

unit - 1.

Introduction to Compilers.

The today's world depends up on programming languages, because all the softwares running on computers was written in some programming languages. But before a program can run, it first must be translated to a suitable form in which it can execute on computer.

The software that does this translation is called as compiler.

Language Pre processors.

Compiler is a program that can read a program in one language i.e., source language, and translate it into another language, known as target language.

An important role of the compiler is to report any errors in source program that it detects during the translation process.

Source Program



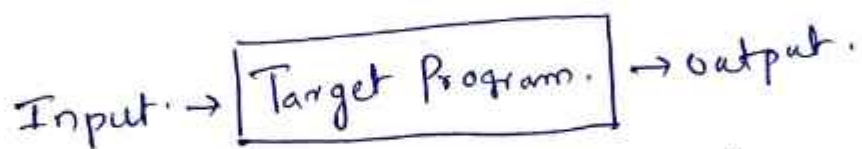
Compiler



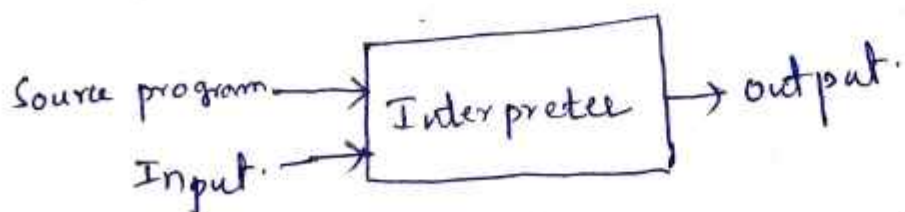
Target Program.

fig. Compiler.

If target program is executable machine program, it can take the inputs from user and produce output.



Interpreter is another common kind of language pre processor, instead of producing target program, interpreter directly executes the operations specified in the source program, and on inputs provided by user.



Machine language Target program produced by a compiler is much faster than interpreter. But interpreter can give better error diagnostics.

Then Compiler because it executes source program statement by statement.

In Addition to a Compiler several other programs may be needed to create an executable target program.

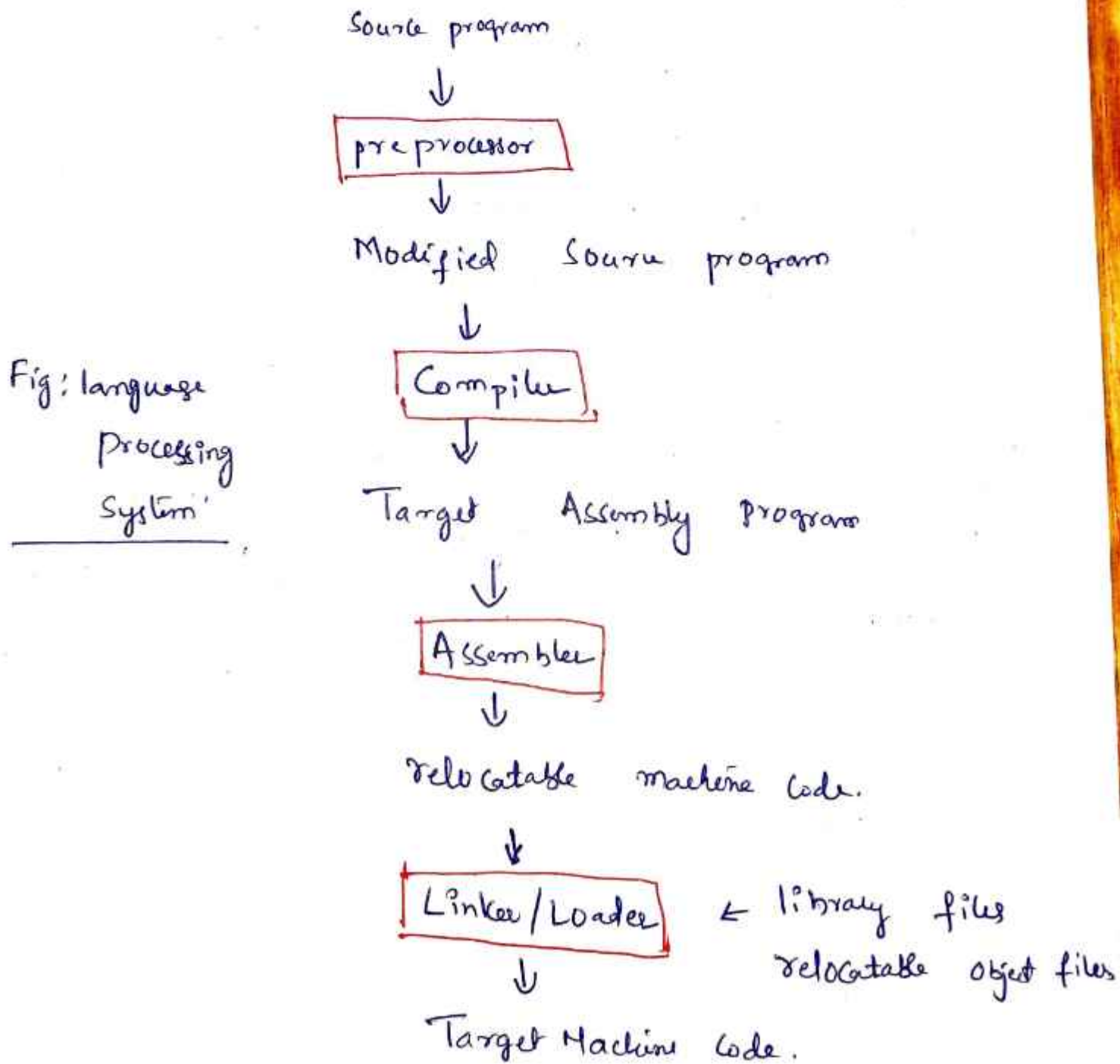
1. A source program may be divided into several modules and stored in separate files. Preprocessor collects these files into one unit. Preprocessor may also expand shortcuts, called as macros, into source language statements.

2. The modified source program is then fed to a Compiler. Compiler produces an Assembly language program as output. Because Assembly language program is easier to debug and to produce output.

3. The Assembly language is then processed by Assembler, That produces machine code as output.

4. Large programs are often compiled as pieces. The relocatable machine code has to be linked together. The linker links these files.

5. The loader puts together all executable object file into main memory for execution.



The Structure Of a Compiler.

Compiler consists of Two parts:

- i. Analysis.
- ii. Synthesis.

Analysis part breakes up Source program into pieces and imposes a grammatical structure on them. It then uses this structure to create intermediate code representation of source program.

If the Analysis part detects that Source program is syntactically or semantically weak, i.e; ~~not~~ following the rules of Grammar, Compiler provides informative messages to user in the format of errors, so that user can take corrective action.

This Analysis part collects the infoⁿ about source program, and stores in a Symbol table, which is passed along with intermediate code to Synthesis part.

The Synthesis part constructs desired target program from intermediate code and info in Symbol Table.

The Analysis part is often called as front end of the Compiler, The Synthesis part is called as backend.

Compilation process, operates as a sequence of phases, each of which transforms one representation to other.

A typical decomposition of Compiler into phases is shown in below figure.

Several phases may be grouped together. The Symbol table which stores the info about the entire source program is used by all phases of the Compiler.

Some Compiler have Optimization phase between the front end and backend.

The purpose of Code Optimization is to do some transformations in intermediate code so that backend can produce a better target Program.

