.



→ The dictionary required for syntax analyzer is Context-free grammar.

→ Parsing:- The process of checking whether the string is accepted by the grammar or not

eg:- $E \rightarrow E + E \mid E * E \mid \text{id}$

$G = (V, T, P, S) \quad t = \{+, *, \text{id}\}$

LMD

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



Parse tree



* Semantic Analyzer:-

→ It checks the meaning of the language construct

→ Input: Parse tree

Output: Annotated parse tree / semantic tree

→ At this phase type checking and type compatibility is done.

→ With each grammar symbol we associate an attribute. The grammar after attaching attribute is called as Syntax Directed Definition.

→ The parse tree with attributes is called as Annotated parse tree

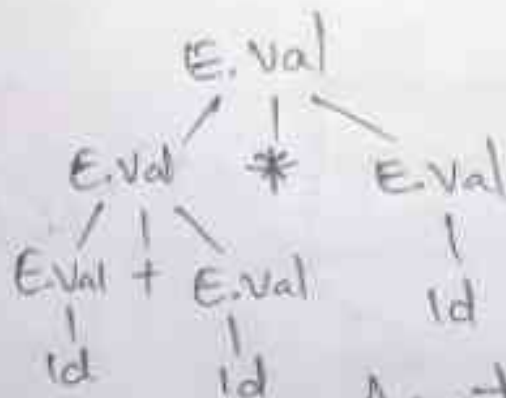
$E \rightarrow E + E \mid E * E \mid Id$

$E.val \rightarrow E.val + E.val$

$E.val \rightarrow E.val * E.val$

$E.val \rightarrow Id$

} Syntax
directed
definition



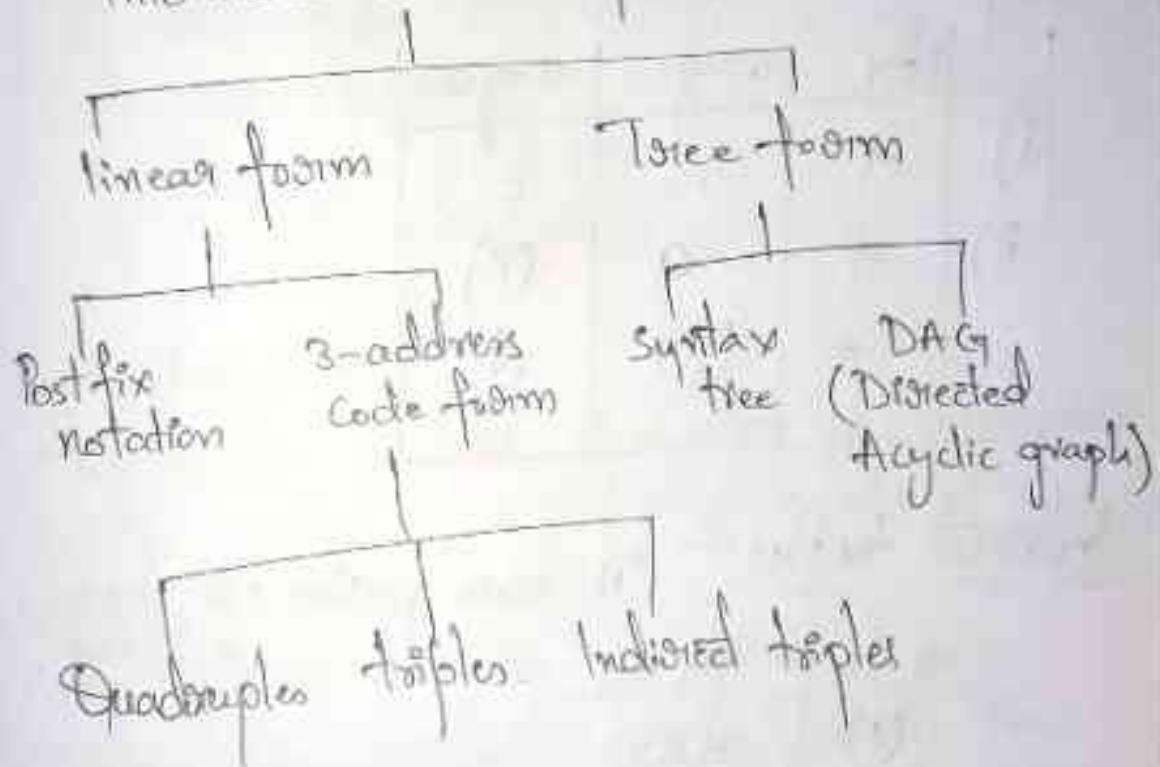
Annotated parse tree.

* Intermediate code generator:-

→ The machine independent front-end part is separated by machine dependent backend part by intermediate code form



Intermediate code form



→ post fix notation:-

→ $a + b * c$

→ $a + bc *$

→ $abc * +$

Quadruples:-

$$d = a + b * c$$

Ops	arg1	arg2	Result
*	b	c	t ₁
+	a	t ₁	t ₂
=	t ₂		d

$$t_1 = b * c$$

$$t_2 = a + t_1$$

$$d = t_2$$

Triples:-

	Ops	arg1	arg2
i)	+	b	c
ii)	+	a	(i)
	=	(ii)	d.

Indirect triples:- It uses pointer to store address.

S.No	ptr to
(1)	1000
(2)	1001

Syntax tree:- $d = a + b * c$



* Code Optimization:-

→ It transforms the code so that it consumes fewer resources and produce more speed

→ Input: Intermediate code form
output: Optimized code

→ Various Code Optimizing techniques are

- i) Common sub-expression elimination
- ii) Dead code elimination
- iii) Copy propagation
- iv) Strength reduction
- v) Loop optimization

* Code Generator:-

→ The main purpose of target code generator is to write a code that machine can understand.

→ The Optimized code is converted into relocatable machine code which then forms the input to linker and loader.

→ Input:- Optimized code

output: Target code

* Symbol table:- It is a data structure being used and maintained by the compiler, consists of the identifiers name along with their types. It helps the compiler to function smoothly by finding identifiers quickly.

* The Translation process:-

- The translator divides the translation process into series of phases
- Each phase manages some particular aspect of translation.

i) $a = b * c + d$. (d is float)

sol a, b, c are integers and
d is float

Symbol table:-

a	
b	
c	
d	

$$a = b * c + d$$



Lexical Analyzer

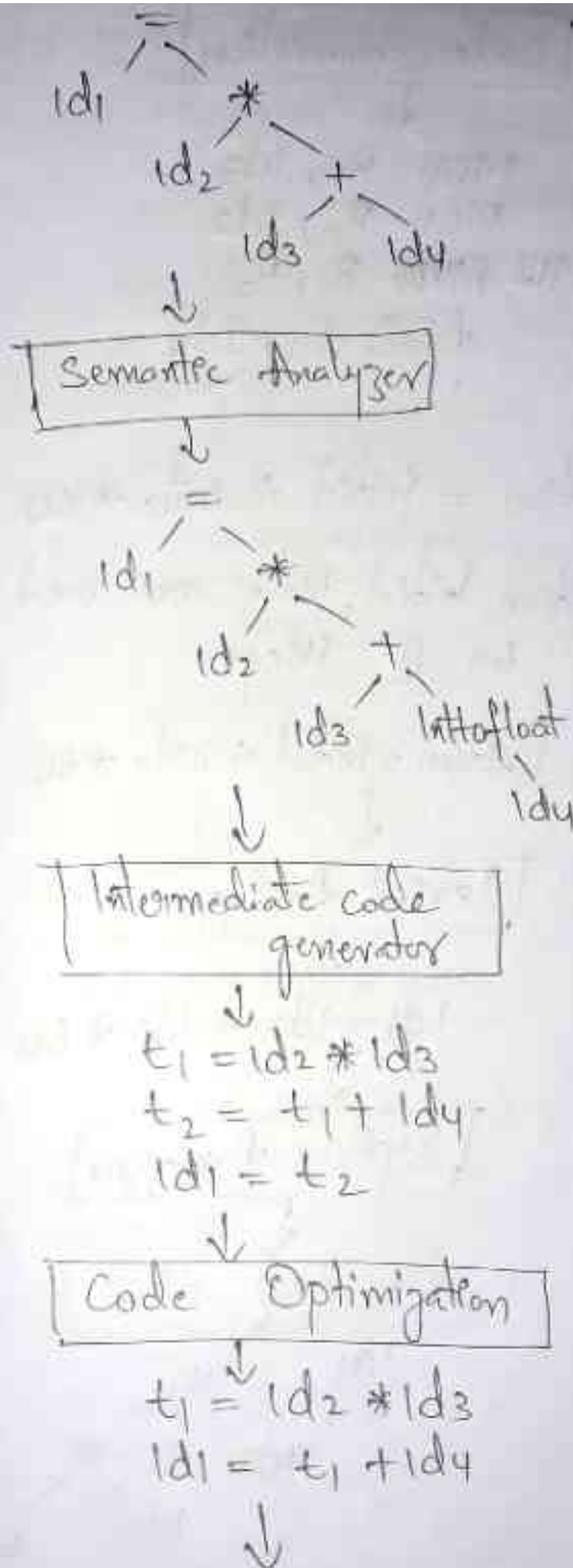


$$Id1 = Id2 * Id3 + Id4$$



Syntax Analyzer





Code Generation

↓
MOV R1, Id3
MOV R2, Id2
~~MUL R1, R2~~
ADD R1, Id4

ii) $\text{Position} = \text{initial} + \text{rate} * 60$

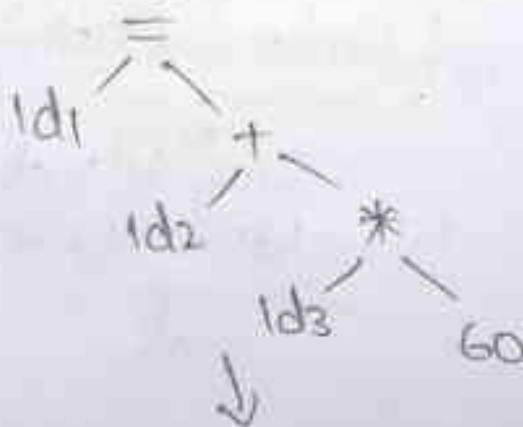
Sol Position, initial, rate are real values
& 60 is integer

$\text{Position} = \text{initial} + \text{rate} * 60$

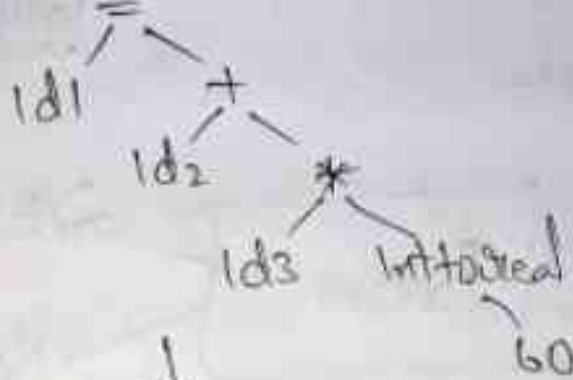
↓
[Lexical Analyzer]

↓
 $\text{Id1} = \text{Id2} + \text{Id3} * 60$

↓
[Syntax Analyzer]



Semantic Analyzer



Intermediate code generation

$t_1 = \text{IntToReal}(60)$
 $t_2 = \text{Id3} * t_1$
 $t_3 = \text{Id2} + t_2$
 $\text{Id1} = t_3$

Code Optimization

$t_1 = \text{Id3} * 60.0$
 $\text{Id1} = \text{Id2} + t_1$

Code generation

MOVF Id3, R2
MULF #60.0, R2
MOVF Id2, R1
ADDF R2, R1
MOVF R1, Id1

* Automatic Lexical Analyzer Generator (LEX).

- LEX is a tool or utility in UNIX-OS
- LEX is a program that generates lexical Analyzer. It is used with YACC parser generator.

YACC: Yet Another Compiler Compiler.

- The lexical Analyzer is a program that transforms an input stream into sequence of tokens.

→ Functions of LEX:

- i) Firstly lexical Analyzer creates a program lex language. The lex compiler runs the lex.l program and produces a c program lex.yy.c



- ii) C Compiler runs lex.yy.c program and produces an object program a.out
- iii) a.out is lexical Analyzer that transforms an input stream into sequence of tokens.



→ LEX File Format:-

→ The lex program is separated into three sections by %/% delimiter.

Declaration Section
%/% Rules Section %/%
Auxiliary procedure Section

i) Declaration Section:-

- It consists of Variables, Constants, manifest that are used in the program.
- It is Optional.

ii) Rules Section:-

- The rules section gives the specification of lex file.
- It describes the pattern or regular expression and corresponding action to be taken.

RE1 pattern1

RE2 pattern2

RE3 pattern3

- It is Compulsory.

iii) Auxiliary procedure section:-

- It consists of all the functions and methods used in a program

Q) Write a lex program to identify keywords, numbers, identifiers, operators, special symbols.

→ J\$ vi lex.l # // Creating lex.l file

```
%%
```

```
[a-zA-Z][a-zA-Z0-9]* printf("Identifier");
```

```
int|float|while|for printf("keywords");
```

```
[+|-|*|/] printf("Operators");
```

```
[0-9]* printf("Number");
```

```
[@|$] printf("Special Symbol");
```

```
%%
```

```
main()
```

```
{ yylex();
```

```
y
```

```
int yywrap()
```

```
{  
}
```

→ J\$ lex dext.l

J\$ cc lex.yy.c

J\$./a.out

Output:-

area	Identifier
foot	keyword
+	Operator
456	Number
@	Special Symbol.

* Role of lexical Analyzer:-



→ lexical Analyzer scans whole programs and divides into small tokens.

→ It removes white spaces & comments.

* lexemes, tokens, patterns:-

<u>lexeme</u>	<u>token</u>	<u>pattern</u>
int	keyword	→ -
circle	Identifier	→ letter (letter/digit)*
10	number	→ [0-9]*

Lexemes:- It is defined as sequence of characters that are matched with the pattern of the token and is identified by lexical analyzer as an instance of token.

Ex:- int, circle, a etc

Token:- Tokens are the sequence of characters having collective meaning. They describe class and category of input string. Token is a pair of token name & its value (attribute)

Pattern:- Patterns are the set of rules that describe the lexeme representing a particular token

Ex:- Pattern for an Identifier $\rightarrow$ letter
(letter/digit)*
Pattern for number $\rightarrow$ (0-9)*

* Input Buffering:-

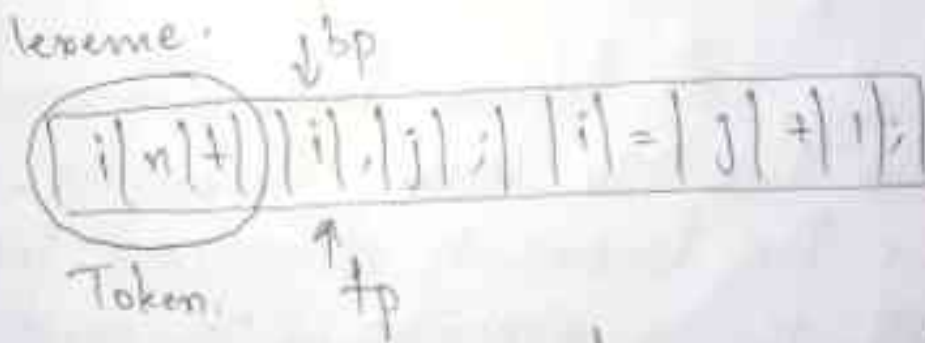
- The lexical Analyzer scans the input from left to right one character at a time. It uses two pointers begin ptr (bp) and forward pointer (fp)



↑ → fp moves forward.

fp

- Initially both pointers point to first character.
- The fp moves ahead to search for end of lexeme. As soon as blank space is encountered, it indicates the end of lexeme.



- To process large amount of source program, specialized buffering techniques are used to reduce the amount of overhead required to process a single input character.

→ There are two techniques:-

- i) Buffer pairs
- ii) Sentinels

i) Buffer Pairs:-

→ Two buffers are involved that are alternatively loaded. Each buffer is of same size n (where $n = 4096$ bytes)

→ Two pointers are maintained.

i) lexeme begin: It marks to the beginning of current lexemes

ii) Forward pointer: It scans ahead until a pattern match is found

→ Initially, both lexeme begin & forward pointer to the same character. Once the lexeme is determined forward is set to the character at its right end.

→ Once the lexeme is recorded as token, lexeme begin is set to character immediately after the lexeme is found.

→ Advancing forward requires that the first test whether we have reached the end of one of the buffers & if so we must reload the other buffer from the inputs & move forward fp to the beginning of newly loaded buffer.

Algorithm Advance forward:-

if fp reaches the first half then
 reload the second half

fp = fp + 1;

~~if~~ else if fp reaches the second
half then reload the first half
 move fp to the beginning of first half

~~if~~ else

fp = fp + 1;

end.

Disadvantage with buffer pairs:-

→ The pointer fp is required to perform two tests while reading the character from the i/p

a) to check the end of the half

b) whether the characters reach the end of the lexeme or not.

ii) Sentinels:-

- > The above disadvantage can be overcome by using Sentinels.
- > In Sentinels we use certain character (not a part of source language) to mark the end of the half generally we use EOF.

Algorithm:-

$fp = fp + 1;$

if $fp = EOF$

if fp reaches first half then
reload the second half.

$fp = fp + 1$

else if fp reaches second half then
reload the first half move
 fp to the begin of first half.

else
terminate lexical Analysis

end;

* T-diagram:-

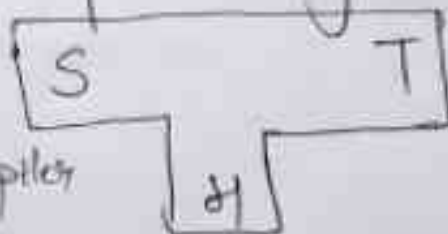
→ A Compiler can be categorized by

i) Source language

ii) Target language

iii) Host language (Implementation language)

This is depicted by T-diagram as



$S \rightarrow$ Source lang.

$T \rightarrow$ Target lang.

$H \rightarrow$ host language

T diagram shows a Compiler
 $S \xrightarrow{C} T$
 H

* Cross Compiler:-

→ It is defined as a Compiler running on one machine and producing the object code for another machine.



→ A compiler written in language H (host language) that translates language S into language T.
 Compiler $S \rightarrow T$

* Bootstrapping:-

→ It is a technique for producing efficient compilers.

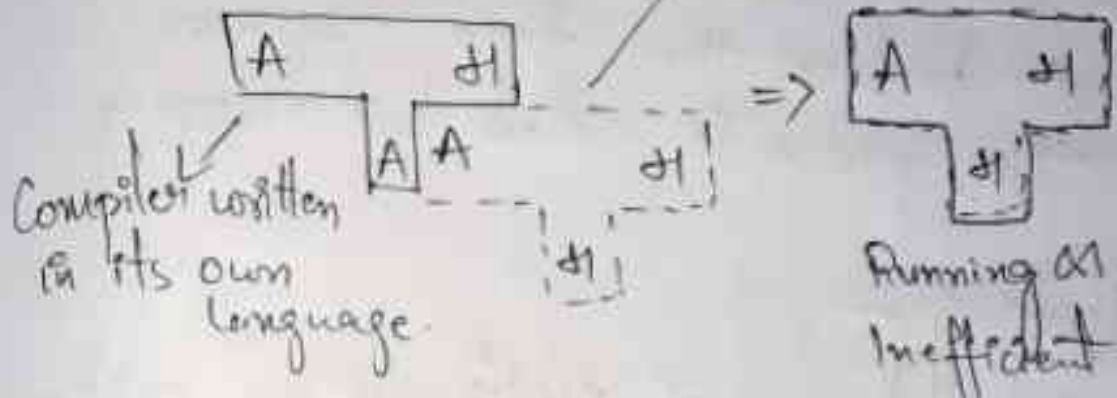
→ Bootstrap compiler is used to compile the compiler and then you can use this compiled compiler to compile everything else as well as future version of itself.