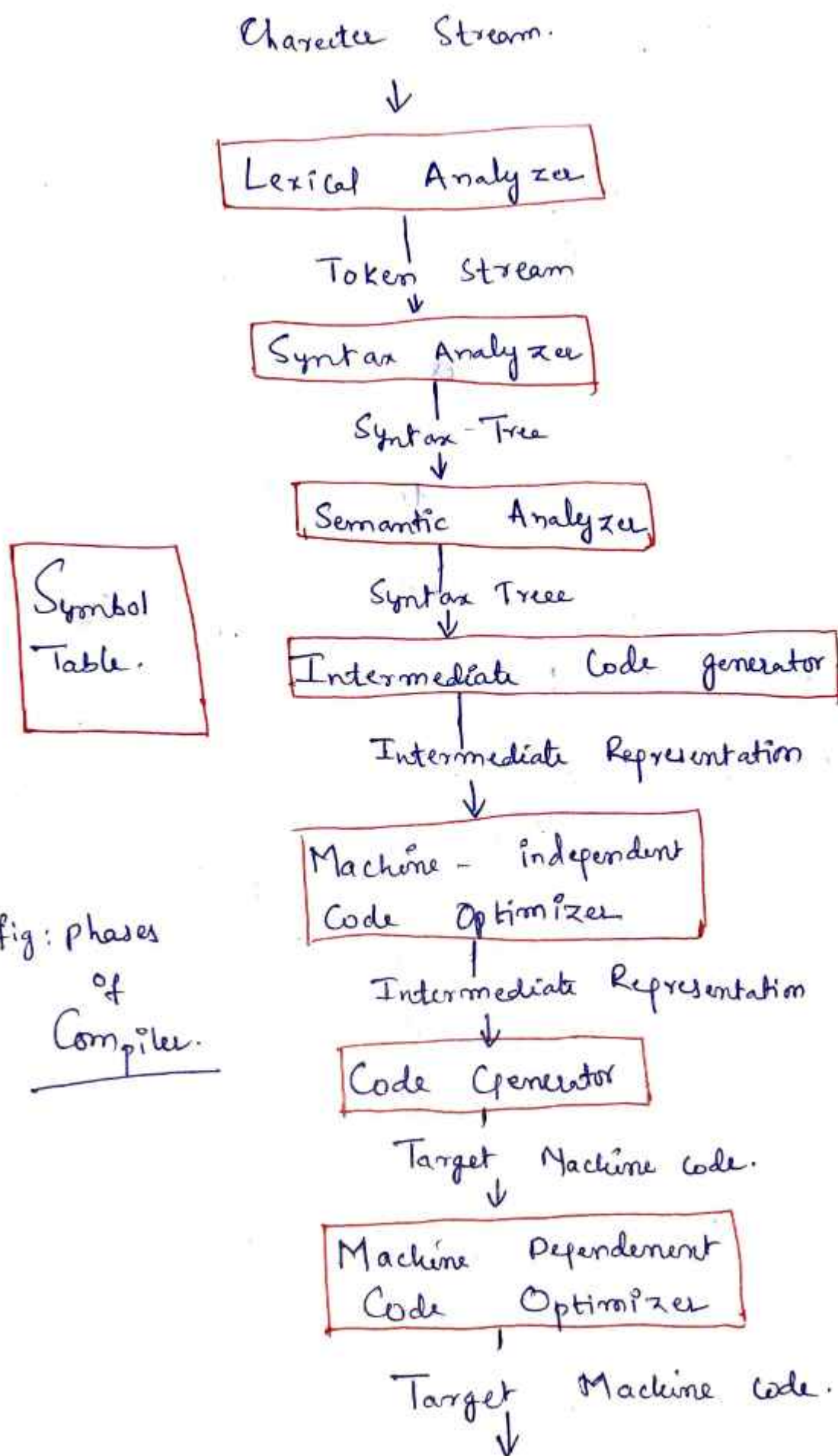


fig: Phases
of
Compiler.

Lexical Analysis

The first phase of compiler is called as Lexical Analysis or scanning. The Lexical Analyzer reads the stream of characters in a source program, and groups them into meaning full sequences called as lexemes.

For each lexeme, the lexical analyzer produces as op a token of the form

$\langle \text{token-name, Attribute value} \rangle$

Which is passed to next phase Syntax Analysis.

In token first component $\langle \text{token-name} \rangle$ is a symbol that is used during syntax analysis, Attribute value points to an entry in the Symbol table for this token. Info in Symbol table is needed for Semantic analysis and code generation.

For ex consider a source program
Contains assignment statement

$\text{position} = \text{initial} + \text{rate} * 60.$

So that

(5)

1. position is a lexeme, that will be mapped to a token $\langle id, 1 \rangle$, where id stands for identifier, 1. points to symbol table entry. This symbol table entry holds info about identifier such as its name and type.
2. The assignment symbol '=' is a lexeme which is mapped to token $\langle = \rangle$, it will not have attribute value, so second component is omitted.
3. initial is a lexeme that is mapped to a token $\langle id, 2 \rangle$, where 2 points to symbol table entry for initial.
4. + is a lexeme, mapped to token $\langle + \rangle$
5. rate is a lexeme, mapped to token $\langle id, 3 \rangle$, where 3 points to symbol table entry for rate.
6. * is a lexeme, mapped to token $\langle * \rangle$
7. 60 is a lexeme, mapped to token $\langle 60 \rangle$

After Lexical Analysis The assignment statement can be represented as sequence of tokens as below

Position = initial + rate * 60

$\langle id, 1 \rangle \langle = \rangle \langle id, 2 \rangle \langle + \rangle \langle id, 3 \rangle \langle * \rangle \langle 60 \rangle$

Syntax Analysis.

The second phase of compiler is Syntax analysis or parsing. Parser uses first component of token produced by Lexical analyzer to create a Tree like Intermediate representation that depicts Grammatical structure of token stream. This Tree structure is known as Syntax Tree.

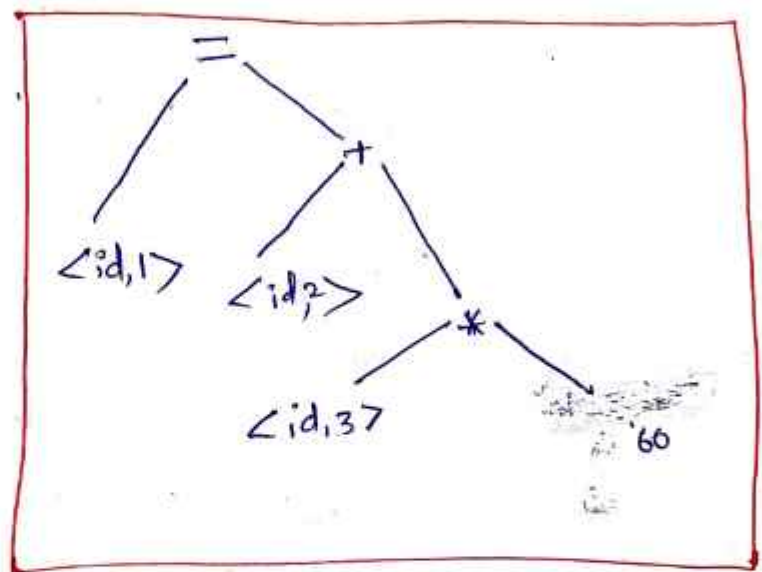
The Tree shows the order in which assignment is performed. For ex in the statement $position = initial + rate * 60$,

The Tree has a node $*$, with left child $\langle id, 3 \rangle$ as its left child, and 60 as its right child. The node $*$ tells that

We must first multiply rate by 60. (6)

→ The node + represent, we should add the result of multiplication, to the value of initial.

→ The root "=" denotes the right child result should be stored in identifier position.



Syntax Tree.

Semantic Analysis.

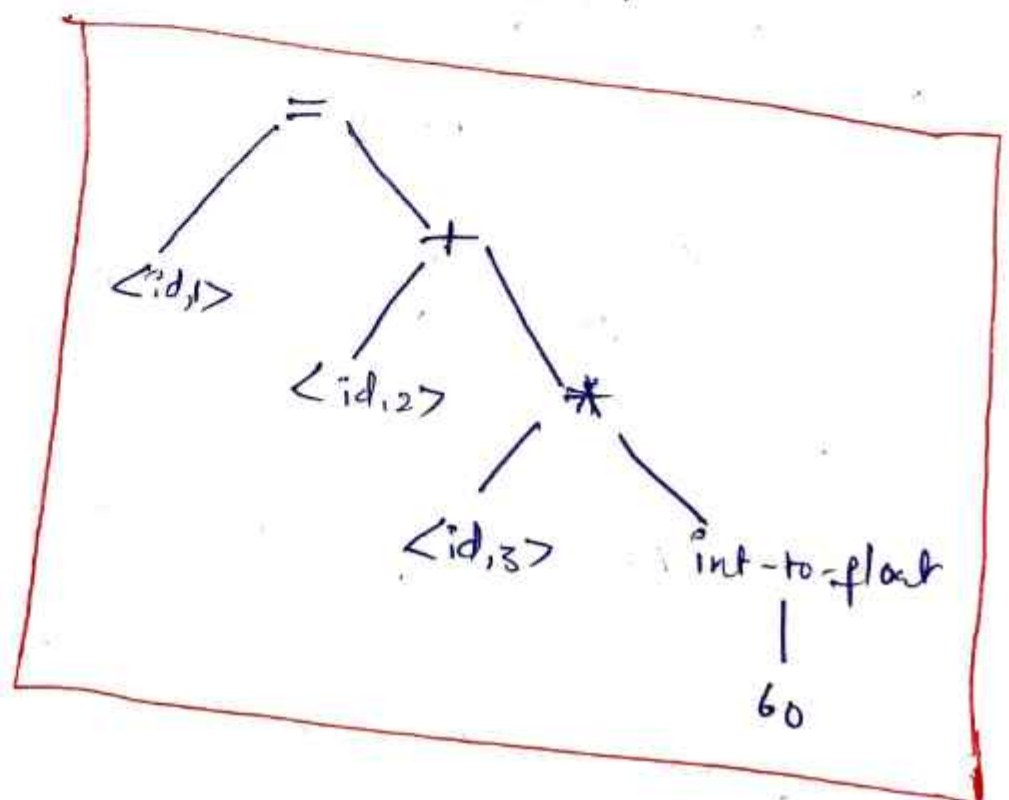
Semantic Analyzer uses the Syntax Tree and the info in symbol table to check the source program for semantic consistency with language definition.

An important part of semantic analysis is type checking, compiler checks each operator has matching operands.

The language Specification may permit type conversions known as Coersion.

For ex, binary Arithmetic Operator may be applied to either on two integers, or on Two floating point numbers. If Operator is applied to floating point number and integer, The compiler will convert integer into a floating point number.

Such Coersion appears in following Syntax tree, We consider that position, initial and rate have been declared to floating point numbers, and 60 is an integer. The Type checker in Semantic analysis Converts 60 to float.



Intermediate Code Generation.

While converting source program into machine code, a compiler may convert to one or more intermediate forms. Syntax Trees are a form of intermediate representations which are used in Syntax and semantic phases.

After Syntax and semantic analysis, compiler will generate an explicit low-level (or) machine like ^{intermediate} representation. It will have two important properties:

- i. It should be easy to produce
- ii. It should be easy to Translate into the target machine code.

One format of intermediate code used is Three address code, which consists of Assembly like instructions with Three operands per instruction. Each operand act like a register.

$t_1 = \text{inttofloat}(60).$

$t_2 = id_3 * t_1$

$t_3 = id_2 + t_2$

$id_1 * t_3$

In Three address code.

1. At most One Operator should be present at right side of instⁿ.
2. To store the result temporary names are used at left side.
3. Some instructions may have fewer than 3 Operands.

Code Optimization.

Machine independent code optimization phase attempts to improve the intermediate code so that better target code will result. Better means faster, shorter, and which consumes less power.

In above intermediate code representation The Optimizer can deduce the conversion of 60 from int to float, by replacing 60 by 60.0 t_3 is used only once, to transmit its value to id1, So optimizer can transform previous code as below.

$$t_1 = id_3 * 60.0$$

$$id_1 = id_2 + t_1$$

Code generation.

→ Code generator takes intermediate representation of source program and maps it into the target language.

→ If target language code is machine code & code registers or memory locations are selected for each variable used in the program.

→ An important task of code generation is assignment of registers to hold variables.

→ For ex. previous optimized intermediate code, using registers R1 and R2, can be translated to machine code as follows.

LDF. R2, id3.

MULF R2, R2, #60.0

LDF R1, id2

ADD F. R1, R1, R2

STF. id1, R1.

→ The first operand in each instⁿ specifies a destination. F in each instⁿ tells us that

it deals with floating point numbers.

Symbol Table Management

- Compiler will record the variable names used in source program, and various infoⁿ about attributes of each name.
- The Attributes may provide infoⁿ about storage allocated for a name, type, scope, and in case of functions, the number of arguments, type of arguments, parameter passing methods and the type returned.
- Symbol table is a data structure which contains a record for each variable, with fields for the attributes of that name.
- Symbol table is designed in such a way that compiler will find the record for each name quickly and stores and retrieves data from that name quickly.

The Science of Building a Compiler.

→ Compiler design is full of beautiful examples, where complicated real world problems are solved easily, mathematically. By using abstraction

→ A compiler must accept all source programs that conform to specification of programming languages. The programs can be large, possibly millions of lines of code.

→ While converting such big programs to machine code the compiler must preserve meaning of the program.

Modelling in Compiler Design and Implementation.

Compiler design is a study of how to design a mathematical model and choose the right algorithm.

Some of most fundamental models are finite-state machines and regular expressions which are used to describe syntactic lexical units of programs (keywords, identifiers etc)., And Context-Free Grammars used to describe syntactic structure of programming languages.

The Science of Code Optimization.

- Optimization in compiler design, attempts to produce code that is most efficient than obvious code.
- Optimization is complex because processor architecture ~~became~~ is very complex.
- Optimization is important because, parallel computers ~~needs~~ substantial optimization. Otherwise their performance suffers, by order of magnitude.

Programming Language Basics.

(10)

The Static / Dynamic Distinction.

→ If a programming language uses a policy that allows a compiler to decide an issue, then we say language uses static policy, or that issue can be decided at compile time.

→ A policy that allows a decision to be made when we execute a program is said to be a dynamic policy or to require a decision at runtime.

scope → scope of a variable x , is the region in which x is declared. A language uses static scope or lexical scope if we can say scope at time of declaration. In Dynamic scope, same use of x could refer to different declaration of x .

Ex: Consider the use of term "static" in java, to declare a variable, as follows:

```
public static int x;
```

Variable is a name of memory location used to hold data value.

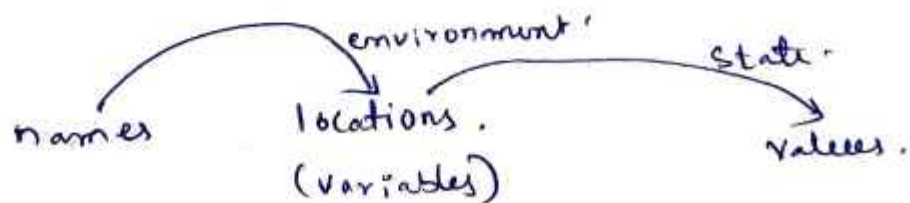
Above statement declares x as a static, i.e., only one variable is created, where all objects access same memory location while accessing this variable.

If static is omitted, all objects refers to different memory locations.

Environments and State

The Association of names with location in memory, and then values can be described by two mapping which changes as the program runs.

1. Environment → Mapping from names to memory location is known as Environment. Since variable refers to location, ^(l-value) Environment is mapping b/w name and Variable.
2. State → It is mapping from memory location to their values. State maps l-value to corresponding r-values in C terminology.



- Most bindings of names to locations is dynamic.
- The bindings of locations to values is generally dynamic as well as static.

for ex:-

#define ARRAYSIZE 1000.

binds the name ~~arr~~ ARRAYSIZE to a value statically.

Static scope and Block Structure.

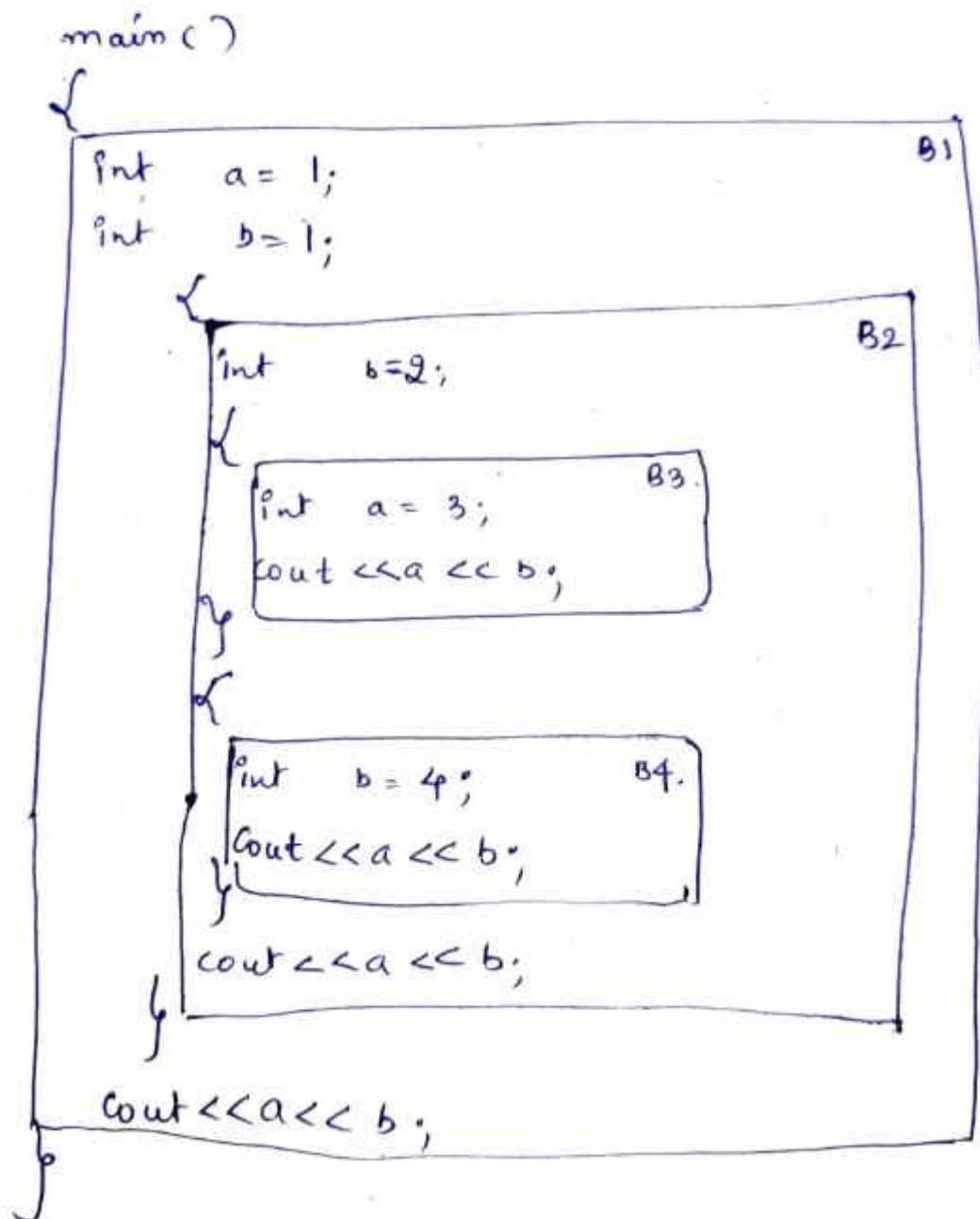
The Scope rules for C are based on program structure. Scope of a variable is implicitly determined by where declaration appears in the program. In languages like C++, C# and Java provides explicit control over scope through keywords like public, private and protected.