→ There are two stages in bootstrapping.

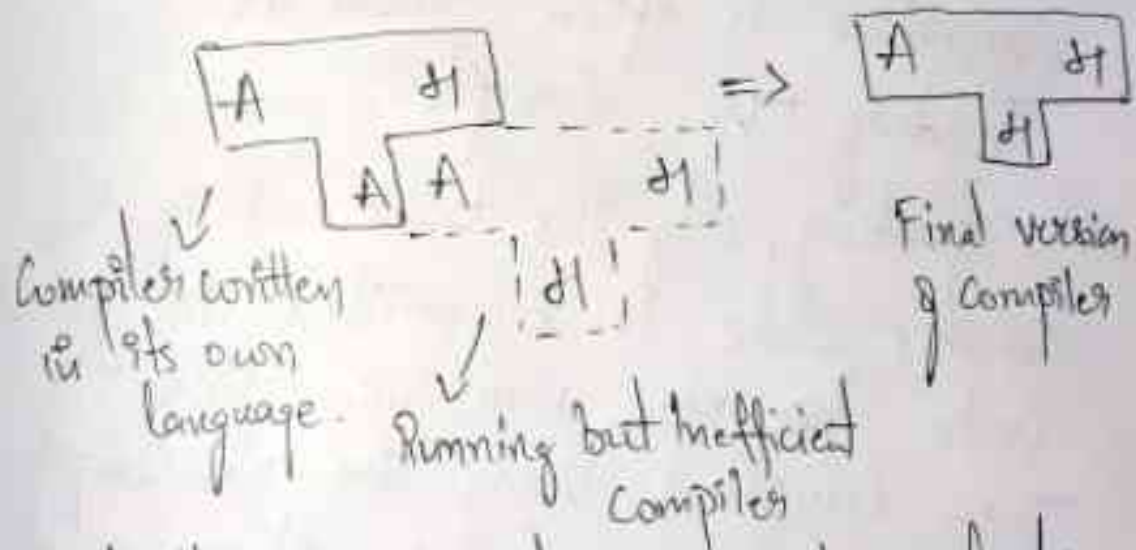
→ Stage 1: - A compiler for subset language is developed which is running and is inefficient

stage 1: subset language

Quick & dirty compiler.



stage 2: (Full language)



→ Initially we write a quick & dirty compiler in assembly language translating only those features of the language that are used in the compiler

→ The quick & dirty Compiler will produce extremely inefficient code. Once we run it, we use it to compile the good Compiler.

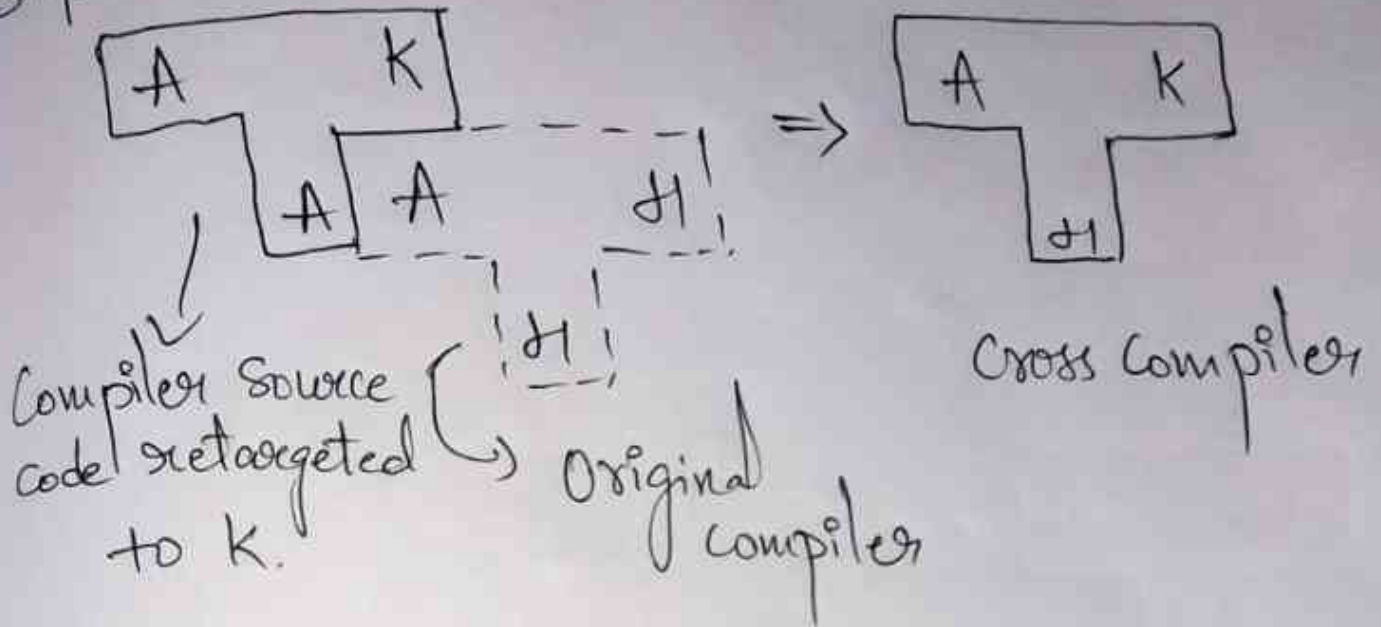
→ Then we recompile the good Compiler to produce the final version. This process is called bootstrapping.

Porting:-

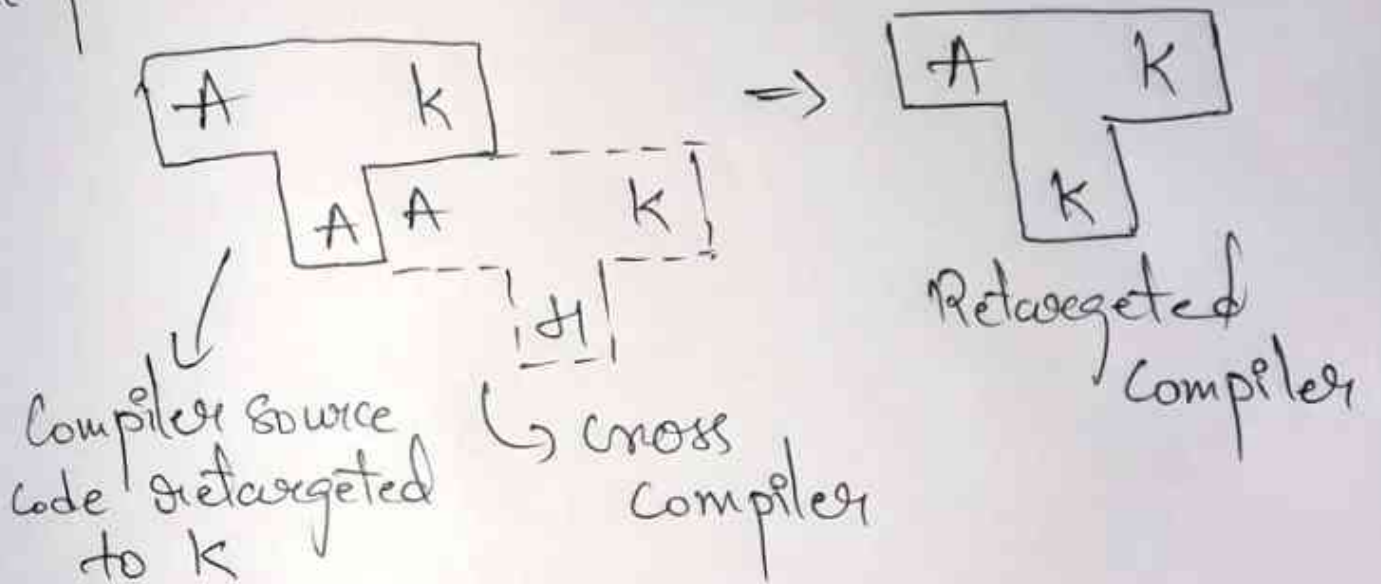
→ The process of modifying an existing Compiler to work on a new machine is often known as Porting the Compiler.

→ In porting, only the back end of the source code is rewritten to generate code for new machine. This is then compiled using the old Compiler to produce a cross Compiler, and the Compiler is again recompiled by the cross Compiler to produce a working version for new machine.

Step 1:-



Step 2:-



* Major Data Structures in Compiler:-

1. Symbol table
2. Literal table
3. Parse tree
4. Syntax tree

* Symbol table:-

- It stores information about program identifiers that will be used across the phases
- Compiler allocate memory to symbol table.
- Symbol table has facilities to manipulate (add/delete) the elements in it
- Identifiers are name of variables, constants, functions etc

Eg:- $X = Y + Z$

Symbol	Category	Type	Attribute	Memory location
X	Identifier	float	#1	1000
=	Operator		Assignment(1)	
Y	Identifier	float	#2	1100
+	Operator		Arithmetic(1)	
Z	Identifier	float	#3	1200

* Literal table:-

- It is also a data structure used to maintain the details of constants & strings used in the program
- It reduces the size of a program in memory by allowing reuse of constants & strings
- It is needed by code generator to construct symbolic address for literals

Eg:- `int n = 60;`

Literal	Category	Attribute	Memory location
60	Const	int	1200

* Parse tree:-

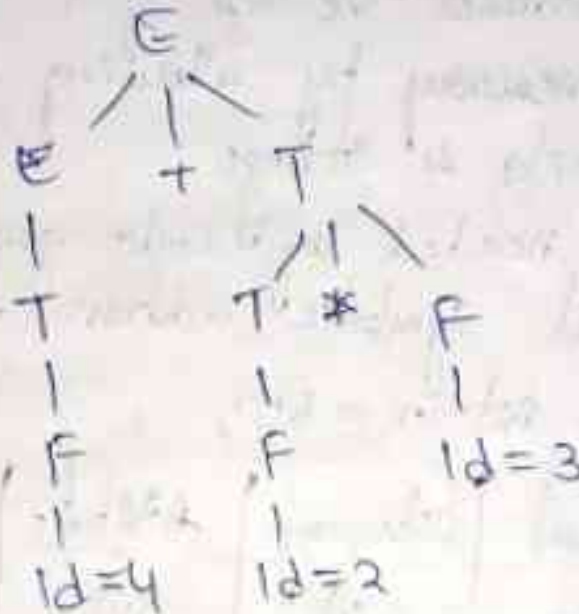
- It is dynamically-allocated pointer based structure.
- It describes the syntactic structure of input.
- In parse tree, the terminal nodes ^{tokens and} represent the ^{tokens and} interior nodes represent expression
- In parse tree nodes are variant type records storing information for different types of data
- It is constructed from bottom-up

Eg:- $E \rightarrow E + T \mid T$

$T \rightarrow T * F \mid F$

$F \rightarrow (E) \mid id$

Suppression:- $4 + 2 * 3$



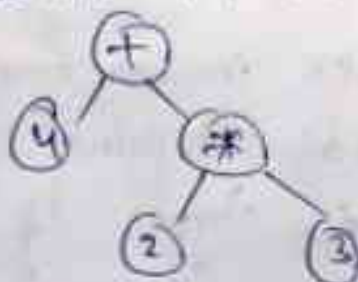
* Syntax tree:-

→ Syntactic structure is ^{also} represented using Syntax tree

→ A Syntax tree is a compressed representation of parse tree

→ In Syntax tree, Operators appear as interior nodes & Operands as children.

Eg:- Expression:- $4 + 2 * 3$



Syntax tree.

* Specification of tokens:-

→ There are three specification of tokens

- i) String
- ii) Language
- iii) Regular expression

* Recognition of tokens:-

→ In this we study how to take pattern for tokens and build a piece of code that examines the input string.

Grammar: $stmt \rightarrow \text{if } expr \text{ then } stmt$
 $\quad \quad \quad | \text{if } expr \text{ then } stmt \text{ else } stmt$

$\quad \quad \quad | \epsilon$

$expr \rightarrow term \text{ diclop } term$
 $\quad \quad \quad | term$

$term \rightarrow id \mid number$

Variables = $\{ stmt, expr, term \}$

Terminals = $\{ \text{if, then, else, diclop, id, number} \}$

Relop : Relational Operators

Relop $\rightarrow <|>|<=|>=|=|<>$

$<>$ represent not equals

$\rightarrow$ Pattern for tokens:-

digit $\rightarrow [0-9]$

digits $\rightarrow [0-9]^+$ or

digits $\rightarrow \text{digit}^+$

number $\rightarrow \text{digits}(\cdot \text{digits})?$

$(E[+|-]? \text{digits})?$

letter $\rightarrow [A-Za-z]$

id $\rightarrow \text{letter}(\text{letter}|\text{digit})^*$

if $\rightarrow \text{if}$

then $\rightarrow \text{then}$

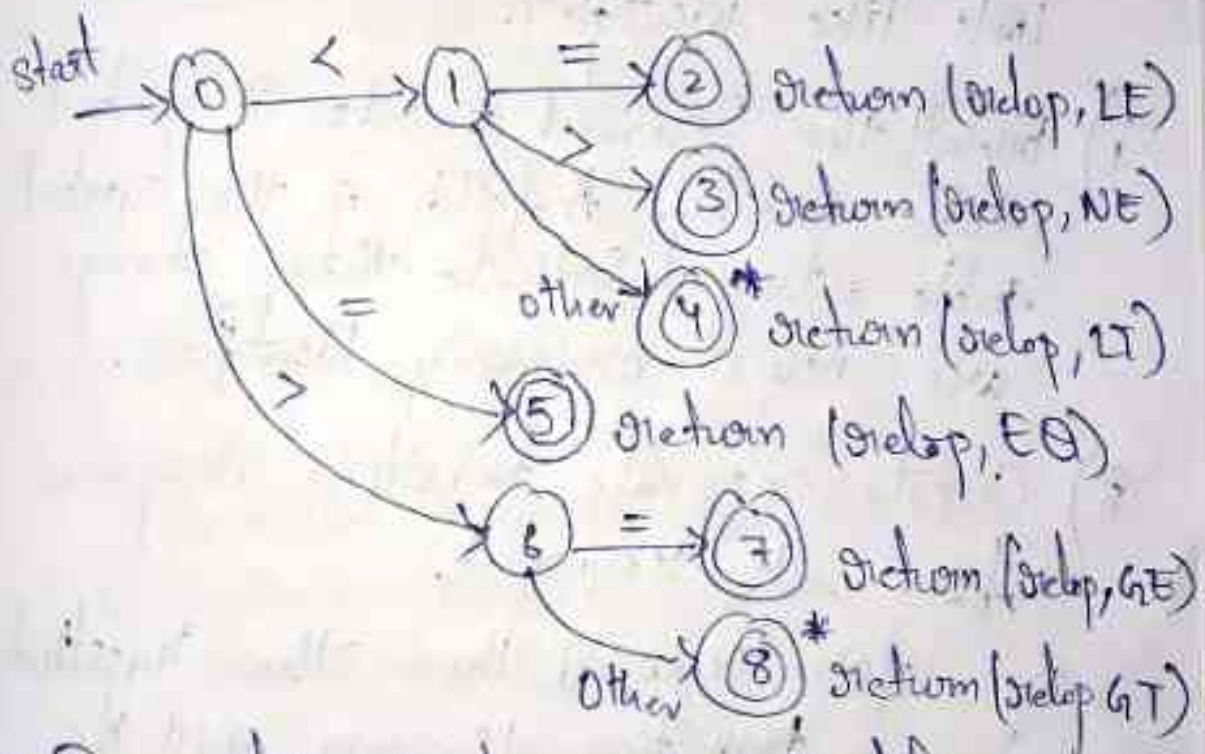
else $\rightarrow \text{else}$

Relop $\rightarrow <|>|<=|>=|=|<>$

Transition Diagram:-

- It is an intermediate step in construction of lexical analyzer.
- It is a collection of states with edges as transitions.

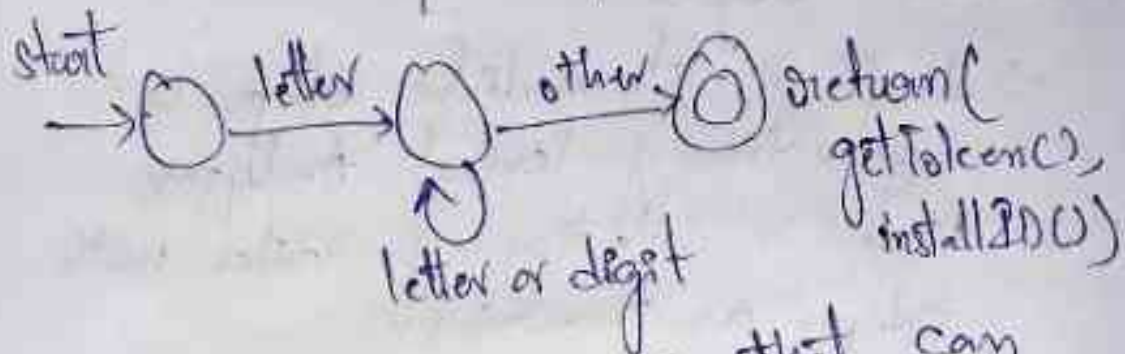
Transition diagram for `<`, `=`, `>`:-



* Recognition of keyword & Identifier:-

- Recognizing keywords and Identifier presents a problem. Usually keyword like `if` or `then` are reserved as so they are not identifiers even though they look like Identifier.

Transition diagram for ids & keyword

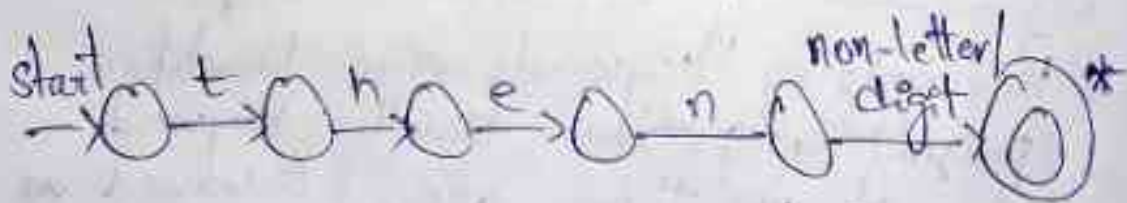


→ There are two ways that can handle reserved words that look like Identifier.

i) Install the Reserved words in symbol table initially. A field of the symbol table entry indicates these strings are never ordinary Identifier.

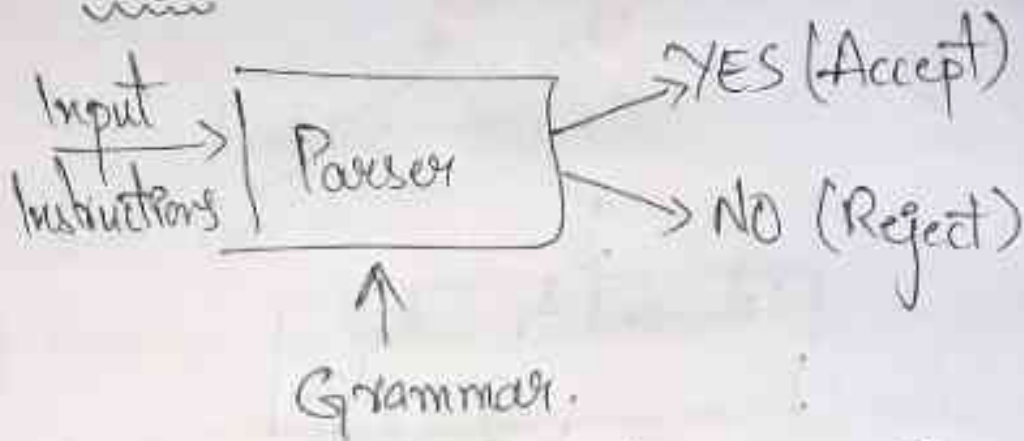
ii) Create separate transition diagram for each keyword.

Eg:- for keyword then the hypothetical transition diagram will be



UNIT:-02 Syntax Analyzer / Parser:-

* Parser:-



→ Every programming language is described by grammar (context-free grammar)

→ A Grammar can be described using 4 tuples $G = (V, T, P, S)$

V = set of Variables / non-terminals

T = set of Terminals

P = Production rule

S = Start symbol

→ Terminals are lower case letters, operators and special symbols
Eg:- a, b, +, -, id

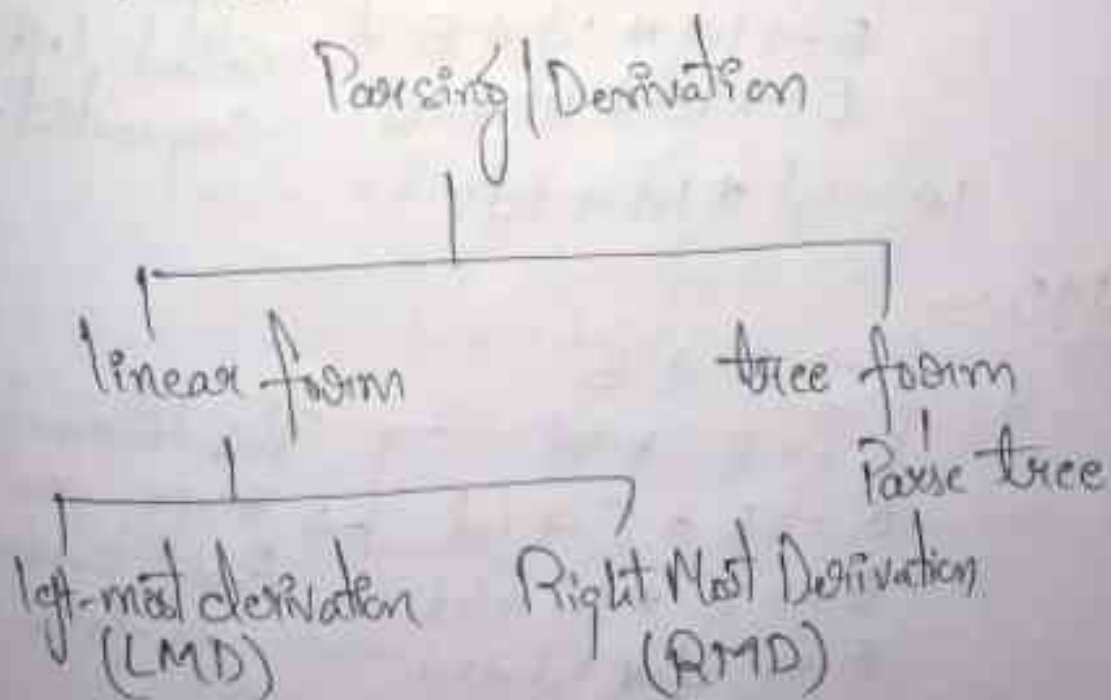
→ Variables are denoted by Upper Case letters or words.

Eg:- A, S, B

→ Variables can be deduced terminals cannot be deduced.

* Derivation:- (Parsing)

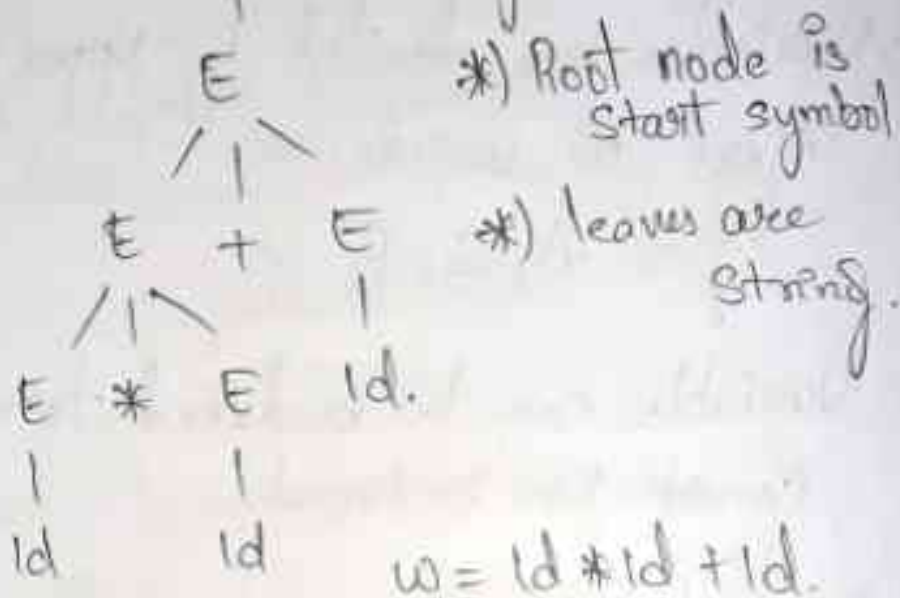
→ The process of obtaining a string from start symbol is called Derivation.



Parse tree:-

$$E \rightarrow E + E \mid E * E \mid Id$$

Parse tree for above grammar



LMD:-

$$\begin{aligned}
 E &\rightarrow E + E \\
 E &\rightarrow E * E + E \\
 E &\rightarrow Id * E + E \\
 E &\rightarrow Id * Id + E \\
 E &\rightarrow Id * Id + Id
 \end{aligned}
 \left. \begin{array}{l} \\ \\ \\ \\ \end{array} \right\} \begin{array}{l} \text{This intermediate} \\ \text{steps are} \\ \text{called left} \\ \text{sequential form} \end{array}$$

$w = Id * Id + Id.$

RMD:-

$$\begin{aligned}
 E &\rightarrow E + E \\
 E &\rightarrow E + Id \\
 E &\rightarrow E * E + Id \\
 E &\rightarrow E * Id + Id \\
 E &\rightarrow Id * Id + Id
 \end{aligned}
 \left. \begin{array}{l} \\ \\ \\ \\ \end{array} \right\} \begin{array}{l} \text{This intermediate} \\ \text{steps are called} \\ \text{Right sequential} \\ \text{form.} \end{array}$$

$w = Id * Id + Id$

* Ambiguous Grammar:-

→ A grammar is said to be ambiguous if there exists two or more LMD / Left most parse tree or two or more RMD / Right most parse-tree generating the same string.

eg:- $E \rightarrow E + E \mid E * E \mid id.$

1) L.M.D

$$E \rightarrow E + E$$

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

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

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

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

$$w_1 = id * id + id.$$

2) L.M.D

$$E \rightarrow E * E$$

$$E \rightarrow id * E$$

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

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

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

$$w_2 = id * id + id.$$

∴ w_1 is equals to w_2 , hence we say the grammar is ambiguous.

* Removing Ambiguity:-

→ Conversion of Ambiguous grammar to Unambiguous grammar

Let $G = (V, T, P, S)$ be Ambiguous grammar then its Unambiguous grammar will be $G' = (V', T, P', S)$

→ Rules:-

- 1) Number of new variables to be added = Number of Operators
- 2) If Operator is left associative replace the right most variable

Ex:- $E \rightarrow E + E \mid E * E \mid id$
is ambiguous grammar.

No. of operators = 2 ($\because +, *$)

So, 2 new variables are added to remove ambiguity.

New variables be T, P

$$E \rightarrow E + E$$

$\hookrightarrow +$ is left associative

So, replace its
right most variable
with new variable

$$\textcircled{1} E \rightarrow E + T$$

$$\textcircled{2} E \rightarrow T$$

Similarly, $E \rightarrow E * E$ written as

$$T \rightarrow T * T$$

$\hookrightarrow$ left associative

So replace
right most variable
with new variable

$$\textcircled{3} T \rightarrow T * P$$

$$\textcircled{4} T \rightarrow P$$

$$\textcircled{5} P \rightarrow \text{Id}$$

So the Unambiguous grammar is

$$E \rightarrow E + T$$

$$E \rightarrow T$$

$$T \rightarrow T * P$$

$$T \rightarrow P$$

$$P \rightarrow \text{Id}$$

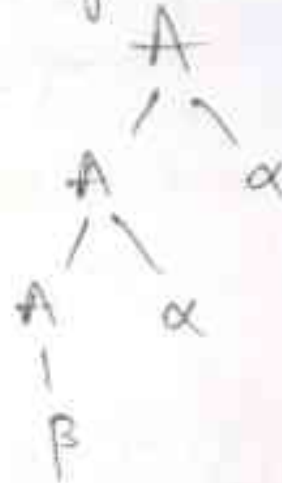
This grammar is free from ambiguity.

Note:- 1) $+$, $-$, $/$, $*$ are left associative.

2) $\wedge$ is right associative

* Left-Recursion:-

→ The grammar of the form
 $A \rightarrow A\alpha / \beta$ is left recursive grammar.

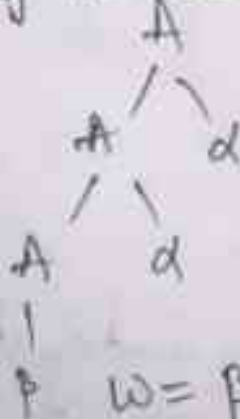


→ Eliminating left recursion
 we transfer the grammar into

$$A \rightarrow \beta A'$$

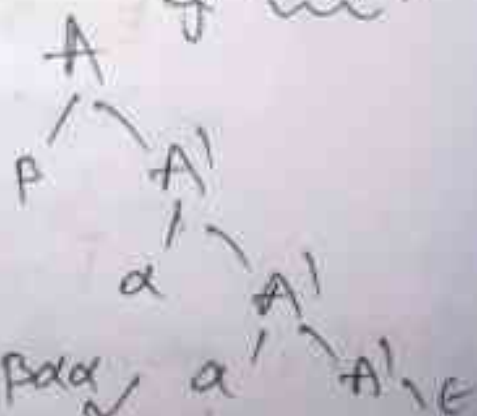
$$A' \rightarrow \alpha A' / \epsilon$$

Left Recursion



$$w = \beta \alpha \alpha$$

after eliminating left recursion



$$w = \beta \alpha \alpha$$

Q1) Eliminate Left Recursion

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

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

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

Ans

$$G = (V, T, P, S)$$

$$N = \{E, T, F\} \quad T = \{+, *, \rangle, \langle, id\}$$

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

left recursion, if A is in the form
 $A \rightarrow A\alpha \mid \beta$

$$A = E$$

$$\alpha = +T$$

$$\beta = T$$

$\begin{aligned} \therefore A &\rightarrow \beta A' \\ A' &\rightarrow \alpha A' \mid \epsilon \end{aligned}$

we get,

$$E \rightarrow TA' \quad (1)$$

$$E' \rightarrow +TE' \mid \epsilon \quad (2)$$

Similarly, $T \rightarrow T * F \mid F$

$$T \rightarrow FT' \quad (3)$$

$$T' \rightarrow *FT' \mid \epsilon \quad (4)$$

Therefore, the grammar will be :

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

$$F \rightarrow (E) / \text{Id.}$$

The above grammar is "free from left recursion."

* Indirect left recursion:-

→ A grammar is said to have indirect left recursion, on substituting any variable in other production we get direct left recursion.

Eg:- $S \rightarrow Aa / b$
 $A \rightarrow Ac / Sd / \epsilon$

It is indirect left recursion as on substituting A in S we get

$$S \rightarrow Sda / Aca / a / b$$

Left recursion

$$S \rightarrow Sda | Aca | a | b$$

left recursion

It is in the form

$$A \rightarrow A\alpha | \beta_1 | \beta_2 | \beta_3$$

we get $A \rightarrow \beta_1 A' | \beta_2 A' | \beta_3 A'$

$$A' \rightarrow \alpha A' | \epsilon$$

$$A \Rightarrow S, \alpha = da$$

$$\beta_1 = Aca, \beta_2 = a, \beta_3 = b$$

we get

$$S \rightarrow AcaA' | aA' | bA'$$

$$A' \rightarrow daA' | \epsilon$$

* left-factoring:-

→ left factoring is a grammar transformation technique. It consists in "factoring out" prefixes which are common to two or more production.

Eg:- $A \rightarrow \alpha\beta_1 | \alpha\beta_2$

After them performing left factoring
on above grammar
we get

$$A \rightarrow \alpha A'$$

$$A' \rightarrow \beta_1 / \beta_2$$

Eg:- $S \rightarrow iEtS / iEtScS / a$
 $E \rightarrow b$

Eliminate left factor

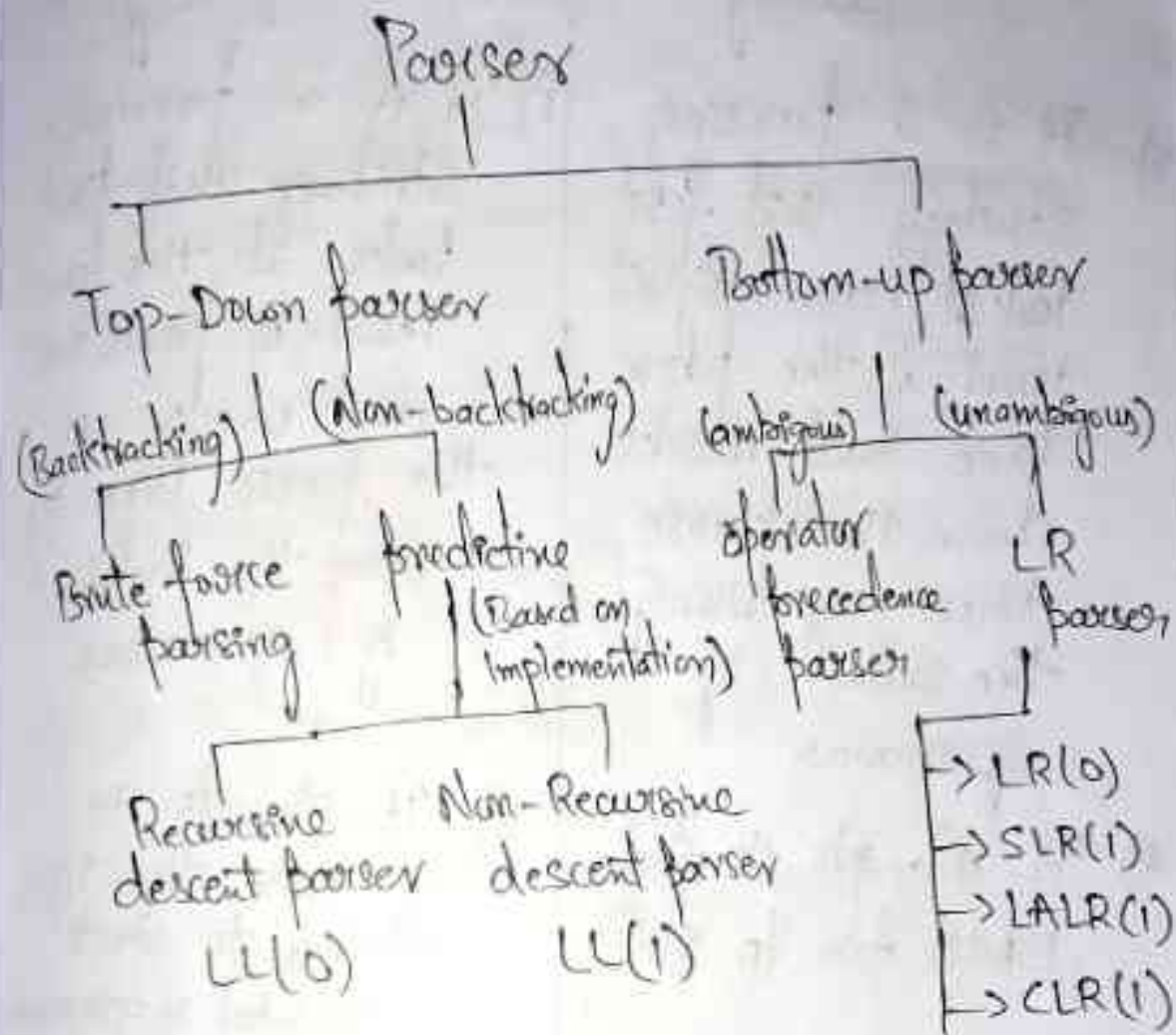
$$S \rightarrow iEtSS' / a$$

$$S' \rightarrow eS / \epsilon$$

$$E \rightarrow b$$

"

* Parsers (Types)



SLR(1) - Simple LR Parser

LALR(1) - Look Ahead LR Parser

CLR(1) - ~~Canonical~~ LR Parser
Canonical

→ LL(0) or LL(1)

↳ It Represents that input is read from left to right
↳ It Represents Left most Derivation

→ LR
↳ It Represents Right most Derivation
↳ It Represents that input is read from left to right.

Top-Down vs Bottom up parsing parsing

1) It is a parsing strategy that first look at the highest level of the parse tree and works down the parse tree by using the rules of grammar

2) It attempts to find LMD for I/p string

3) We start parsing from top to bottom

4) It use LMD

5) It's main decision is to select what production rule to use in order to construct the string

1) It is a parsing strategy that first looks at the lowest level of parse tree and works up the parse tree by using the rules of grammar

2) It attempts to produce the i/p string to start symbol of grammar

3) we start parsing from bottom to up

4) It use RMD

5) It's main decision is to select when to use a production rule to produce the string to get starting symbol

6) It handles unambiguous grammar & grammar free from left recursion

7) It involves production

8) Parsing table size is less

9) It is less powerful

10) Error detection is Easy

6) It handles both ambiguous & unambiguous grammar

7) It involves Reduction

8) Parsing table size is more

9) It is more powerful

10) Error detection is difficult.

* Brute force parsing:-

→ It is a top-down parser which involves backtracking.

→ Brute force strategy:-

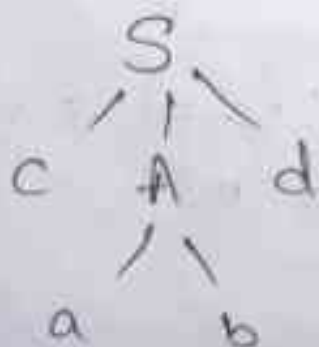
> Start with the start symbol and replace it with first Alternative.

> If there is a variable & it has many alternatives, then select the first alternative and compare leafs with given lp string. If they are not equal then select second alternatives. If all the alternatives are completed backtracks go back and substitute the second alternative. Continue the process until string is Accepted or Rejected.