Static scope can be given by using blocks. Block is a grouping of declarations and statements. C uses braces "{" and "}" to denote a block. In some languages like Algol we use "begin" or "end" to denote a block.

Blocks can be nested inside other blocks. We say a variable declaration D, belongs to a block B, if it is declared in B. If the declaration D of a name x is in B, then scope of x is B, as well as in any nested block. x will not have scope in

nested block if it is redeclared in that block.

Ex:- Consider the program shown in below figure, which has four blocks. With several definitions of A and B. Each declaration initializes its variable to number of block to which it belongs.



(12)

For ex, Consider The declaration of 'int a=1' in B_1 , its scope is entire block B_1 , except the nested block of B_1 , where 'a' is redeclared. B_2 is immediately nested under B_1 , but does not contain redeclaration of ~~a~~, a, B_3 , which is nested under B_1 , has redeclaration of ~~a~~, a, so it is out of scope of declaration of a in B_1 , B_4 does not contain redeclaration of a, so it is in scope of B_1 's declaration.

Consider print statement of B_4 , since B_4 has redeclared $b=4$, so it prints a b value as 4, and it does not contain a, its surrounding block B_2 also not contains a's declaration, B_2 's surrounding block B_1 , contains a's declaration, so a's value is printed as 1.

<u>Declaration</u>	<u>Scope.</u>
int a=1	$B_1 - B_3$
int b=1	$B_1 - B_2$
int b=2	$B_2 - B_4$.
int a=3	B_3 .
int b=4	B_4 .

Explicit Access Control.

If p is an object of class with a member x , then $p.x$ refers to access of object p to class member x . And If a class C with member x , its subclass C' can also access x , except if it redeclared x .

The Object Oriented languages like C++, Java provides explicit access control through keywords public, private, protected. These keywords provides encapsulation by restricting access.

→ private limits the access to the members declared within the class and "friend" classes.

→ protected members can be accessed from subclasses.

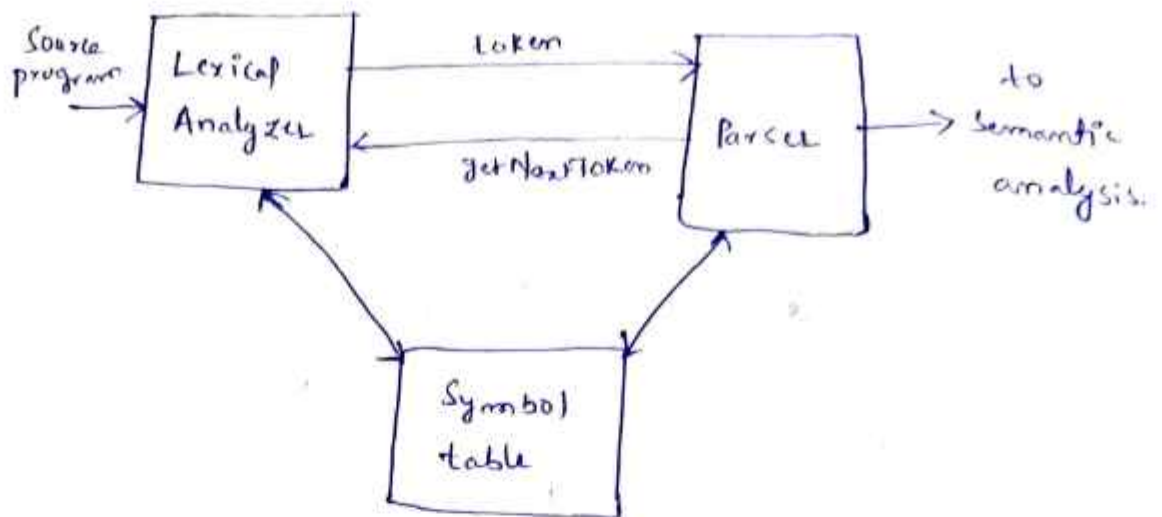
→ public members can be accessed from outside the class.

Chapter-2.

Lexical Analysis.Role of the Lexical Analyzer.

As it is the first phase compiler its main task is to read input characters of source program, group them into tokens of lexemes, and produce an output sequence of tokens from each lexeme. The stream of tokens is sent to lexical analyzer for syntax analysis. Lexical analyzer also interacts with symbol table. When lexical analyzer discovers a lexeme as an identifier, it enters the lexeme into symbol table. These interactions are shown in following figure.

Commonly the interaction is started by parser. The call from parser "get Next Token" causes the lexical analyzer to read characters from input until it identifies next lexeme, and produces token from it, which it returns to parser.



As the Lexical analyzer is part of compiler it will also do some other tasks besides identification of lexemes.

1. It will remove comments and white spaces (blank, newLine, tab).
2. It will correlate error messages generated by compiler with source program. It will associate a line number in source program with error message.
3. If the source program uses macros, the expansion of macros will be performed by Lexical Analyzer.

Lexical Analyzer consists of two processes.

- a). Scanning → Simply reads input, does not contend with tokenization.
- b). Lexical Analysis → It is complex, generates tokens from output of scanner.

Lexical Analysis Versus Parsing.

The reasons for separating Lexical Analysis from parsing are as follows.

1. The design of Compiler will be simplified.
For ex., if comments and white spaces are not removed in Lexical Analyzer, the parser's task will become more complex.
2. Compiler's Efficiency is improved.
3. Compiler's portability is enhanced.

Tokens, Patterns, Lexemes.

→ Token is a pair consists of token name and optional attribute value. Token name represents kind of lexical unit. Token name works as input to parser.

→ pattern denotes type / form of lexeme of a token. For ex., if token is a keyword, the pattern will be character that forms key word. For identify the pattern is more complex.

→ lexeme is a sequence of character in source program that matches pattern for a token.

For ex^y, in below statement

```
printf("Total = %.d\n", score);
```

printf, score matches the pattern for token id, and "Total = %.d\n", is a lexeme matching literal.

Token	Informal Description	Sample lexeme.
if	characters i, f	if.
else	characters e, l, s, e	if else
Comparison.	< or > or <= or >= or == or !=	<=, !=
id	letter followed by letters or digits	pi, score, d2.
literal	Anything that ^{that} surrounded by " ' s.	"hello".

Examples of tokens.

Attributes for tokens.

When more than one token lexeme matches a pattern, we need additional infoⁿ about ϕ matched lexeme and pattern.

(15)

→ For ex^o for pattern of token "number", both 0 and 1 matches. So it is important to know for code generator to know which lexeme is found in source program. Thus in many cases lexical analyzer sends token name and attribute value pair to parser that describes lexeme represented by token. The most important example is the token id. When we need to associate several info with identifier like its lexeme, token, type, and where it can found in symbol table. Thus appropriate value for identifier is a pointer to symbol table entry for that identifier.

For ex^o, The id token names and associated attribute values for the Fortran statement,

$$E = M * C ** 2.$$

are given as follows.

- <id, pointer to symbol table entry for E>
- <assign-op>
- <id, pointer to symbol table entry for M>
- <mult-op>
- <id, pointer to symbol table entry for C>
- <exp-op>
- <number, integer value 2>.

Input Buffering.

(16)

The task of reading lexeme from source program can be speed up by using Buffers.

It is a difficult task because we need to see next character, before we can sure that we have right lexeme. For ex, we can not be sure that we've have seen end of identifica until we see a character that is not a letter or digit, and any thing which is not a part of lexeme id.

Buffer pairs.

Buffers we used to reduce the overhead involved in reading input characters. This scheme consists of two buffers.

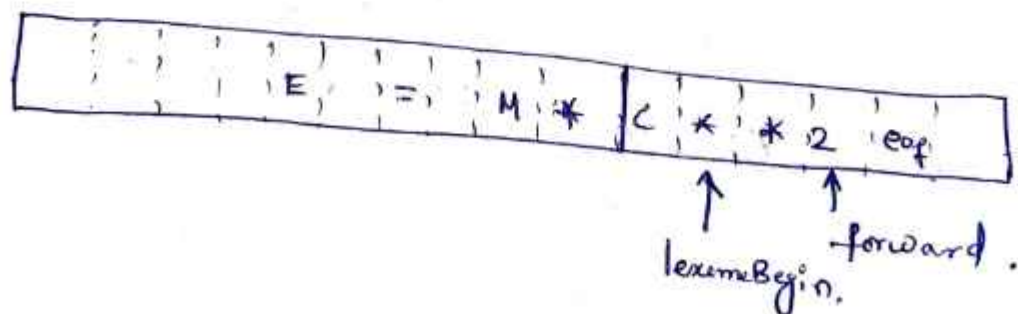
Each buffer size is same, say N , i.e., size of disk block e.g. 4096 bytes. Using a read command we can read N characters into buffer at a time, rather than using one system call per character. If fewer than N characters remain in i/p file, special character "eof" marks the end of source file, which is different from characters in source program.