Eg:-

Q) Show the acceptance of string cad for the grammar by brute force method. $S \rightarrow cAd$
 $A \rightarrow ab|a$
 $w = cad$

100%



$w = cabd \neq cad$
 It backtracks

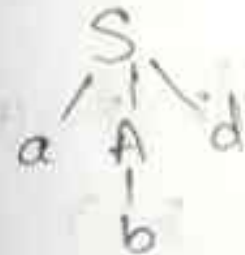


$w = cab$ ✓

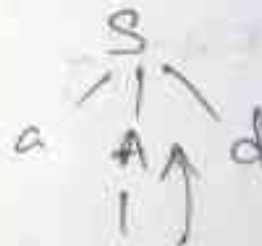
String is Accepted

- 9) Show the acceptance of String
 $w = addc$ for $S \rightarrow aAd/aB$
 $A \rightarrow b/c$
 $B \rightarrow ccd/ddc$

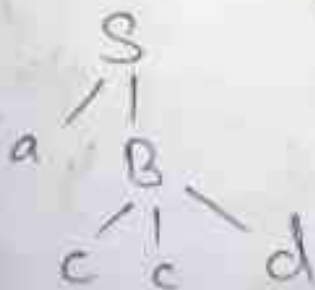
Sol $w = addc$



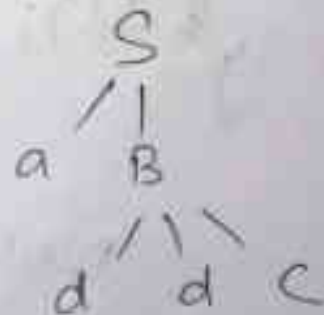
$abd \neq addc$
Backtracks



$acd \neq addc$
Backtracks



$acd \neq addc$
backtracks



$w = addc$ ✓

String is Accepted

* Recursive descent Parser LL(0):-

→ In Recursive descent parser for each variable, we associate or write a procedure or function

Q) Construct LL(0) for following grammar

- i) $E \rightarrow E + T / T$
 $T \rightarrow T * F / F$
 $F \rightarrow (E) / id$

1) $E \rightarrow E + T / T$ 2) $T \rightarrow T * F / F$ 3) $F \rightarrow (E) / id$

```
E()  
{  
  E();  
  +  
  T();  
} else  
{  
  T();  
}
```

```
T()  
{  
  T();  
  *  
  F();  
} else  
{  
  F();  
}
```

```
F()  
{  
  (  
  E();  
  )  
} else  
{  
  id  
}
```

```
main()  
{  
  E(); // start symbol  
}
```

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

11/8/20 LL(0) or Recursive descent parser

1) $E \rightarrow TE'$ 2) $E' \rightarrow +TE' | \epsilon$

```

E()
{
    T();
    E'();
}

E'()
{
    +
    T();
    E'();
}

```

3) $T \rightarrow FT'$ 4) $T' \rightarrow *FT' | \epsilon$ 5) $F \rightarrow (E) | id$

```

T()
{
    F();
    T'();
}

T'()
{
    *
    F();
    T'();
}

F()
{
    (
    E();
    )
    else
    id
}

```

```

main()
{
    E();
}

```

* Non-Recursive Descent Parser LL(1) parser

→ LL(1) Parser is a top-down parser without backtracking strategy.

→ LL(1)

- 1 represent the number of look-ahead, which means how many symbols are we going to see when you want to make decision.
- and L represent that we will be using left most derivation
- 1st L represent that scanning of input is done from left to right.

→ LL(1) parser will have

i) Parsing table — To check the grammar is in LL(1) or not

ii) Driver Routine — It is used to check whether the string is accepted or not

* Algorithm to Construct LL(1) Parsing table:-

- 1) If grammar is left recursive then eliminate the left recursion. The grammar must be free from left-recursion, ambiguity and left-factoring.
- 2) Identify null and non-null productions
- 3) Calculate $\text{First}()$ for non-null production and $\text{Follow}()$ for null production.
- 4) Draw the parsing table.

* $\text{First}()$:-

→ If there is a variable, and from that variable, if we try to derive all the strings then beginning terminal symbol is called the First .

Rules:-

- i) If x is a terminal, then $\text{First}(x) = \{x\}$
- ii) If $x \rightarrow \epsilon$, is a production rule then add ϵ to $\text{First}(x)$
- iii) If $x \rightarrow y_1 y_2 y_3 \dots y_n$ is a production,
1) $\text{First}(x) = \text{First}(y_1)$ If 2) $\text{First}(y_1)$ contains ϵ then $\text{First}(x) = \{ \text{First}(y_1) - \epsilon \} \cup \{ \text{First}(y_2) \}$

3. If $\text{First}(Y_i)$ contains ϵ for all $i=1$ to n then add ϵ to $\text{First}(X)$

Eg:- 1) Production rules of Grammar

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

FIRST Sets:-

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

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

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

$$\text{FIRST}(T') = \{ *, \epsilon \}$$

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

ii) Find FIRST sets for the grammar

$$S \rightarrow ACB \mid Cbb \mid Ba$$

$$A \rightarrow da \mid BC$$

$$B \rightarrow g \mid \epsilon$$

$$C \rightarrow h \mid \epsilon$$

$$\text{FIRST}(S) = \text{FIRST}(ACB) \cup \text{FIRST}(Cbb) \cup \text{FIRST}(Ba)$$

$$= \text{FIRST}(A) \cup \text{FIRST}(C) \cup \text{FIRST}(B)$$

$$= \{d, g, h, b, a, \epsilon\}$$

$$\text{FIRST}(A) = \{d, g, h, \epsilon\}$$

$$\text{FIRST}(B) = \{g, \epsilon\}$$

$$\text{FIRST}(C) = \{h, \epsilon\}$$

* FOLLOW():-

→ What is the terminal symbol which follows a variable in the process of derivation.

Rules:-

i) $\text{FOLLOW}(S) = \{ \$ \}$ // $S \rightarrow$ start symbol

ii) If $A \rightarrow pBq$ is a production
then everything in $\text{FIRST}(q)$
except ϵ will be in
 $\text{FOLLOW}(B)$

iii) If $A \rightarrow pB$ is a production,
then everything in $\text{FOLLOW}(A)$
will be in $\text{FOLLOW}(B)$

iv) If $A \rightarrow pBq$ is a production and
 $\text{FIRST}(q)$ contains ϵ then
 $\text{FOLLOW}(B)$ contains

$$\{ \text{FIRST}(q) - \epsilon \} \cup \text{FOLLOW}(A)$$

Eg:-

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

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

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

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

$$\begin{aligned} \text{FOLLOW}(T') &= \text{FOLLOW}(T) \\ &= \{ +, \$,) \} \end{aligned}$$

$$\begin{aligned} \text{FOLLOW}(F) &= \{ \text{FIRST}(\text{FOLLOW}(T')) - \epsilon \} \cup \\ &\quad \{ \text{FOLLOW}(T') \} \cup \\ &\quad \{ \text{FOLLOW}(T) \} \\ &= \{ *, +, \$,) \} \end{aligned}$$

Therefore.

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

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

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

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

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

Q) Construct LL(1) Parser or check whether the given grammar is in LL(1) or not

$$\begin{aligned} 1) \quad E &\rightarrow E + T \mid T \\ T &\rightarrow T * F \mid F \\ F &\rightarrow (E) \mid \text{Id} \end{aligned}$$

Sol Variables = $V = \{ E, T, F \}$
 $T = \{ +, *, (,), \text{Id} \}$

The grammar is left recursive

Step 1: Eliminate left recursion from grammar we get,

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

$$V = \{ E, E', T, T', F \}$$

$$T = \{ +, *, (,), \text{Id} \}$$

step 2:-

Non-Null productions

$$E \rightarrow TE'$$

$$E' \rightarrow +TE'$$

$$T \rightarrow FT'$$

$$T' \rightarrow *FT'$$

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

Null productions

$$E' \rightarrow \epsilon$$

$$T' \rightarrow \epsilon$$

step 3: - find FIRST() for non-null production
and FOLLOW for null production

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

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

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

$$\text{FIRST}(T') = \{ * \}$$

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

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

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

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

Step 4:- Using FIRST() & FOLLOW() computed above fill the parsing table.

Parsing table:-

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

→ The above parsing table does not contain multiple entries in single cell therefore the grammar given is in LL(1).

$$\begin{aligned} \text{ii)} \quad S &\rightarrow iEtS \mid iEtSeS \mid a \\ E &\rightarrow b \end{aligned}$$

$$N = \{ S, E \}$$

$$T = \{ i, t, e, a, b \}$$

Step 1: For Constructing LL(1) the grammar must be free from ambiguity, left recursion and left factor.

The above grammar contains left factor.
Eliminating left factor.

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

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

$$E \rightarrow b$$

$$N = \{ S, S', E \}$$

$$T = \{ i, t, e, a, b \}$$

Step 2:- Identify null & non-null productions

Non-Null

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

$$E \rightarrow b$$

$$S' \rightarrow eS$$

Null

$$S' \rightarrow \epsilon$$

Step 3:- Calculate FIRST() for non-null production & FOLLOW() for null productions.

$$\text{FIRST}(S) = \{i, a\}$$

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

$$\text{FIRST}(S') = \{e\}$$

$$\text{FOLLOW}(S') = \text{FOLLOW}(S)$$

$$= \{\text{FIRST}(S) - \epsilon\} \cup \text{FOLLOW}(S)$$

$$\text{FOLLOW}(S') = \{e, \$\}$$

Step 4:- Construct Parse table using above FIRST() & FOLLOW()

Parse table:-

	i	e	a	t	b	\$
S	$S \rightarrow iEtSS'$	-	$S \rightarrow a$	-	-	-
S'	-	$S' \rightarrow eS$ $S' \rightarrow e$	-	-	-	$S' \rightarrow e$
E	-	-	-	-	$E \rightarrow b$	

The parse table contains multiple entries $[S'e]$
Therefore it is not in LL(1).

* Driver Routine :-

→ It is used to check whether the string is accepted by the grammar or not.

→ It contains

- i) Stack
- ii) Input tape



→ Rules :-

i) If x is a top of stack & a is an input symbol if $x = a = \$$ then string is accepted.

ii) If $x = a \neq \$$, then pop x and advancing input pointer

iii) If $x \neq a$, then look into parse table and write RHS of the production in reverse order.

Q) Check whether the string
 $Id + Id * Id$ is accepted by

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

Using LL(1) parse table.

Note:- The LL(1) parse table is
 constructed above in LL(1) parser
 topic

Stack	Input String
\$E	Id + Id * Id \$
\$E' ↑ T	↑ Id + Id * Id \$
\$E' ↑ T' ↑ F	↑ Id + Id * Id \$
\$E' ↑ T' ↑ Id	↑ Id + Id * Id \$
\$E' ↑ T' ↑	↑ + Id * Id \$
\$E' ↑	↑ + Id * Id \$

\$E'T↑

\$E'T↑

\$E'T'F↑

\$E'T'id↑

\$E'T'↑

\$E'T'F*↑

\$E'T'F↑

\$E'T'id↑

\$E'T'↑

\$E'↑

\$↑

↑id*id\$

↑id*id\$

↑id*id\$

↑id*id\$

↑*id\$

↑*id\$

↑id\$

↑id\$

↑\$

↑\$

↑\$

→String is Accepted

Therefore the given string id+id*id is Accepted by the grammar.

* Bottom-up Parser:- (Shift-Reduce Parser)

→ It can be defined as an attempt to reduce the input string w to the start symbol of grammar by tracing out the rightmost derivations of w in reverse order.

Eg:- $E \rightarrow E + E \mid E * E \mid id$

$id * id + id$

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

$E * E + id$

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

$E + E$ ($E \rightarrow E + E$)

$E \rightarrow \text{start symbol}$

→ The intermediate steps in the process of derivation is sequential form.
If the derivation is right most then it is right sequential form.

→ The task in bottom up parser is to find the substring that would be reduced by appropriate variable such a substring is called as handle.

→ The Operations Involved in bottom-up parser are

- i) Shift
- ii) Reduce
- iii) Accept
- iv) Error

Eg:- $E \rightarrow E + E \mid E * E \mid id$
 $w = id + id * id.$

Stack	Input	Action
\$	$id + id * id \$$	
$\$ id$	$+ id * id \$$	Shift
$\$ E$	$+ id * id \$$	Reduce
$\$ E +$	$id * id \$$	Shift
$\$ E + id$	$* id \$$	Shift
$\$ E + E$	$* id \$$	Reduce
$\$ E$	$* id \$$	Reduce
$\$ E *$	$id \$$	Shift
$\$ E * id$	$\$$	Shift

$\$ E * E \quad \$$ reduce
 $\$ E \quad \$$ reduce

↳ start symbol.

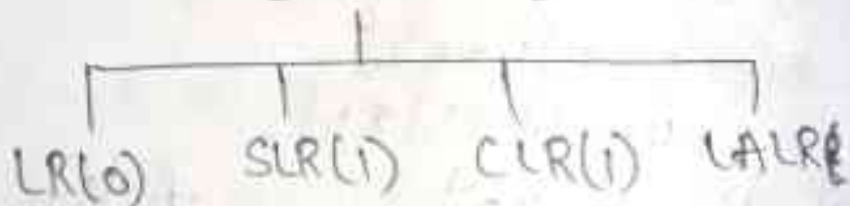
∴ The string is Accepted.

→ Bottom up parser is used for both ambiguous and unambiguous grammar

→ Bottom-up parser



LR Parsers



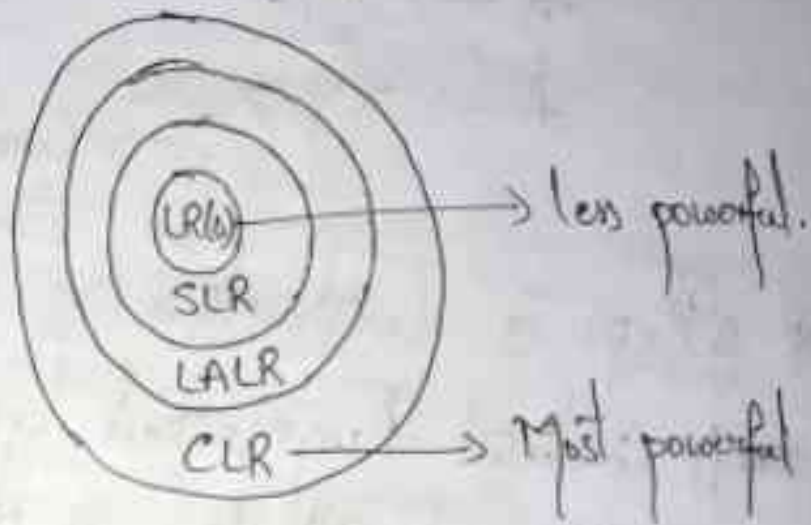
* L-R parser

→ LR parsers are bottom up parser which handles unambiguous grammar

L → left to right Scanning g/lp

R → RMD in reverse order

Types:- SLR - Simple LR
CLR - Canonical LR
LALR - Look Ahead LR



→ In LR parsing, the rows corresponds to items and columns are divided into two parts: 1) Terminal part called as Action part
2) Variable part called goto part.

→ Items: - Introducing a $\cdot$ (dot) on RHS forms an item. It signifies what is parsed and what is yet to parse.

Eg:- $A \rightarrow \cdot XYZ \rightarrow x$ is parsed
 $A \rightarrow X \cdot YZ \rightarrow xy$ is parsed
 $A \rightarrow XY \cdot Z \rightarrow$ & z is yet to parse
 $A \rightarrow XYZ \cdot \rightarrow$ Complete item

* Steps in Construction of LR parser:-

- 1) Construct Augmented grammar
- 2) Canonical Collection of LR(0) items
- 3) Construction of DFA
- 4) Construction of Parse table
- 5) Driver Routine.

Augmented Grammar:-

$\rightarrow$ Including an additional start symbol forms an Augmented grammar.

Rules for Canonical Collection of LR:-

- i) Write all production in Augmented grammar in LR(0) & place dot at beginning of RHS
- ii) If there is a variable after dot place all productions of that variable in that item & place a Dot at beginning of RHS.

Note:- In SLR(1) of parsing table, we use FOLLOW() to place reduce operations.

Parse table

Q) Construct SLR(1) for below grammar

i) $E \rightarrow E + T \mid T$

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

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

Step 1:- Augmented grammar

$$E' \rightarrow \cdot E$$

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

$$E \rightarrow \cdot T$$

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

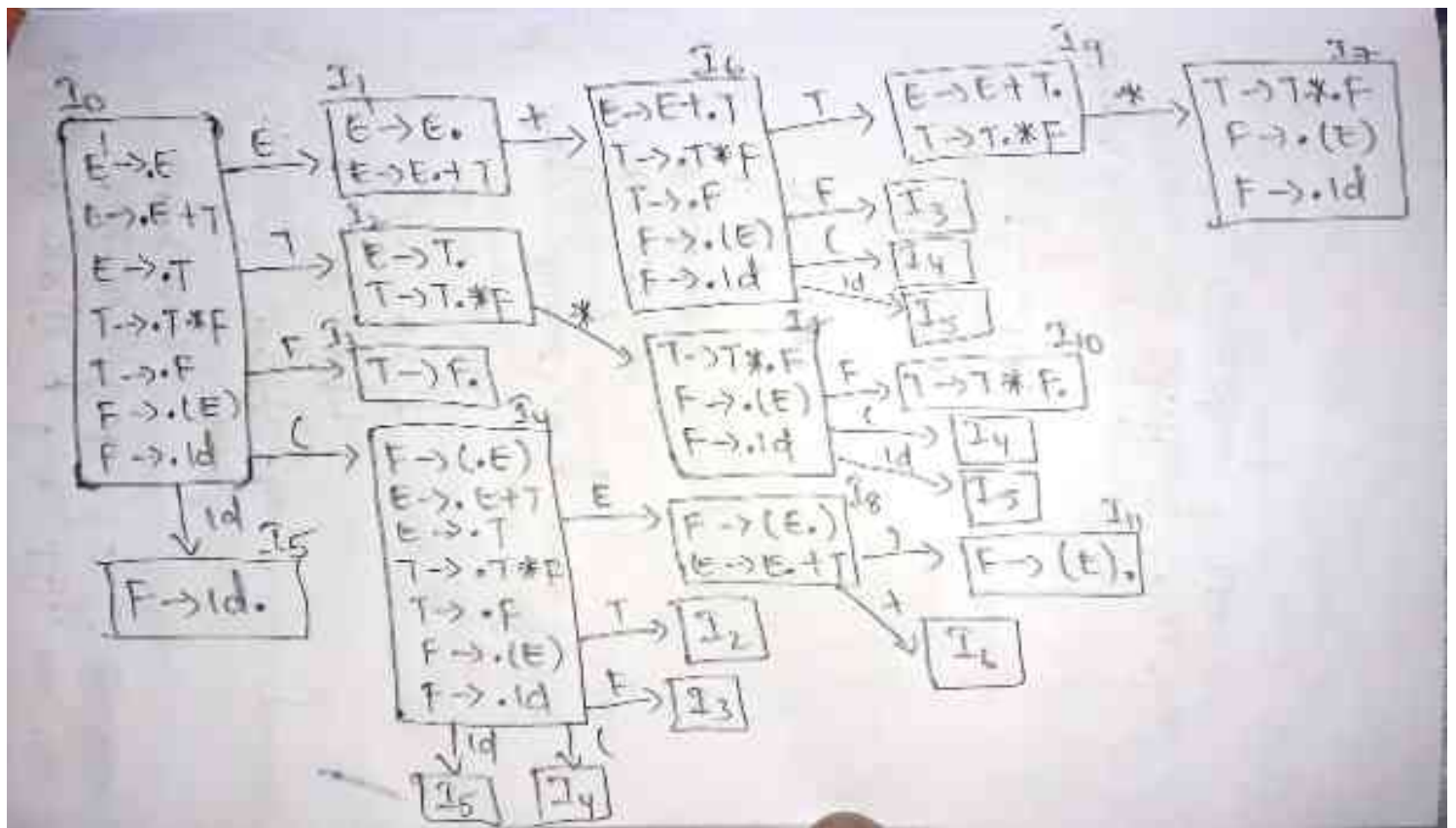
$$T \rightarrow \cdot F$$

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

$$F \rightarrow \cdot id$$

Step 2 & 3:-

Canonical Collection of LR(0) items
& Constructing DFA



Action		Goto	
		Items in which it is completely parsed	FOLLOW(LHS)
$E \rightarrow E + T$	θ_{11}	I_9	FOLLOW(E) $= \{ \$, +,) \}$
$E \rightarrow T$	θ_{12}	I_2	
$T \rightarrow T * F$	θ_{13}	I_{10}	FOLLOW(T) $= \{ \$, +, *,) \}$
$T \rightarrow F$	θ_{14}	I_3	
$F \rightarrow (E)$	θ_{15}	I_{11}	FOLLOW(F) $= \{ \$, +, *,) \}$
$F \rightarrow Id$	θ_{16}	I_5	

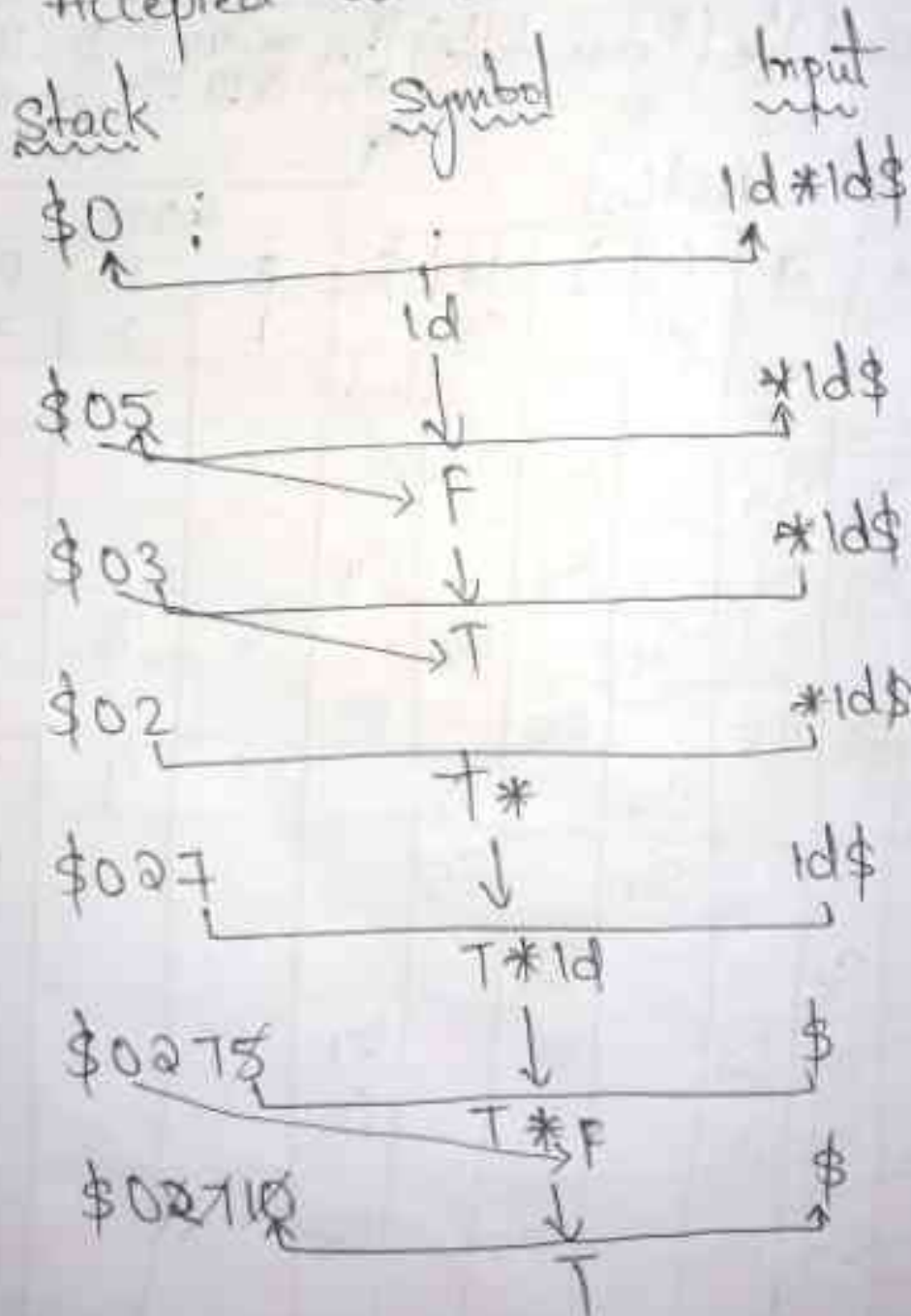
Parsing table (Parse table)
 For reduce $\xrightarrow{use}$ FOLLOW()
 For shift $\xrightarrow{use}$ Canonical Items

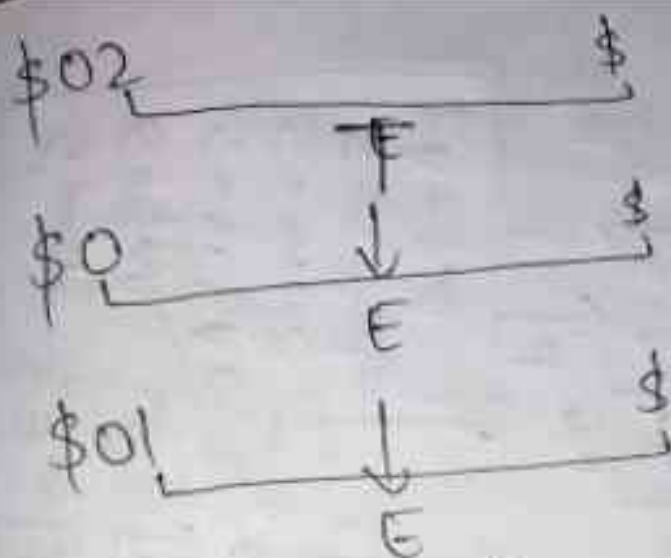
Action						GOTO			
	+	*	(	)	Id	\$	E	T	F
			S_4		S_5		1	2	3
I_0						Accepted			
I_1	S_6								
I_2	θ_{12}	S_7		θ_{12}		θ_{12}			
I_3	θ_{14}	θ_{14}		θ_{14}		θ_{14}			
I_4			S_4		S_5		8	2	3
I_5	θ_{16}	θ_{16}		θ_{16}		θ_{16}			
I_6			S_4		S_5			9	3
I_7			S_4		S_5				10
I_8	S_6			S_{11}					
I_9	θ_{11}	S_7		θ_{11}		θ_{11}			
I_{10}	θ_{13}	θ_{13}		θ_{13}		θ_{13}			
I_{11}	θ_{15}	θ_{15}		θ_{15}		θ_{15}			

→ The above grammar is in
 LR(1), since no cell
 contains multiple entries.

* Driver Routine:-

→ Using the above parse table we
 check whether the string is
 Accepted or not.





$\therefore$ String is Accepted
(Since, I_1 on $\$$ is Accepted)

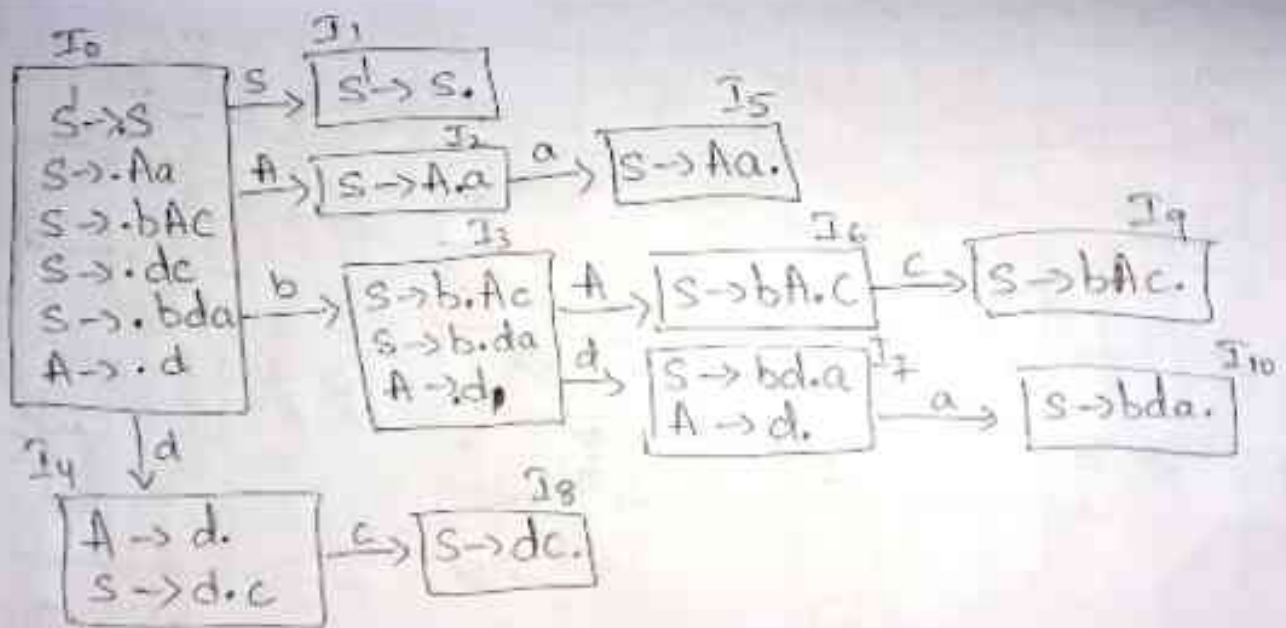
ii) $S \rightarrow A|bAc|dc|bda$
 $A \rightarrow d$. Check whether the grammar is LR(0) or not.

Ans $S \rightarrow A|bAc|dc|bda$
 $A \rightarrow d$

Augmented Grammar:-

$S' \rightarrow .S$
 $S \rightarrow .Aa$
 $S \rightarrow .bAc$
 $S \rightarrow .dc$
 $S \rightarrow .bda$
 $A \rightarrow .d$

Canonical Collection of LR(0) items.



		Item in which it is completely parsed	FOLLOW(LAHS)
$S \rightarrow Aa$	δ_{11}	I_5	FOLLOW(s) = \$ \$ y
$S \rightarrow bAc$	δ_{12}	I_9	
$S \rightarrow dc$	δ_{13}	I_8	FOLLOW(A) = \$ a, c y
$S \rightarrow bda$	δ_{14}	I_{10}	
$A \rightarrow d$	δ_{15}	$I_4 \& I_7$	

Parse table:-

	a	b	c	d	\$	S	A
I_0		S_3		S_4		1	2
I_1					Accept		
I_2	S_5						
I_3				S_7			6
I_4			S_8				
I_5					δ_{11}		
I_6			S_9				
I_7	δ_{15}, S_{10}		δ_{15}				
I_8					δ_{13}		
I_9					δ_{12}		
I_{10}					δ_{14}		

$\therefore I_7, a_1$ contains both shift & reduce operation
therefore the given grammar is not in SLR(1).

* Canonical LR parser:- CLR(1)

→ In CLR(1) parser we use symbols called look aheads to place reduce operation in parse table.

→ With each production rule we associate a look ahead.

→ The look ahead for start symbol is \$

→ If $A \rightarrow \alpha \cdot B \beta$, a then all

productions starting ^{look ahead} with B will have look ahead as FIRST(βa)

→ In SLR(1) we used LR(0) canonical items, but in CLR(1) and LALR(1) we make use of LR(1) canonical items.

$$\{ \text{LR(1) item} = \text{LR(0) item} + \text{look ahead} \}$$

→ The look ahead is used to determine that where we place the final item.

→ The look ahead always add \$ symbol for augmented production.

*) Conflicts in LR parser:-

- i) Shift Reduce Conflict
- ii) Reduce Reduce Conflict

Q) Construct CLR(1) parse table for the following grammar.

- i) $S \rightarrow CC$
 $C \rightarrow eC$
 $C \rightarrow d$

sol I) Augmented Grammar:-

$S' \rightarrow .S, \$$

$S \rightarrow .CC, \$$

$C \rightarrow .eC, e/d$

$C \rightarrow .d, e/d$

FIRST(C)

$C \rightarrow \gamma \text{FIRST}(C, \$)$

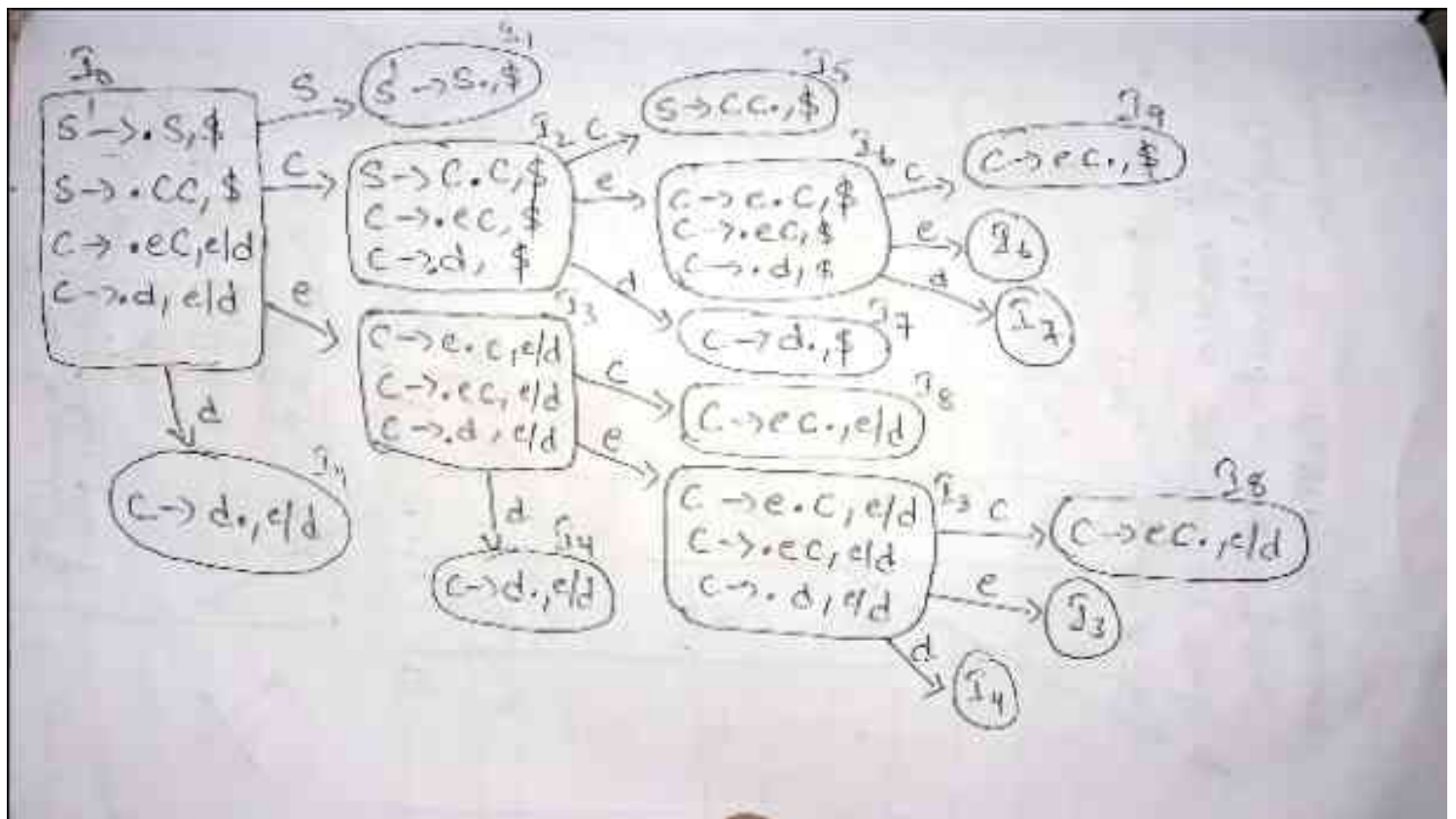
$\gamma \text{FIRST}(e, \$)$

$\gamma \text{FIRST}(d, \$)$

$\rightarrow e/d$

II) Construct LR(1) Canonical Items

III) for the above grammar



		Item in which completely passed
$S \rightarrow CC$	θ_1	I_5
$C \rightarrow eC$	θ_2	I_8, I_9
$C \rightarrow d$	θ_3	I_4, I_7

Parse table:-

	ACTION			GOTO	
	e	d	\$	S	C
I_0	S_3	S_4		1	2
I_1			Accept		
I_2	S_6	S_7			5
I_3	S_3	S_4			8
I_4	θ_3	θ_3			
I_5			θ_1		
I_6	S_6	S_7			9
I_7			θ_3		
I_8	θ_2	θ_2			
I_9			θ_2		

The above grammar is CLR(1) grammar
 Since Parse table does not contain
 shift reduce conflict.

$$\begin{aligned} \text{ii)} \quad S &\rightarrow L = R \mid R \\ L &\rightarrow * R \mid d \\ R &\rightarrow L \end{aligned}$$

Sol Augmented Grammar

$$S \rightarrow \cdot L = R, \$$$

$$S \rightarrow \cdot R, \$.$$

$$L \rightarrow \cdot * R$$

$$L \rightarrow \cdot d$$

$$R \rightarrow \cdot L$$

Look Ahead for L :-

$$S \rightarrow \cdot L = R, \$$$

It is in the form $A \rightarrow \alpha \cdot B \beta \mid a$ $B = L \quad \beta = = R \quad a = \\$

$$\begin{aligned} \text{Look Ahead of } L &= \text{FIRST}(\beta a) \\ &= = / \$ \end{aligned}$$

Look Ahead for R :-

$$S \rightarrow \cdot R, \$$$

$$\text{Look Ahead for R} = \$$$

Now,

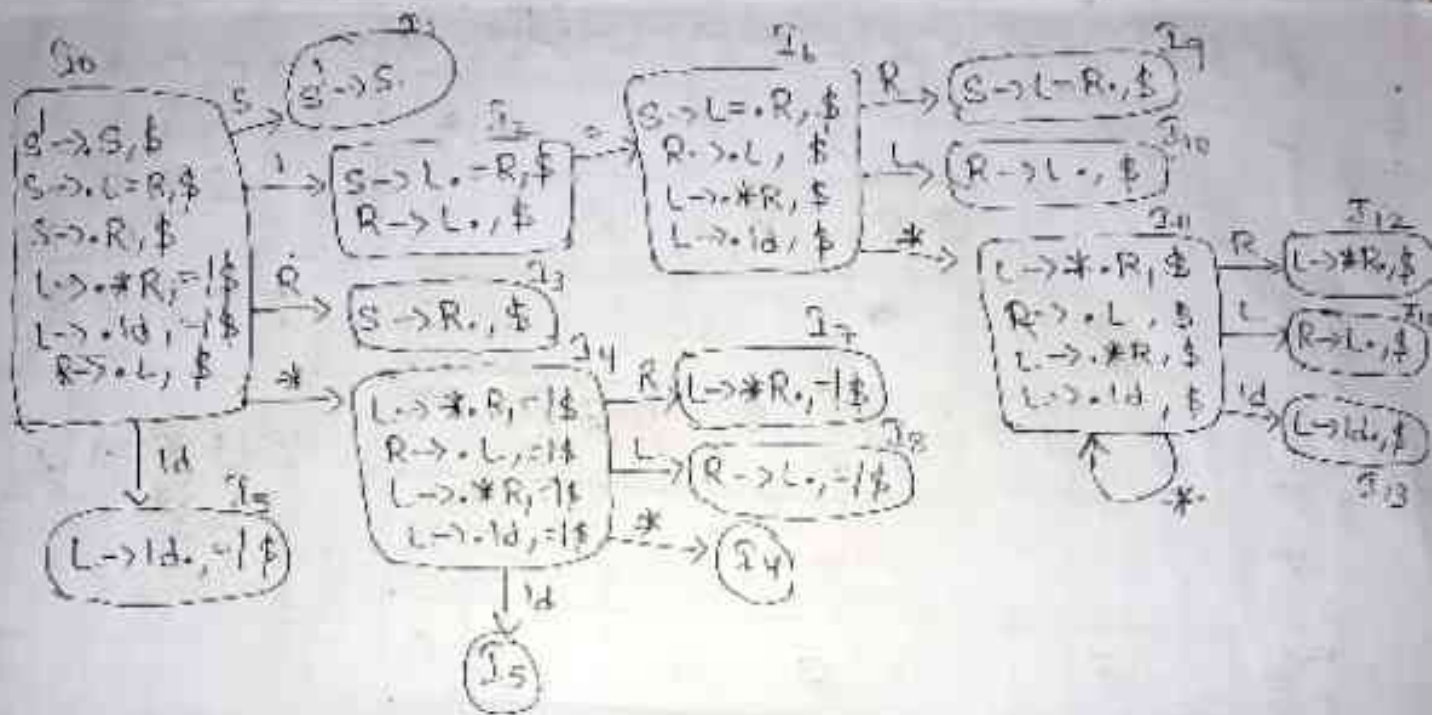
$$S \rightarrow \cdot L = R, \$$$

$$S \rightarrow \cdot R, \$$$

$$L \rightarrow \cdot * R, = / \$$$

$$L \rightarrow \cdot d, = / \$$$

$$R \rightarrow \cdot L, \$$$



		Item in which it is completely processed
$S \rightarrow L = R$	θ_{11}	I_9
$S \rightarrow R$	θ_{12}	I_3
$L \rightarrow * R$	θ_{13}	I_4, I_{12}
$L \rightarrow Id$	θ_{14}	I_5, I_{13}
$R \rightarrow L$	θ_{15}	I_2, I_8, I_{10}

Parse Table:-

	*	=	Id	\$	S	L	R
I_0	S_4		S_5		1	2	3
I_1				Accept			
I_2		S_6		θ_{15}			
I_3				θ_{12}			
I_4	S_4		S_5			8	7
I_5		θ_{14}		θ_{14}			
I_6	S_{11}					10	9
I_7		θ_{13}		θ_{13}			
I_8		θ_{15}		θ_{15}			
I_9				θ_{11}			
I_{10}				θ_{15}			
I_{11}	S_{11}		S_{13}			10	12
I_{12}				θ_{13}			
I_{13}				θ_{14}			

The above grammar is LR(1)
 Since there is no shift reduce
 conflict.