Two pointers to input are maintained

1. pointer "LexemeBegin" marks the beginning of current Lexeme, whose length we are attempting to determine.
2. Pointer "forward" scans until the pattern ~~token~~ is matched.

Once the lexeme is matched, forward is set to character at right end. Then after the lexeme is found, and passed to parser as token, "lexeme begin" is set to next character found immediately after the lexeme just found.

When "forward" pointer has reached end of the buffer, we must reload other buffer from the input, and move forward to beginning of newly loaded buffer.



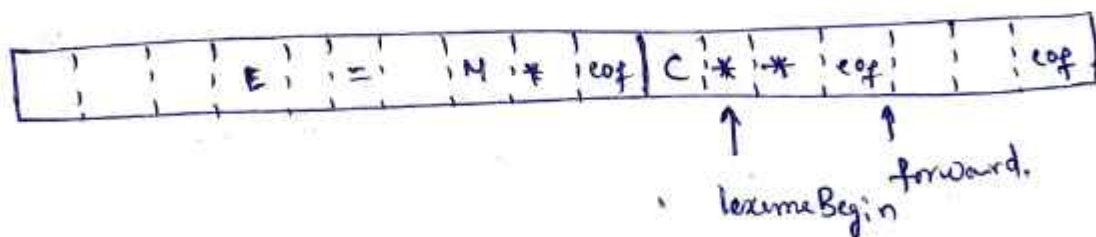
Sentinels.

For each character we read, we make two tests, one for end of the Buffer, and second one to determine, which character we read.

We can combine the buffer end test, and with test for new character read by using a special character known as Sentinel character at the end. Sentinel is a special character which is not part of source program, and this is represented by eof.

Eof is placed at the end of the buffer, as well as at the end of entire i/p.

As shown in below figure.



Recognition of Tokens.

The input string is compared with patterns, and then tokens are generated. For ex⁴. Consider below grammar

stmt $\rightarrow$ if expr then stmt.

| if expr then stmt else stmt.

| ϵ .

expr $\rightarrow$ term relop term.

| term.

term $\rightarrow$ id

| number.

The above grammar is a simple fragment of branching and conditional statements. This syntax is similar to pascal.

$\Rightarrow$ relop is a comparison operator like "=" for equals, and "<>" for not equals in pascal.

Terminals of the grammar is which are the names of tokens. if, then, else, relop, id and number are

The patterns for these tokens are given by regular definitions as follows:

digit $\rightarrow [0-9]$

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 $\rightarrow \text{else}$

rel op $\rightarrow <|>|<=|>=|=|<>$

The keywords if, then, else are the lexemes. They are not identifiers, but their lexemes matches patterns of identifiers.

Another job of lexical analyzer is to remove white spaces, by using token "ws" which is defined as follows.

ws $\rightarrow (\text{blank}/\text{tab}/\text{newline})^+$

When any of these are recognized as "ws", we do not return it to parser, rather the next character after whitespace is returned, i.e; Next token will be forwarded to parser.

Transition Diagrams

An intermediate step in lexical analyzer is we first convert patterns into flowcharts called as "Transition Diagram".

Transition diagram have a collection of nodes or circles called as states. Each state representations a condition that could occur during the process of scanning the input for matching the lexeme to pattern. State is a summary between lexemeBegin pointer and forward pointer.

Edges directs from one state of Transition diagram to other. Every edge is labeled by one or more symbols.

If we are in state 's', and next input symbol is 'a', we look for an edge out of s labeled by a. If such an edge is found, we advance forward pointer, and enter into the state where edge leads. If every state has only one edge the transition diagram is deterministic otherwise non deterministic.

Some important conventions about Transition diagram are.

1. Some states are called as Accepting or final. These states indicate that lexeme has been found. The accepting/final state is represented by double circle. If any action to be taken after this state like returning a token and attribute value to the parser, will be attached to accepting state.
2. If it is necessary to retract / move backward forward pointer to one position, we should place a * near accepting state.
3. One state is designated as "start" state, "initial state", which is indicated by a label start entering from nowhere.

Ex:- Following diagram shows a transition that recognizes lexemes matching

the token "relop" (relational operators).

→ We begin at state 0.

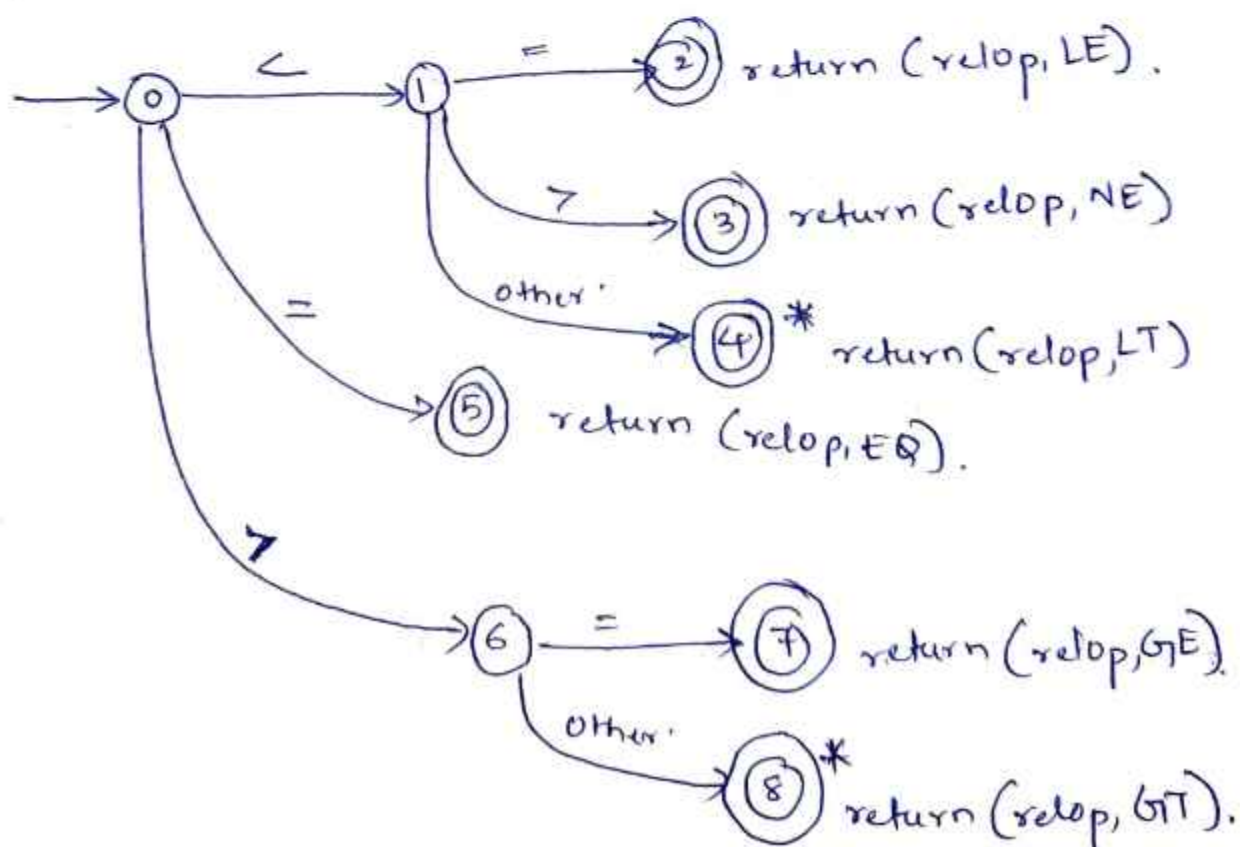
→ If input symbol is <, then among the lexemes that matches the pattern for relop we look for <, <=, <=.

→ we therefore goto state 1, and looks for next character.

→ If next character is =, we recognize <= and enter into state 2. and return the token relop with attribute LE.

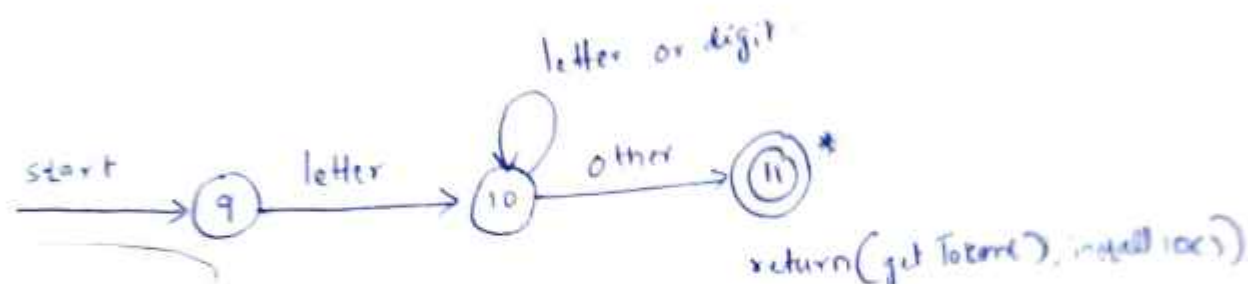
→ If in state 1, the next character is >, Then we have <>, and enter into state 3, and returns not-equals operators token.