* Look-Ahead LR parser: - LALR(1)

→ In LALR we combine the items having common core but different look aheads.

→ LALR(1) parsing is same as CLR(1) parsing, only difference in the parsing table.

Q) Construct LALR(1) parse table for the following grammar.

i) $S \rightarrow CC$
 $C \rightarrow eC$
 $C \rightarrow d$

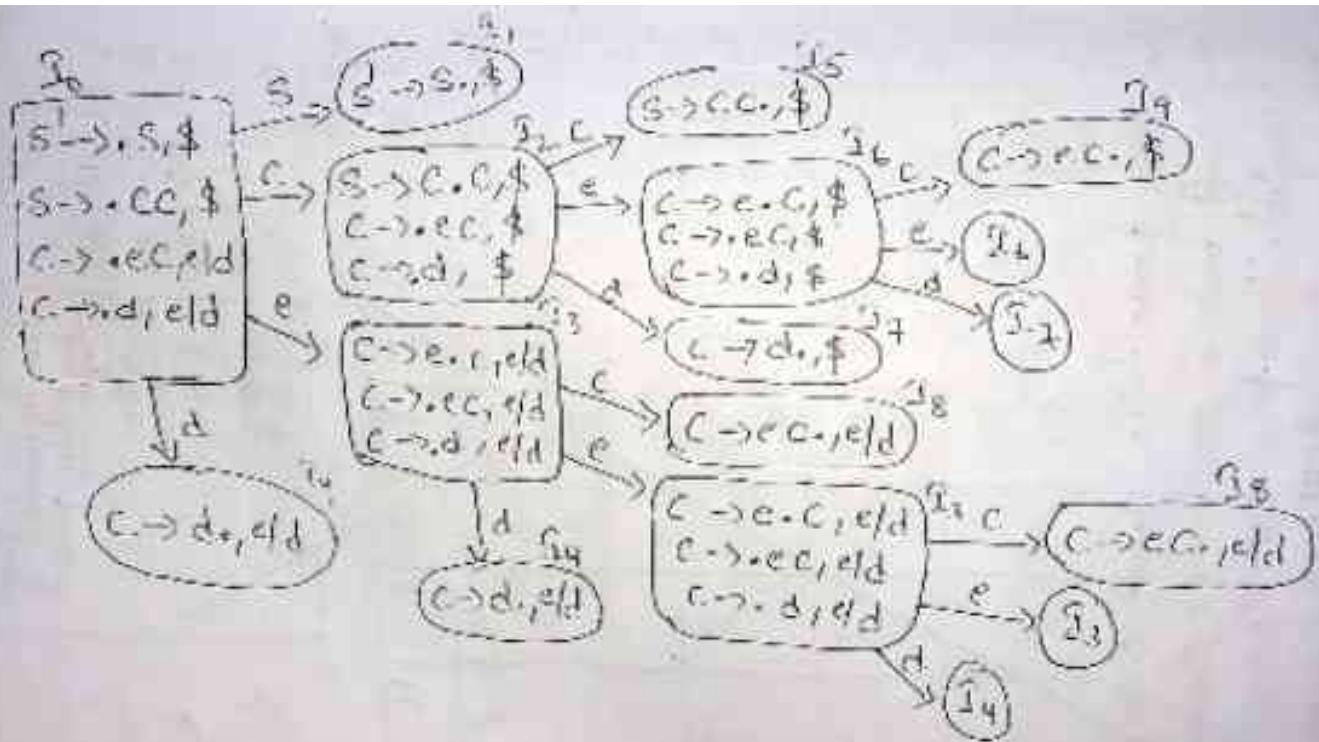
sol Augmented Grammar: -

$S' \rightarrow \cdot S, \$$
 $S \rightarrow \cdot CC, \$$
 $C \rightarrow \cdot eC$
 $C \rightarrow d$

FIRST(C)

$C \rightarrow \gamma \text{FIRST}(C\$)$
 $\rightarrow e/d$

Canonical Collection of LR(1) items



		Item by which completely passed
$s \rightarrow cc$	θ_1	I_5
$c \rightarrow ec$	θ_2	I_8, I_9
$c \rightarrow d$	θ_3	I_4, I_7

→ The above Canonical LR(1) items have common core part but different look aheads. Therefore combine common core part and form Construct parse table.

$$\rightarrow I_3 \cup I_6 \cong I_{36} = I_3$$

$$I_4 \cup I_7 \cong I_{47} = I_4$$

$$I_8 \cup I_9 \cong I_{89} = I_8$$

Parse table

	a	e	d	\$	S	C
I_0		S_3	S_4		1	2
I_1				Accept		
I_2		S_3	S_4			5
I_3		S_3	S_4			
I_4		r_3	r_3	r_3		
I_5				r_1		
I_8		r_2	r_2	r_2		

The above parse table doesn't contain shift reduce conflict. Therefore The grammar is LR(1).

$$\begin{aligned}
 \text{ii) } S &\rightarrow Aa/bAc/Bc/bBa \\
 A &\rightarrow d \\
 B &\rightarrow d
 \end{aligned}$$

Augmented grammar.

$$\begin{aligned}
 S' &\rightarrow \cdot S, \$ \\
 S &\rightarrow \cdot Aa, \$ \\
 S &\rightarrow \cdot bAc, \$ \\
 S &\rightarrow \cdot Bc, \$ \\
 S &\rightarrow \cdot bBa, \$ \\
 A &\rightarrow d, a \\
 B &\rightarrow d, c
 \end{aligned}$$

Look Ahead for A:-

$$S \rightarrow \cdot Aa, \$ \quad (A \rightarrow \alpha \cdot B\beta, a)$$

$$\begin{aligned}
 &\text{FIRST}(\beta a) \\
 \text{Look Ahead for A : } &\text{FIRST}(a\$) \\
 &= a
 \end{aligned}$$

look Ahead for B:-

$$S \rightarrow \cdot Bc, \$$$

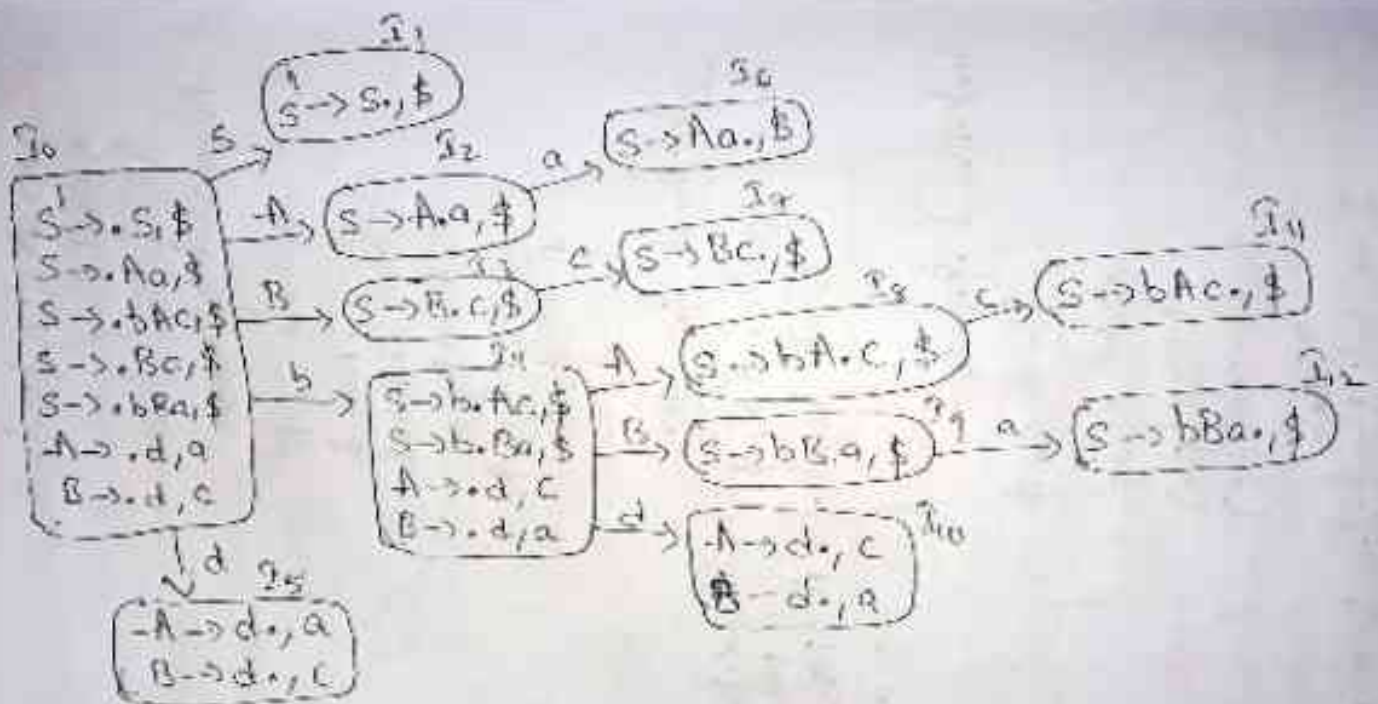
$$A \rightarrow \alpha \cdot B\beta, a$$

$$\rightarrow \text{FIRST}(\beta a)$$

$$\rightarrow \text{FIRST}(c\$)$$

$$\rightarrow ac$$

$$\text{Look Ahead for B} = ac$$



Items in which it is completely parsed.		
$S \rightarrow Aa$	$\$1$	I_6
$S \rightarrow bAc$	$\$2$	I_{11}
$S \rightarrow Bc$	$\$3$	I_7
$S \rightarrow bBa$	$\$4$	I_{12}
$A \rightarrow d$	$\$5$	$I_5 \& I_{10}$
$B \rightarrow d$	$\$6$	$I_5 \& I_{10}$

I_5 & I_{10} has same core part & different look aheads so we combine I_5 & I_{10}

$$I_5 \cup I_{10} \approx I_{510} = I_5$$

Parse table:-

	a	b	c	d	\$	S	A	B
I_0		S_4		S_5		1	2	3
I_1				Accept				
I_2	S_6							
I_3			S_7					
I_4				S_5			8	9
I_5	$\$5, \6		$\$5, \6					
I_6				$\$1$				
I_7				$\$3$				
I_8			S_{11}					
I_9	S_{12}							
I_{11}				$\$2$				
I_{12}				$\$4$				

The above parse table contains multiple entries i.e. reduce reduce conflict. So the grammar is not in LALR(1).

Note:- The above grammar is not in LALR(1) but it is in CLR(1).

* Parser Generator YACC:-

- YACC stands for Yet Another Compiler Compiler
- YACC is an automated tool to produce a parser for given grammar
- YACC is a program designed to compile a LALR(1) grammar.
- The input of YACC is the rule or grammar and output is C program.
- YACC file has .y extension.
- YACC working:-

Parser.y → YACC Compiler → y.tab.c

y.tab.c → C Compiler → a.out

i/p tokens → a.out → output
(parse)

→ YACC input file is divided into three parts.

/* definitions */
%%
/* rules */
%%
/* auxiliary routines */

- The definition part includes information about tokens used in syntax definition.
- The rules part contains grammar definition in modified BNF form. Actions is C code in {} and can be embedded inside.
- The Auxiliary routines part is only C code. It includes definitions for every function needed in rules part.
- It contains the main() function and main function must call yyparse().

* Syntax Error handling & Recovery:

→ In this phase of compilation, all possible errors made by the user are detected and reported to the user in the form of error message.

→ This process of locating errors are reported to user is called Error handling process.

→ Functions of Error handler:-

- 1) Detection
- 2) Reporting
- 3) Recovery

→ Common Errors that can occur

- i) lexical errors include misspellings of identifiers, keyword or operators.
- ii) Syntactic errors include misplaced semicolons or extra or missing braces.
- iii) Semantic errors include type mismatches between operators & operands.
- iv) logical errors can be anything from incorrect reasoning.

→ The precision of parsing methods allow syntactic errors to be detected very efficiently.

→ Several parsing methods, such as LL & LR methods, detect an error as soon as possible.

→ The error handler in a parser has goals that are simple to state but challenging to realize:

i) Report the presence of errors clearly & accurately.

ii) Recover from error quickly enough to detect subsequent errors.

iii) Add minimal overhead to the processing of correct programs.

* Error Recovery Strategies:-

→ Once the error is detected following strategies are used to recover from that error.

i) Panic Mode Recovery:-

→ In this method, successive characters from input are removed one at

a time until a designated set of synchronizing tokens is found. Synchronizing tokens are delimiters such as ; or y.

- Advantage of this method is that it's easy to implement and guarantees not to go in infinite loop.
- Disadvantage is that a considerable amount of input is skipped without checking it for additional errors.

2) Phrase-level Recovery:-

- In this method, when a parser encounters an error, it performs necessary corrections on remaining input so that the rest of the input statements allow the parser to parse ahead.
- The correction can be deletion of extra semicolons, replacing comma by semicolon or inserting missing semicolon.
- While performing correction, utmost care should be taken for not going in infinite loop.
- Disadvantage is that it finds difficult to handle situation where actual errors occurred before point of detection.

3) Error Productions:-

- If User has knowledge for common errors that can be encountered then, these errors can be incorporated by augmenting the grammar with error productions that generate erroneous constructs.
- If this is used then, during parsing, appropriate error message can be generated and parsing can be continued.
- It is difficult to maintain

4) Global Correction:-

- The parser examines the whole program and tries to find out the closest match for it which is error free.
- The closest match program has less number of insertions, deletions and changes of tokens to recover from erroneous input.
- Due to high time and space complexity, this method is not implemented practically.

* Parsing Ambiguous grammar:-

→ LR parser can be used to parse ambiguous grammar. LR parser resolves the conflicts (shift/reduce or reduce/reduce) in parsing table of ambiguous grammar based on certain rules (precedence or associativity) of grammar.

WYAAC:-

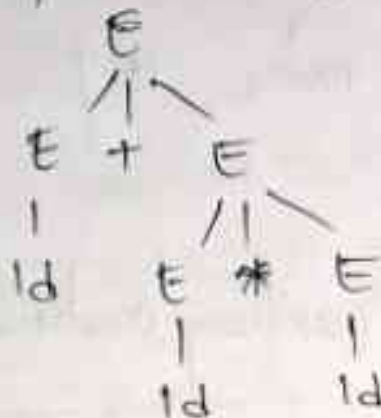
→ When Shift-reduce conflict occurs, it is resolved by giving priority to shift move over reduce move. If the string is accepted for shift move, then reduce move is removed, otherwise shift move is removed.

→ When Reduce-reduce conflict occurs, it is resolved by giving priority to first reduce move over second reduce move. If the string is accepted for first reduce move, then second reduce move is removed, otherwise first reduce move is removed.

UNIT:03 Syntax-Directed Translation

* CFG:- $E \rightarrow E + E \mid E * E \mid Id$

Parse tree for above grammar



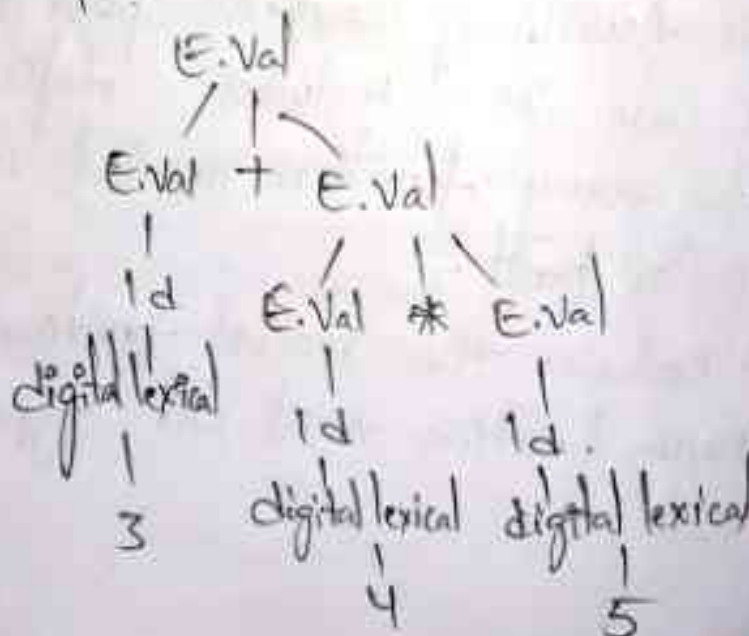
Syntax directed definition

$E.val \rightarrow E.val + E.val$

$E.val \rightarrow E.val * E.val$

$E.val \rightarrow Id$

Annotated parse tree



* Syntax Directed Translation:-

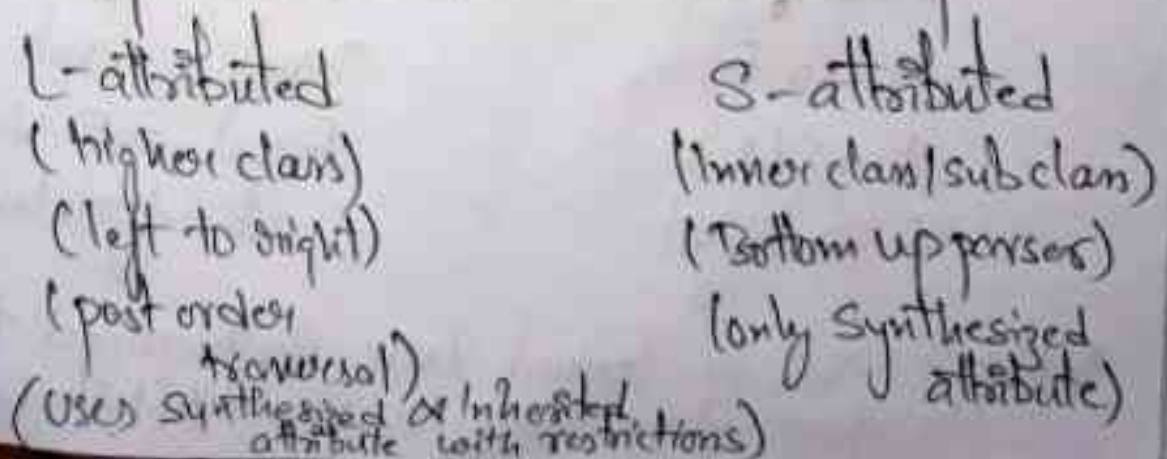
→ Syntax Directed Translation are augmented rules to the grammar that facilitates semantic analysis. SDT involves passing information bottom-up / top-down the parse tree in the form of attributes attached to nodes.