→ On Any other character, The lexeme is <, and we enter into state 4. and returns that info. State 4 has * to indicate that we must retract the input one position.



Recognition of Reserved Words and Identifiers

Recognizing keywords and identifiers is somewhat different from recognizing operators. Keywords like `if` and `then` are reserved. They are not identifiers. Following Transition diagram identifies keywords `if`, `then`, etc.



There are Two Ways, to recognize reserved words, that looks like identifiers.

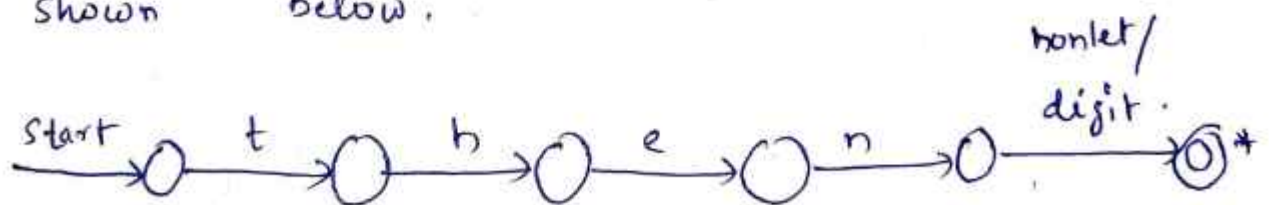
Initially, the reserved words in symbol table. A field of the symbol table tells that they are not identifiers. Another field tells which token they represent.

When we find an identifier, installID(), places the identifier in Symbol table. if it is not there already, and returns a pointer to the Symbol table entry found.

→ Any identifier not in the Symbol table during lexical Analyzer can not be a Reserved Word., so its token is id.

→ getToken(), examines the Symbol table entry for lexeme found, and returns the token name either id or one of the keyword tokens that initially installed in the table.

2. Create a separate transition diagram for each keyword, an ex for keyword "then" is shown below.



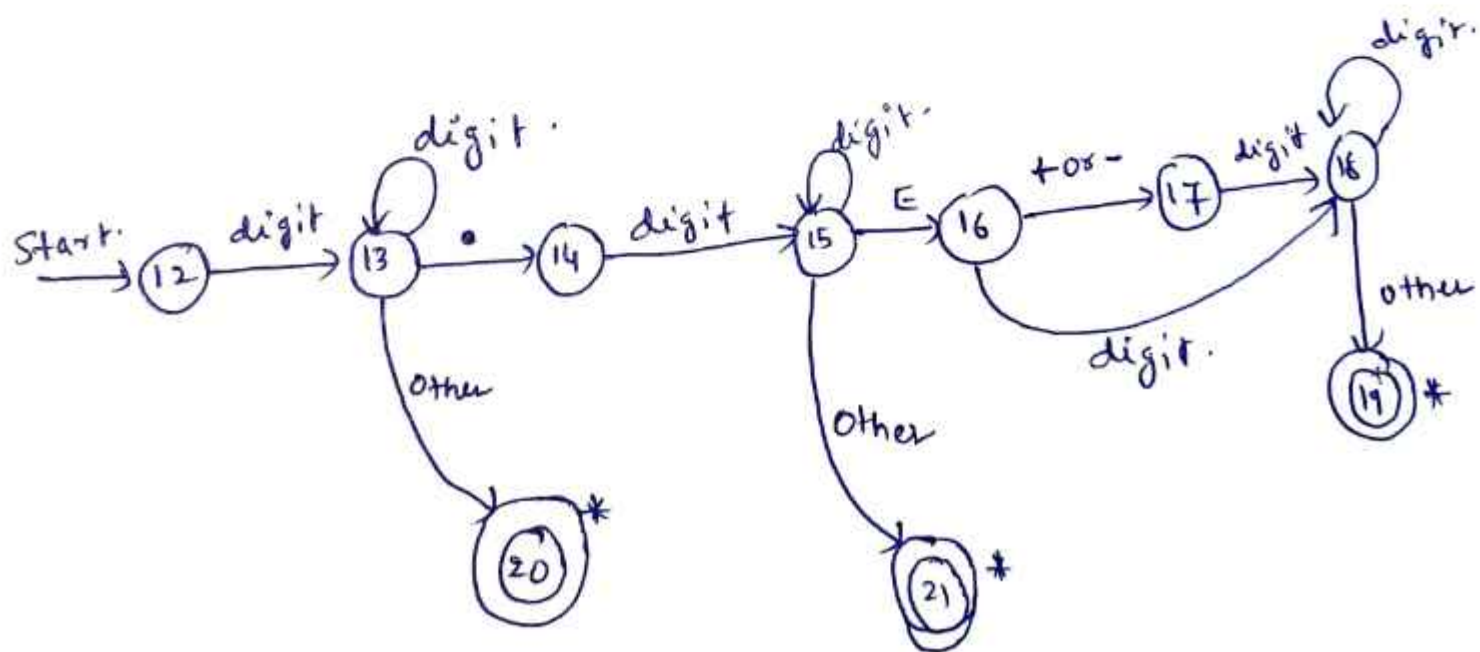
→ Transition diagram consists of states representing successive letters of keywords seen. followed by test for "non-letter-or-digit"

(21)

i.e; Any character that can not be a continuation of identifier. It is necessary to check the end of identifier, or else we return "then", where the correct token was "id" for lexeme "thenentvalue". If this approach is used, we should prioritize tokens so that reserved word tokens are recognized in preference to id, when lexeme matches both patterns.

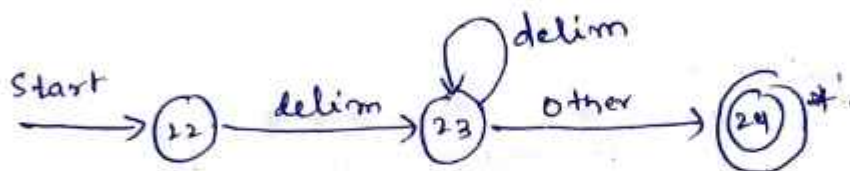
→ This approach is not used in our example, so the states in Transition diagram are numbered.

Transition Diagram for Unsigned Numbers.



A Transition Diagram for White Space.

In following diagram we look for one or more white spaces, represented by delim. These delim characters could be blank, tab, newline and characters which are not considered to be part of any token.



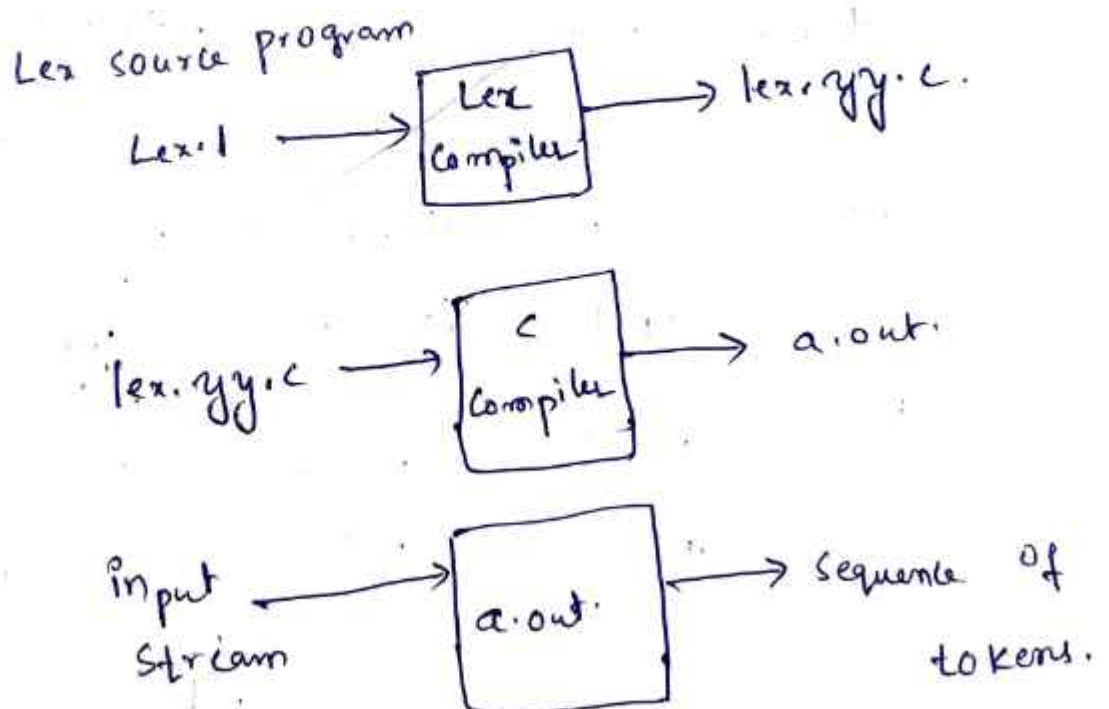
The Lexical Analyzer Generator Lex.