→ In SDT along with grammar we associate some informal notation and these notations are called Semantic Rules.

Grammar + Semantic rule = SDT

→ In SDT, every non-terminal can get one or more than one attribute or sometime 0 attribute depending upon the type of attribute. The value of these attribute is evaluated by semantic rules associated with production rule.

Syntax Directed Translation Definition



Attributes

Synthesized
(Evaluated from
its children)

Inherited
(Evaluated from
parent, sibling)

* Example of Synthesize Attribute

$$E \rightarrow E_1 + E_2 \quad \{ E.val = E_1.val + E_2.val \}$$

* Example of Inherited Attribute

$$A \rightarrow XYZ \quad \{ Y.val = 2 * A.val \}$$

* Annotated parse tree :-

→ A parse tree showing the values of attribute at each node is called as Annotated parse tree

→ The process of computing the attributes value at nodes is called annotating of the parse tree

Dependency Graph:-

→ A Graph representing dependency b/w the attributes is the dependency graph.

* Syntax Directed Definition :- Example
(Synthesized)

Production

$L \rightarrow E n$

$E \rightarrow E + T$

$E \rightarrow T$

$T \rightarrow T * F$

$T \rightarrow F$

$F \rightarrow (E)$

$F \rightarrow \text{digit}$

Semantic Rule

$\text{print}(E.\text{val})$

$E.\text{val} = E.\text{val} + T.\text{val}$

$E.\text{val} = T.\text{val}$

$T.\text{val} = T.\text{val} * F.\text{val}$

$T.\text{val} = F.\text{val}$

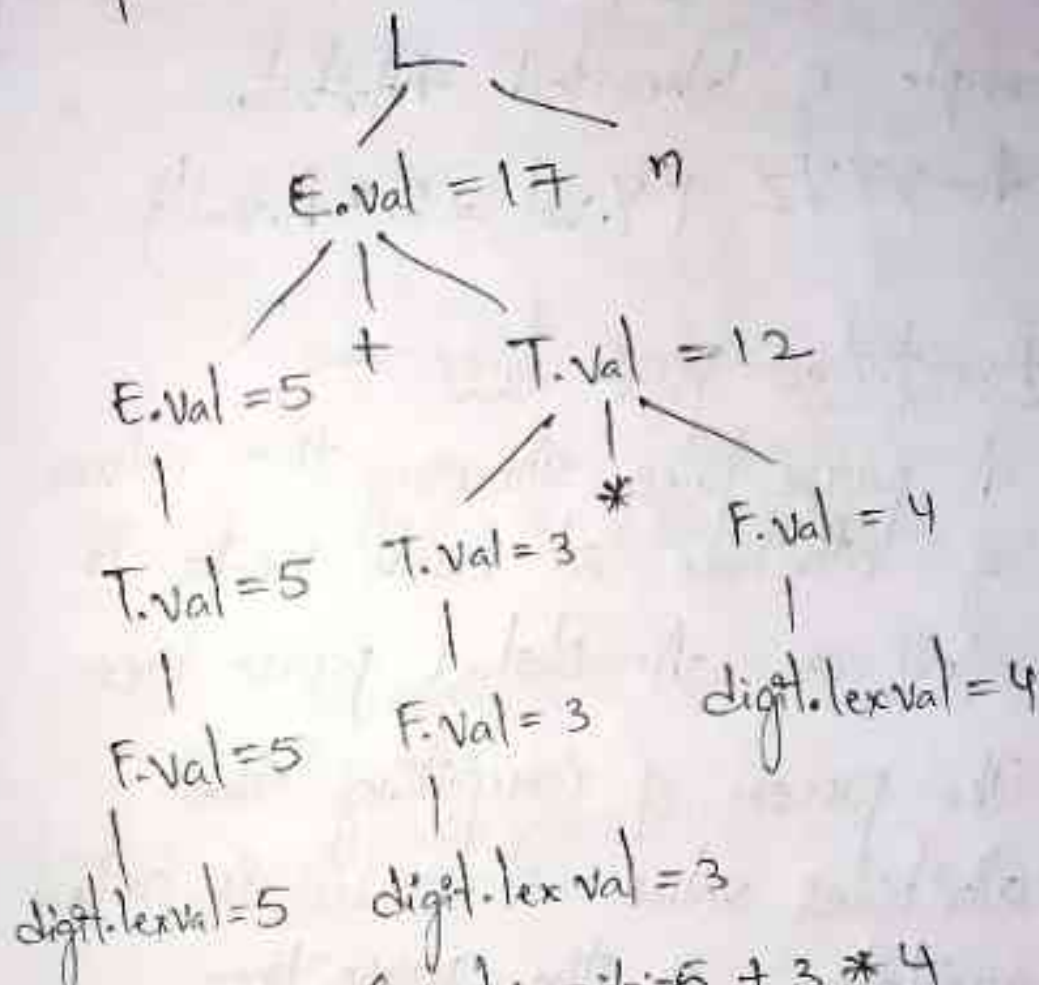
$F.\text{val} = E.\text{val}$

$F.\text{val} = \text{digit}.\text{lexval}$

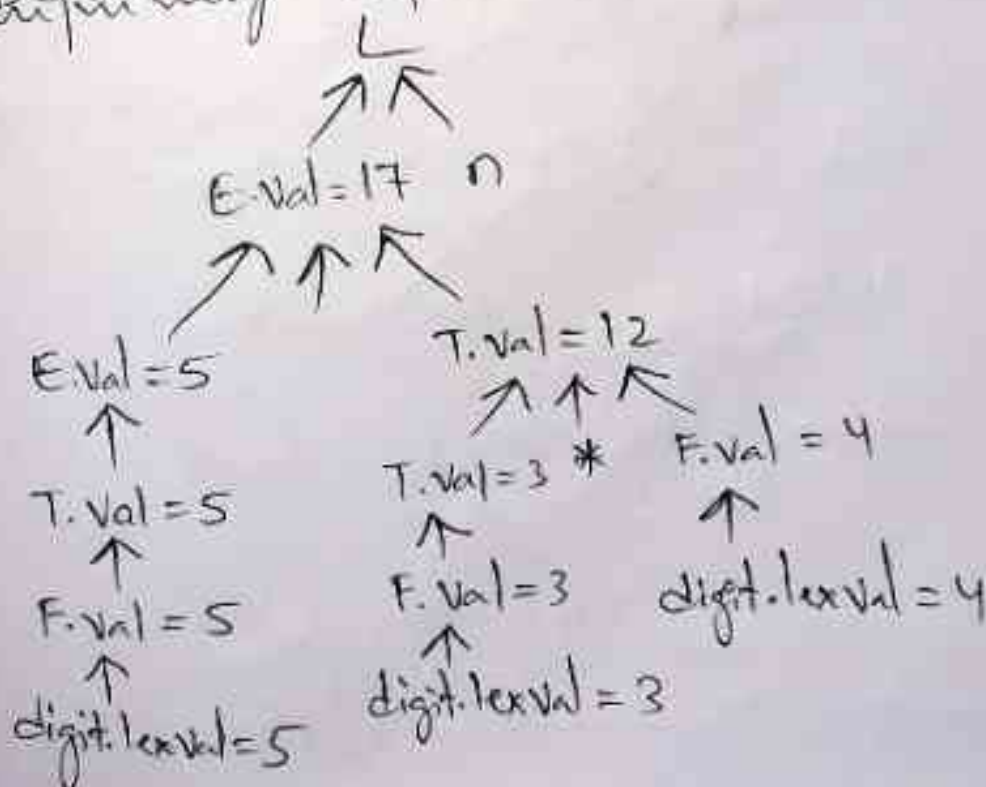
Note :- val & lexval are synthesized attribute
Therefore, above SDD is S-Attributed.

* Annotated parse tree

Input:- $5 + 3 * 4$



* Dependency Graph: - $5 + 3 * 4$



* Syntax Directed Definition : (Inherited) Attribute

(i) Production

$D \rightarrow TL$

$T \rightarrow \text{int}$

$T \rightarrow \text{real}$

$L \rightarrow L, id$

$L \rightarrow id$

Semantic Rule

$L.in = T.type$

$T.type = \text{integer}$

$T.type = \text{real}$

$L.in = L.in$

$addtype(id.entry, L.in)$

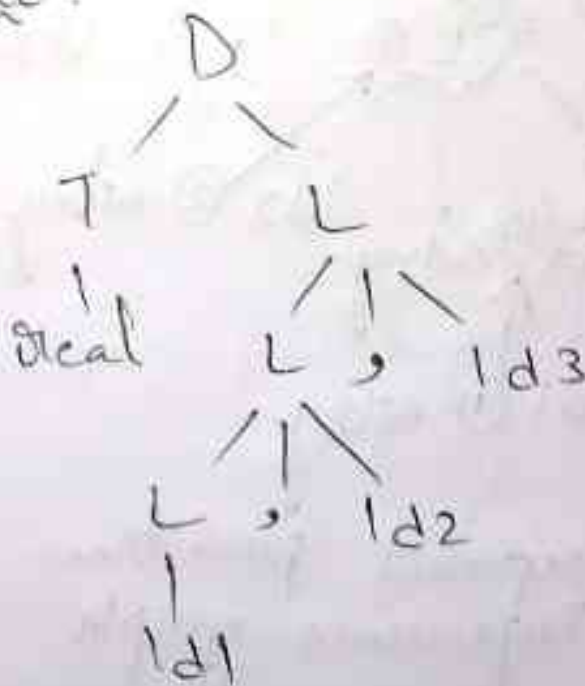
$addtype(id.entry, L.in)$

→ Symbol T is associated with synthesized attribute type

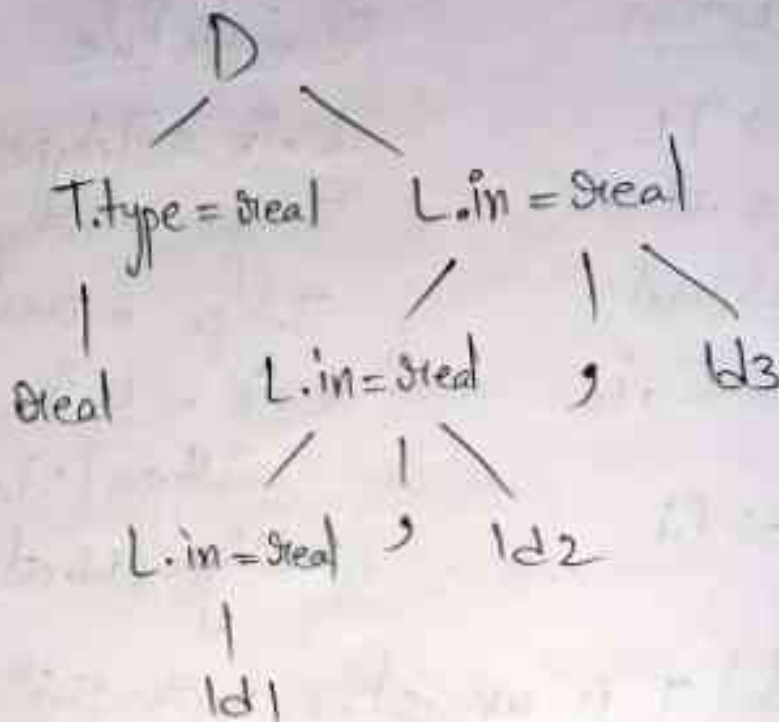
→ Symbol L is associated with an inherited attribute in

Input:- real p, q, r

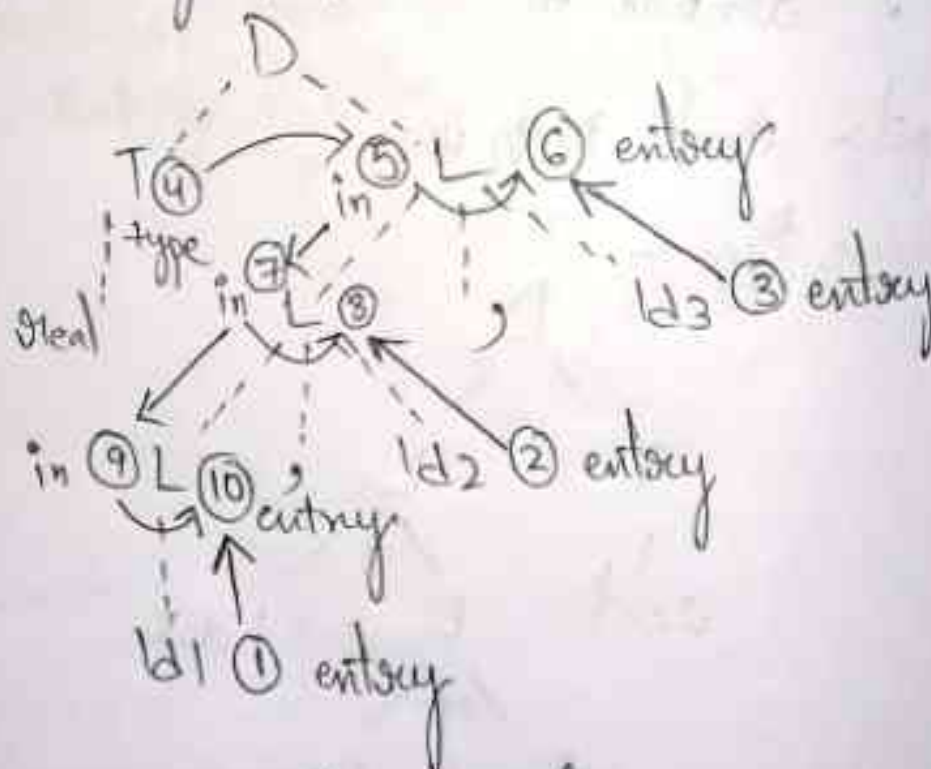
Parse tree :-



Annotated parse tree:-



Dependency graph:-



----- : represent parse tree
 —————> : Dependence graph

(ii) Production

Semantic Rule

1) $T \rightarrow FT'$

$$T'.in = F.val$$
$$T.val = T'.syn$$

2) $T' \rightarrow * FT_1'$

$$T_1'.in = T'.in * F.val$$
$$T'.syn = T_1'.syn$$

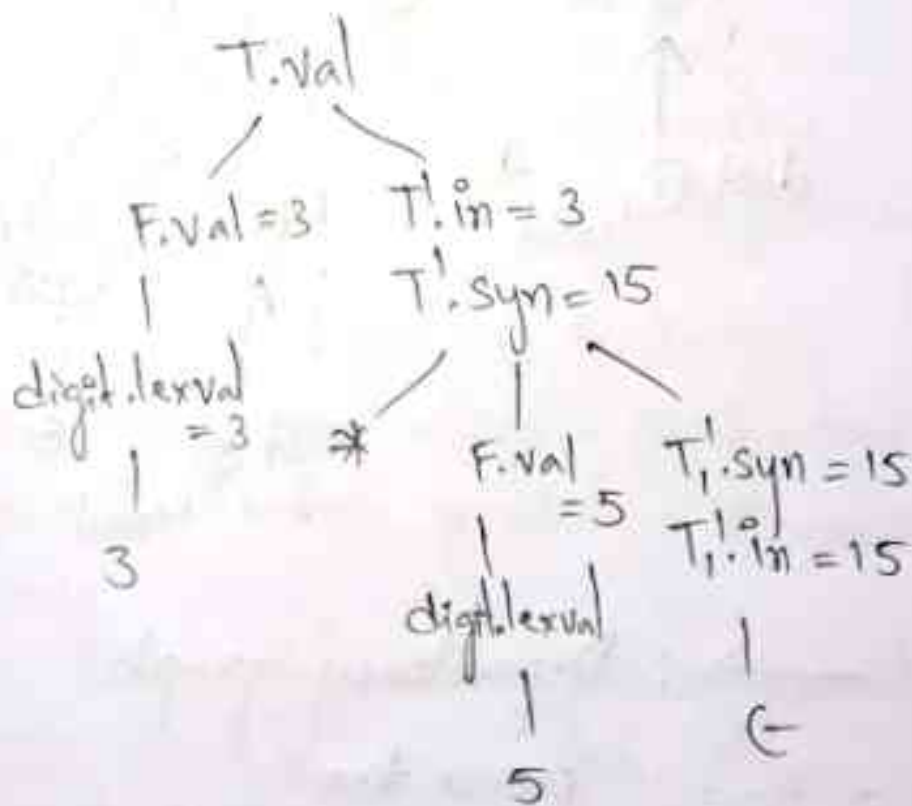
3) $T' \rightarrow \epsilon$

$$T'.syn = T'.in$$

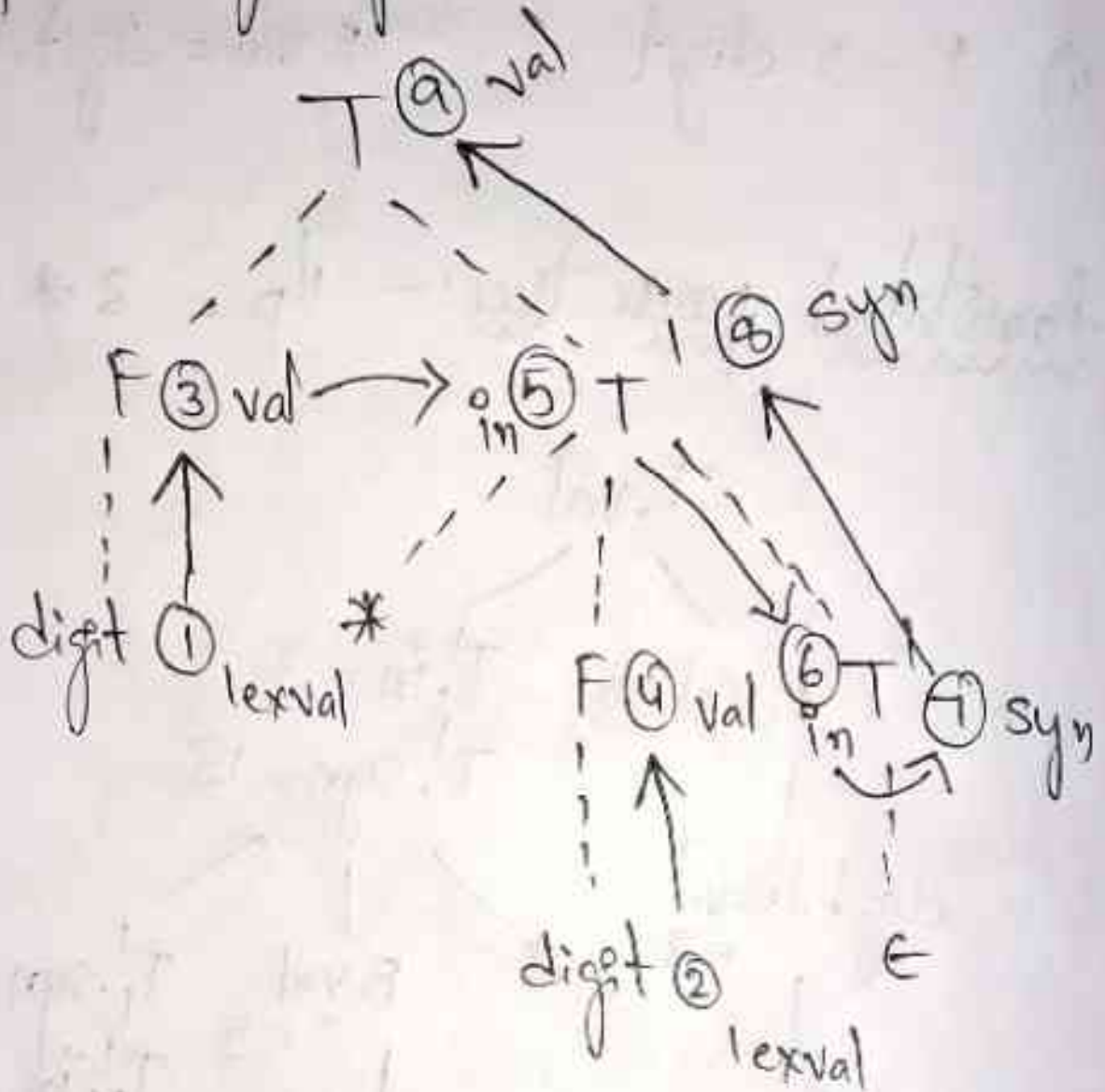
4) $F \rightarrow digit$

$$F.val = digit.lexval$$

* Annotated parse tree:- $1/p = 3 * 5$



Dependency graph:-



→ : Dependency graph

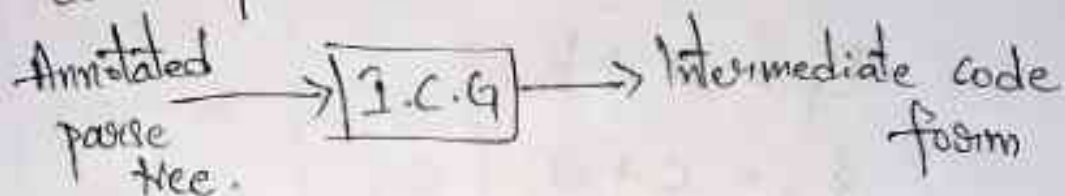
--- : parse tree.