Lex or Flex, allows to specify a lexical analyzer, which describes patterns for Tokens.

The input notation for Lex tool is known as the Lex language, and the tool is known as Lex Compiler. Lex Compiler transforms input patterns into transition diagram and generates code in file called as lex.yy.c.

Use of Lex.

Following figure shows how Lex is used.



- An input file lex.l which is written in Lex language describes lexical Analyzer to be generated.
- Lex compiler transforms Lex.l to a C program, which is named as lex.yy.c
- This C file is compiled by C compiler, to a file a.out.
- This a.out file takes the input stream of source program and generates sequence of tokens.
- a.out is a function in C, that returns an integer, which is code for token names. The Attribute value, i.e; either a numeric code, or pointer to symbol table entry is placed in a global variable "yyval." which is shared between Lexical Analyzer and parser. So a.out returns ~~attribute~~ Token name and attribute Value.

Structure of Lex Programs.

Lex program has following pattern

declaration

%.

Transition rules.

%.

Auxiliary Functions.

→ Declaration section consists of declaration of variables and constants.

→ The Transition rules will have the form

pattern of action {

→ pattern uses definitions of declarations in declaration section. The actions fragments of code written in C.

→ The Third section holds additional functions used in actions.

Lexical Analyzer created by Lex, when called by parser, reads input string one character at a time, until longest prefix matches any one of the pattern P_i , it then ~~executes~~ executes Action A_i . A_i returns a single Value Token Name to parser, and a shared Variable $yyval$ passes additional info about lexeme found, if needed.

If P_i describes white spaces, lexical Analyzer proceeds to find additional lexemes which matches to a pattern, and correspondingly Action is returned to parser.

Example. Following is a lex program That recognizes tokens of below figure and returns Tokens Found.

Lexeme	TokenName	Attribute Value.
Any ws	—	—
if	if	—
Then	Then	—
else	else	—
Any id	id	pointer to table entry
Any Number	number	pointer to table entry.

<	relop	LT
<=	relop	LE
=	relop	EQ
<>	relop	NE
>	relop	GT
>=	relop	GE

Finite Automata.

Lex is a tool, which converts input programs to lexical Analyzer. The heart of transition is known as Finite Automata.

Finite automata are graphs, like Transition diagrams, with following differences.

1. Finite Automata are recognizers, They say "yes" or "no" about each possible string.
2. Finite Automata is of two types.

i. Non Deterministic Finite Automata (NFA).

A symbol can label many edges out of same state, and ϵ is an empty string, is a possible label.

b). ii. Deterministic Finite Automata (DFA)

Each state for each symbol as input, has exactly one edge with symbol leaving that state.

Both DFA, and NFA can recognize some languages known as regular languages.

Non Deterministic Finite Automata (NFA)

NFA consists of.

1. A finite set of states S .
2. A set of input symbols Σ , known as input alphabet. ϵ is an empty input string, which is never a member of Σ .
3. A Transition function δ ; which gives for each state and symbol in Σ a set of next states.
4. A start state s_0 , from S . (or initial state).
5. A set of states F , known as accepting states (or final states).

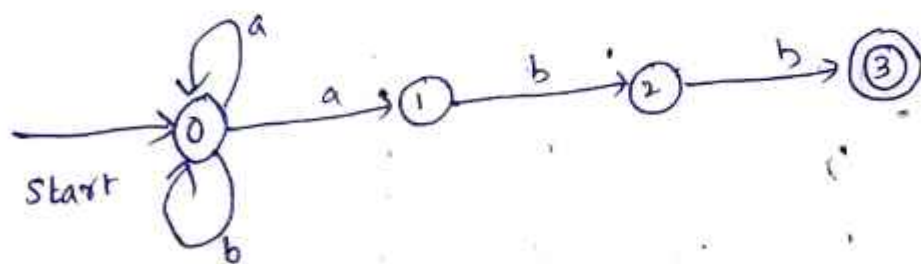
We can represent NFA or DFA by using transition graph, where nodes are states, labeled edges represent transition function.

This graph is same like Transition

diagram but with following differences. (26)

- The same symbol, can label edges from one state to several different states.
- An edge may represent by ϵ , the empty string, in addition to symbols of input string.

Example :- An NFA, recognizing the language of regular expression $(a/b)^*abb$ can be given as follows. This example describes all strings of a's and b's ends with abb



This NFA, accepts all strings ends with abb.

Transition Tables.

NFA can also be represented by Transition table, whose rows correspond to states, and columns correspond to input symbol and ϵ . The entry in the table for a given state and input is the target state. If the transition function has no info about state input pair, we put ϕ in the table.

The Transition table for previous NFA can be given as follows.

State	a	b	ϵ
0	$\{0, 1\}$	$\{0\}$	ϕ
1	ϕ	$\{2\}$	ϕ
2	ϕ	$\{3\}$	ϕ
3	ϕ	ϕ	ϕ

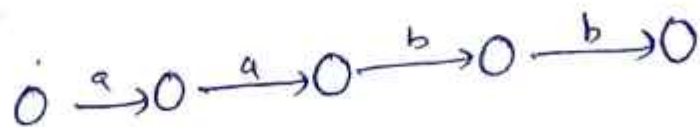
Accepting input Strings.

→ NFA accepts input string x , if only there is some path in the transition graph from the start state to one of the Accepting states. Labels with ϵ along the path are ignored.

Ex:- The string $aabb$, of previous NFA is accepted, The path is from state 0 to 3 demonstrating the fact as follows.

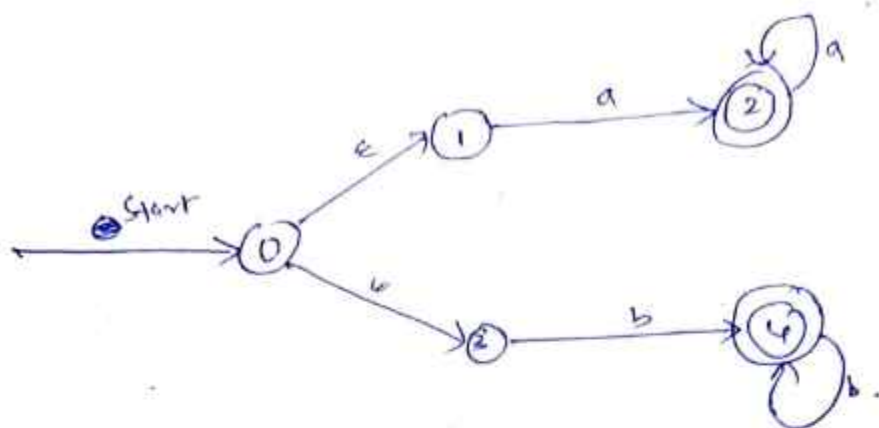


Several path for some string may lead to different states. for ex ||



This path leads to state 0, which is not final state, so string on this path is not accepted.

Ex:- 2) following NFA accepts $L(a^x | b^y)$.



The string "aaa" is accepted because of path
 $0 \xrightarrow{a} 1 \xrightarrow{a} 2 \xrightarrow{a} 2$

Deterministic Finite Automata (DFA)

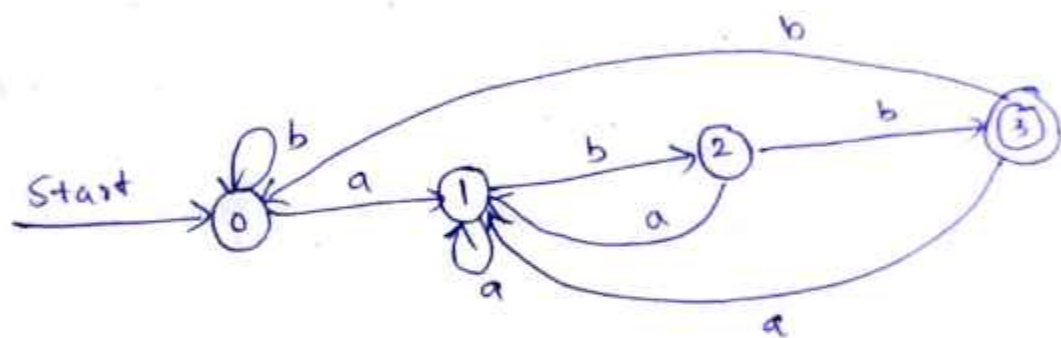
DFA is special case of NFA, where

1. There are no moves on input ϵ , and
2. For each state s , and input symbol a , There is exactly one edge labeled by 'a' out of s .

In Transition Table of DFA, each entry is single state, so states are represented by ~~set~~ without using curly braces, which are used to represent sets.

Every regular expression and NFA can be converted to DFA accepting same language. DFA is used to build or simulate lexical Analyzers.