* Applications of Syntax Directed Translation

- It is used for semantic analysis and SDT is basically used to construct parse tree with grammar and semantic action.
- In grammar, need to decide who has the highest priority will be done first and in Semantic Action, will decide type of action done by grammar.
- Applications:—
 - i) SDT are used for Evaluating Arithmetic Expressions
 - ii) Conversion from Infix to postfix
 - iii) Conversion from Infix to prefix
 - iv) It is also used in Binary to decimal conversion
 - v) In Counting number of Reduction
 - vi) It is used in creating a syntax tree
 - vii) SDT is used to generate Intermediate Code
 - viii) In Storing information into symbol table
 - ix) SDT is also used for type checking.

UNIT:-04 Intermediate Code generation

- Input to Intermediate code generation phase is Annotated parse tree and output is any Intermediate code form.



→ Syntax tree:-

It is a tree representation of an expression. In syntax tree nodes are operators and leaves are operands.

→ 3-Address code:-

The instructions expressed by three addresses. The various 3-address code representation are

- i) Quadruple
- ii) Triple
- iii) Indirect triple

Q) Write 3-address code for the expressions.

i) $y = (a + b) * (c + d)$

→ Three address code

$$t_1 = a + b$$

$$t_2 = c + d$$

$$t_3 = t_1 * t_2$$

$$y = t_3$$

Quadruple:-

Opri	arg1	arg2	Result
+	a	b	t_1
+	c	d	t_2
*	t_1	t_2	t_3
=	t_3	-	y

Tuple:-

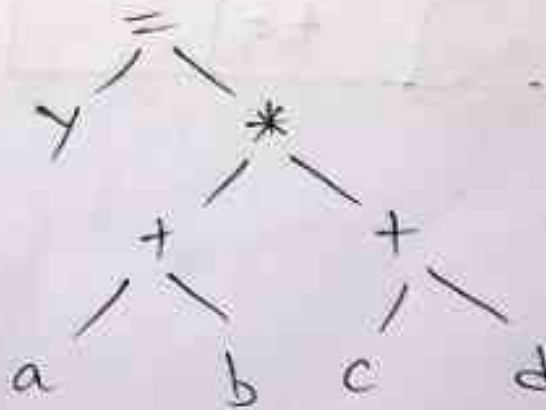
→ In tuple representation, we don't use temporary variable

	op	arg1	arg2
(0)	+	a	b
(1)	+	c	d
(2)	*	(0)	(1)
(3)	=	2	Y

Indirect tuple:-

Address	Pointer to tuple
(0)	100
(1)	101
(2)	102
(3)	103

Syntax tree:-



ii) $a = b * -c + b * -c$

↳ unary minus

→ 3-address code:-

$$t_1 = -c$$

$$t_2 = b * t_1$$

$$t_3 = -c$$

$$t_4 = b * t_3$$

$$t_5 = t_2 + t_4$$

$$a = t_5$$

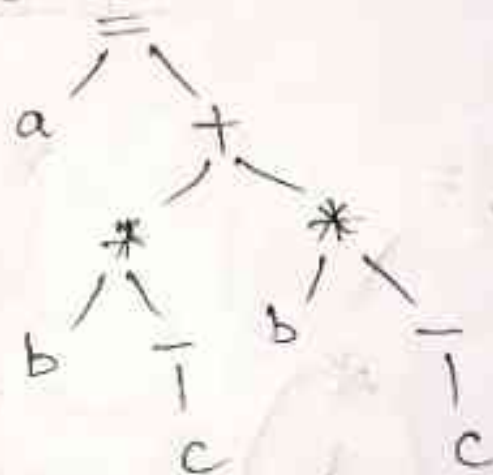
Quadruple:-

Opri	arg1	arg2	Result
-	c		t ₁
*	b	t ₁	t ₂
-	c		t ₃
*	b	t ₃	t ₄
+	t ₂	t ₄	t ₅
=	t ₅		a

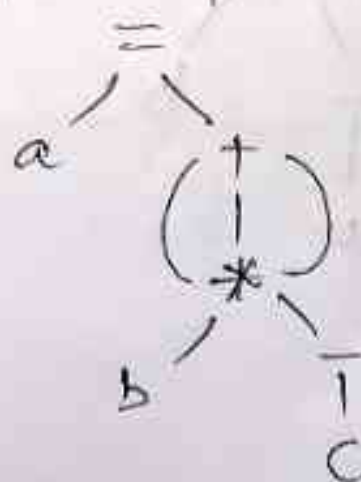
Tuple:-

	op ₁	arg ₁	arg ₂
(0)	-	c	(0)
(1)	*	b	(0)
(2)	-	c	(2)
(3)	*	b	(2)
(4)	+	(1)	(3)
(5)	=	a	(4)

Syntax tree:-



Directed Acyclic Graph (DAG):-



In DAG we combine common sub expressions

iii) DAG for i) $x = -a + b * -a + b$

→ 3-address code

$$t_1 = -a$$

$$t_2 = t_1 + b$$

$$t_3 = -a$$

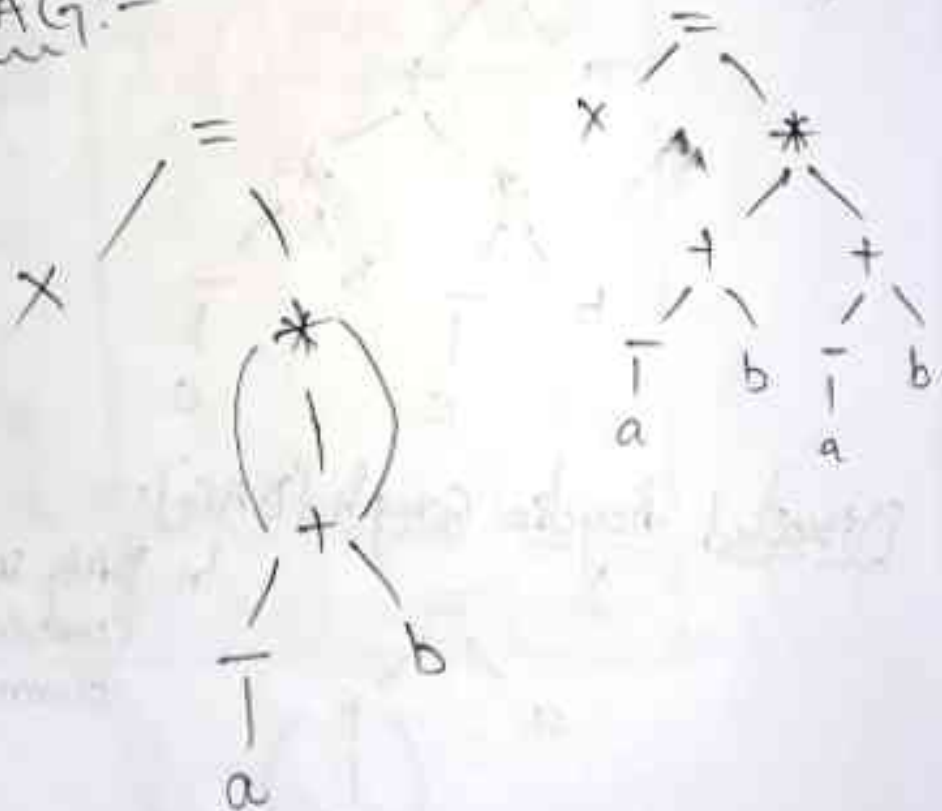
$$t_4 = t_3 + b$$

$$t_5 = t_2 * t_4$$

$$x = t_5$$

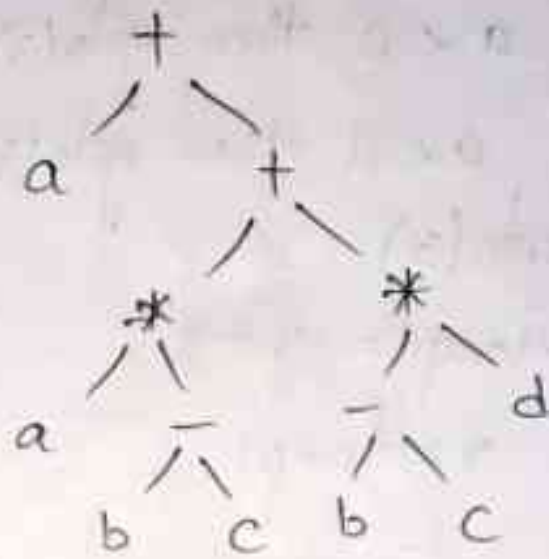
Syntax tree

DAG:-

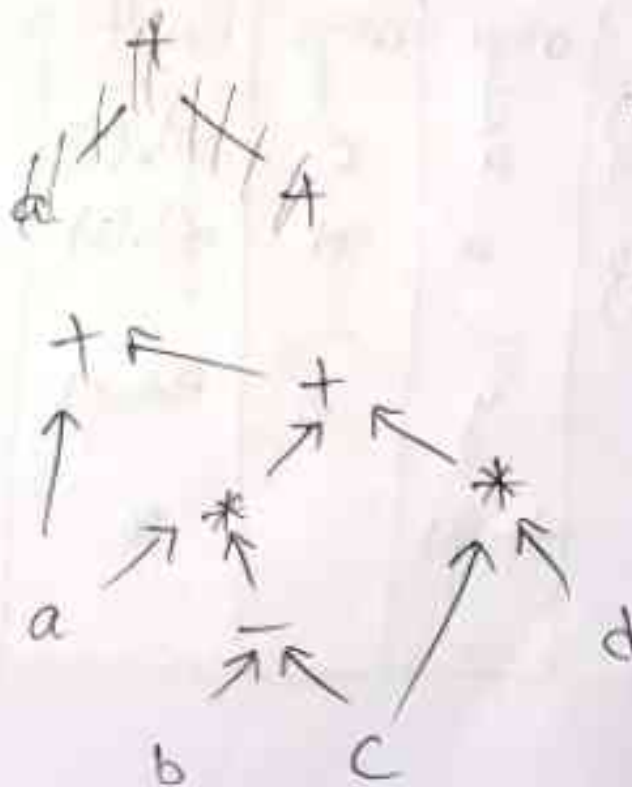


ii) DAG for $a + a * (b - c) + (b - c) * d$

→ Syntax tree :-



DAG :-



Q) If $a < c \parallel a < d$ then $x = y + z$.
 Give quadruple, triple, indirect triple

→ (0) if $a < c$ then goto(3)
 (1) if $a < d$ then goto(3)
 (2) goto(5)
 (3) $temp1 = y + z$
 (4) $x = temp1$
 (5) —

Quadruple:-

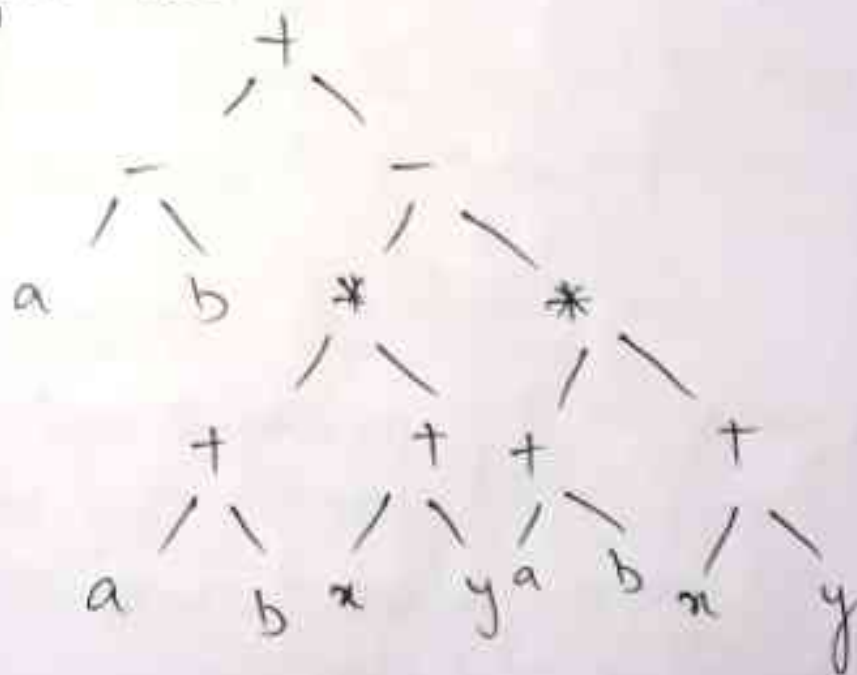
	op ₁	arg ₁	arg ₂	Result
(0)	less than ($<$)	a	c	goto(3)
(1)	less than ($<$)	a	d	goto(3)
(2)	—	—	—	—
(3)	+	y	z	temp1
(4)	=	temp1		x
(5)	—	—	—	—

Table:-

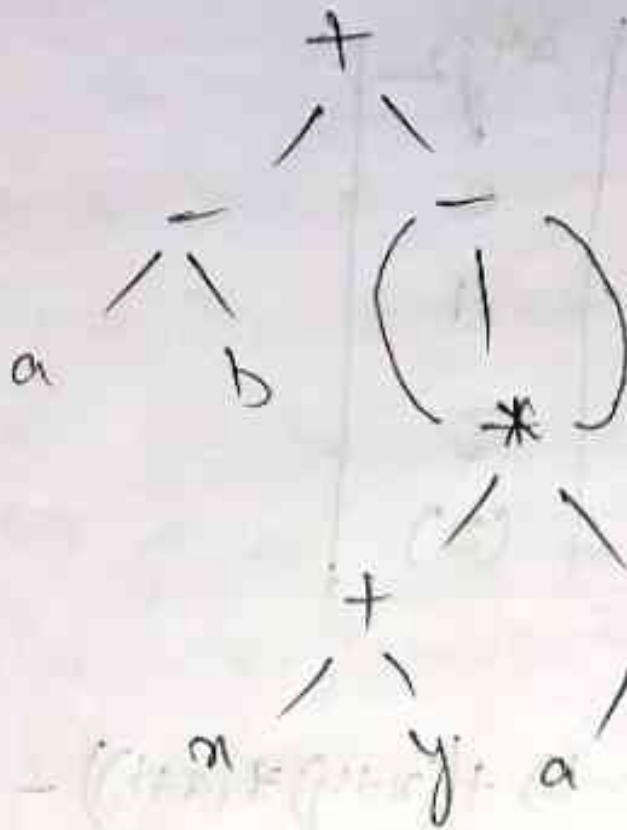
	op _{or}	areg ₁	areg ₂
10)	<	a	c
11)	<	a	d
(2)	+	y	3
(3)	=	x	(2)

a) DAG for $((a-b) + ((x+y) * (a+b)) - ((a+b) * (x+y)))$

→ Syntax tree



DAG:-



$(a + b) * (a + b)$ DAG

* Type Expression :-

→ The idea is to associate each language construct with an expression ~~data~~ describing its type and so called type expressions.

→ Type expressions are defined inductively from basic types and constants using type constructors.

> A basic type is a type expression
Ex:- int, float, char

> Typename is a type expression
int rollno.

> A type constructor applied to type expression is a type expression.

→ Type Expression for Array :-

The type expression for an array is given as $\text{Array}[I, T]$

$I \rightarrow$ Index of array

$T \rightarrow$ data type of array

Ex:- `int a[20];`

Type expression $\rightarrow \text{Array}[0-19, \text{int}]$

→ Type Expression for Cartesian product:-

Type Expression for Cartesian product is given by $T_1 \times T_2$, where T_1 & T_2 are data type

→ Struct:- (Record).

The struct given by keyword struct is also a type expression.

The type expression is given as product of types of field or member of structure.

Struct Student

```
{ int arr[20];  
  float marks;  
}
```

The type expression is given as

struct (arr * array[0-19, int] * marks * float)

→ Pointer:-

array(1-5, array(1-5, int))

The type expression for pointer is given as $\text{pointer}(T)$ where T is the data type

→ Function:-

The type expression for function is given as Domain $\rightarrow$ Range

int sum(int, int) { }

↓ √
Range Domain

Type Expression: $\text{int} \times \text{int} \rightarrow \text{int}$

a) Write type expression for an array of pointers to real where array index ranges from 10 to 100

→ Type Expression $\rightarrow \text{Array}[0-99, \text{Pointer}(\text{real})]$

a) Write type expression of function whose domain are function from integer to pointer of integer & whose ranges are records consisting of integers & characters

→ Typ Exp: $\text{int} \rightarrow \text{pointer}(\text{int}) \rightarrow \text{int} \times \text{char}$

* Intermediate code for procedure:-

Q) Suppose that a is an array of integers and that f is a function from integer to integer. Then the assignment

$$n = f(a[i])$$

Three-address code:-

1) $t_1 = i * 4$

2) $t_2 = a[t_1]$

3) Param t_2

4) $t_3 = \text{call } f, 1$

5) $n = t_3$

→ line ① & ② compute the value of expression $a[i]$ in t_2 .

→ line ③ makes t_2 an actual parameter for the call on line ④ of f with one parameter.

→ line ⑤ assign the value returned by function call.

Unit 3:

Symbol table

<https://www.geeksforgeeks.org/symbol-table-compiler/#:~:text=Symbol%20Table%20is%20an%20important,%2C%20classes%2C%20objects%2C%20etc.>

Unit 4:

1. Storage Organization

<https://www.javatpoint.com/storage-organization>

2. Activate Record

<https://www.javatpoint.com/activation-record>

3. Storage Allocation

<https://www.javatpoint.com/storage-allocation>

4. Access to non local data on stack

https://www.brainkart.com/article/Access-to-Nonlocal-Data-on-the-Stack_8172/

5. Parameter passing

<https://www.geeksforgeeks.org/runtime-environments-in-compiler-design/>