Ex:- following figure represents DFA accepting language $(a|b)^*abb$, same as previous NFA.



Given string 'a babb' is accepted along the

path $0 \xrightarrow{a} 1 \xrightarrow{b} 2 \xrightarrow{a} 1 \xrightarrow{b} 2 \xrightarrow{b} 3$.

Regular Expression to Automata.

Regular Expressions is a choice notation to describe lexical Analyzers and other pattern processing softwares.

Implementation of ϵ requires simulation of a DFA. Because NFA has choice of input symbol or ϵ , NFA's simulation is less straightforward than DFA. Thus it is often important to convert NFA to DFA that accepts same language.

Conversion of an NFA to DFA.

→ After conversion each state of DFA, corresponds to a set of NFA states. This technique is known as subset construction. This can be given by using following Algorithm.

Algorithm: Subset Construction of a DFA from NFA.

Input : An NFA N .

Output : A DFA D , accepting same Language N .

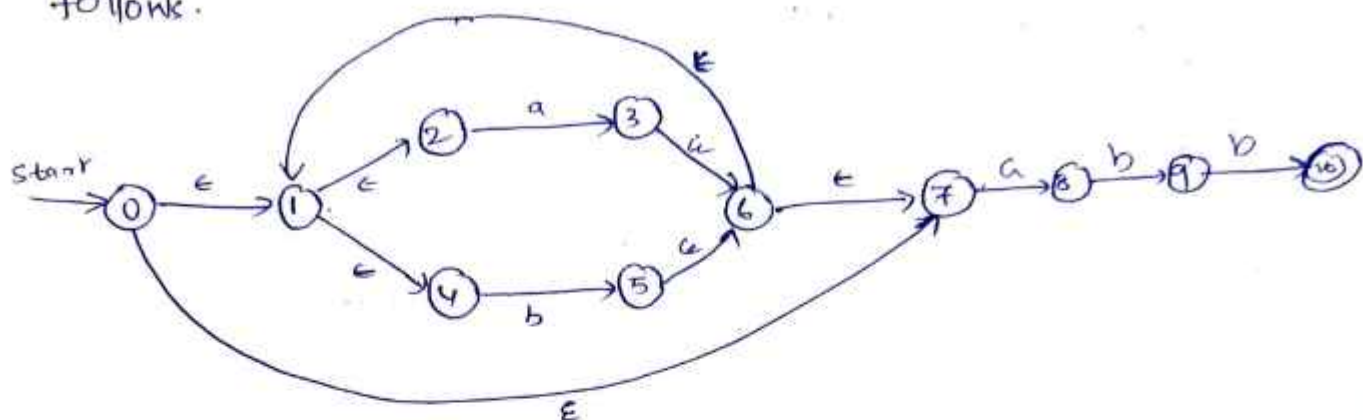
Method :- $\rightarrow$ This algorithm constructs a Transition table D_{tran} for DFA.

$\rightarrow$ Each state of DFA, also contains set of states of NFA, which has moves on a given input string. First we need to deal with null transitions of NFA.

Following figure shows different functions that describes computations on states of NFA. s denotes single state of NFA and T is set of states of NFA.

<u>Operation</u>	<u>Description</u>
$\epsilon\text{-closure}(s)$	$\rightarrow$ Set of NFA states reachable on ϵ -transitions alone.
$\epsilon\text{-closure}(T)$	$\rightarrow$ Set of NFA states, reachable from set of states in set T on ϵ -Transitions.
$\text{move}(T, a)$	$\rightarrow$ Set of states from T on a ,

→ Following figure shows NFA, accepting $(a/b)^*abb$. It can be converted to DFA as follows.



→ ϵ -closure(0) is $\{0, 1, 2, 4, 7\}$, These are the states which are reachable from 0, on ϵ -transition. State 0, can be reached from itself using ϵ -transition. These set of states are equivalent to a set A, in DFA.

→ The input alphabet is $\{a, b\}$, we calculate/find ϵ -closure(move(A, a)) i.e; $D_{trans}[A, a]$, and ϵ -closure(move(A, b)) i.e; $D_{trans}[A, b]$.

→ In state A, only 2 and 7, have Transition on a to 3 and 8 respectively. Thus $Move(A, a) = \{3, 8\}$, and ϵ -closure($\{3, 8\}$) = $\{1, 2, 3, 4, 6, 7, 8\}$

so,

$$D_{\text{trans}}[A, a] = \epsilon\text{-closure}(\text{move}(A, a)) = \epsilon\text{-closure}(\{3, 8\}) \\ = \{1, 2, 3, 4, 6, 7, 8\},$$

Call this state as B.

$$D_{\text{trans}}[A, a] = B.$$

$\Rightarrow$ Now $D_{\text{trans}}[A, b]$, among states of A, only 4 has transition on b to 5. So

$$D_{\text{trans}}[A, b] = \epsilon\text{-closure}(\{5\}) = \{1, 2, 4, 5, 6, 7\}.$$

Call this state as C.

$$D_{\text{trans}}[A, b] = C.$$

$\Rightarrow D_{\text{trans}}[B, a] = \epsilon\text{-closure}(\{3, 8\})$, which is similar to B itself.

$$\text{so } D_{\text{trans}}[B, a] = B$$

$\Rightarrow D_{\text{trans}}[B, b] = \epsilon\text{-closure}(\{5, 9\}) = \{1, 2, 4, 5, 6, 7, 9\}$

$$\text{let } D_{\text{trans}}[B, b] = D.$$

$$D_{\text{trans}}[C, a] = \epsilon\text{-closure}(\text{move}(C, a)) = \epsilon\text{-closure}(3, 8) \\ = \{1, 2, 3, 4, 6, 7, 8\} = B.$$

$$D_{\text{trans}}[C, a] = B.$$

$$D_{\text{trans}}[C, b] = \epsilon\text{-closure}(\text{move}(C, b)) = \epsilon\text{-closure}(5, 4) \\ = \{1, 2, 4, 5, 6, 7\} = C.$$

$$D_{\text{trans}}[D, a] = \epsilon\text{-closure}(\text{move}(D, a)) = \epsilon\text{-closure}(3, 8) \\ = \{1, 2, 3, 4, 6, 7, 8\} \\ = B.$$

$$D_{\text{trans}}[D, b] = \epsilon\text{-closure}(\text{move}(D, b)) = \epsilon\text{-closure}(5, 10) \\ = \{1, 2, 4, 5, 6, 7, 10\} = E.$$

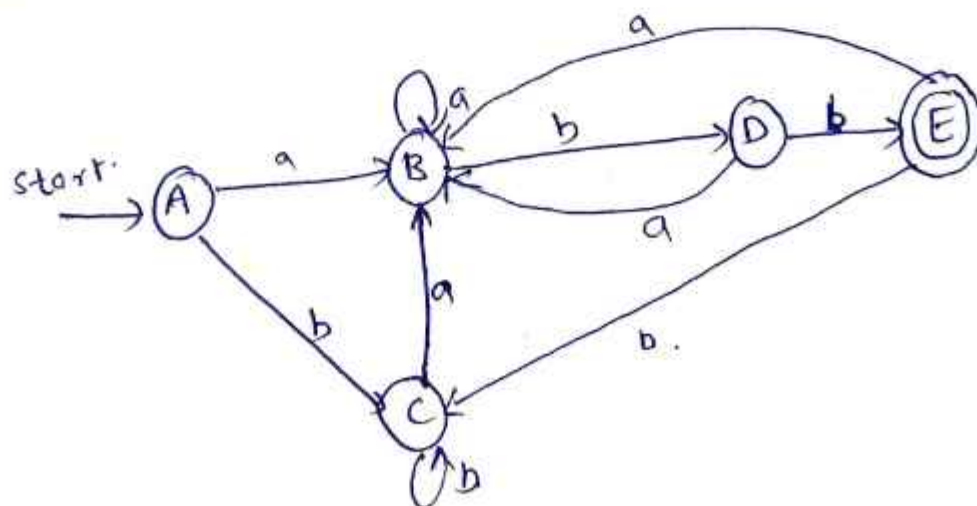
$$D_{\text{trans}}[E, a] = \epsilon\text{-closure}(\text{move}(E, a)) = \epsilon\text{-closure}(3, 8) \\ = \{1, 2, 3, 4, 6, 7, 8\} = B.$$

$$D_{\text{trans}}[E, b] = \epsilon\text{-closure}(\text{move}(E, b)) = \epsilon\text{-closure}(5) \\ = \{1, 2, 4, 5, 6, 7\} = C.$$

State A is the start state, and State E, which contains state 10 of NFA, is only Accepting state. The final D Trans can be given as.

NFA State	DFA state	a	b
$\{0, 1, 2, 4, 7\}$	A	B	C
$\{1, 2, 3, 4, 6, 7, 8\}$	B	B	D
$\{1, 2, 4, 5, 6, 7\}$	C	B	C
$\{1, 2, 4, 5, 6, 7, 9\}$	D	B	E
$\{1, 2, 4, 5, 6, 7, 10\}$	E	B	C

The DFA for above NFA, using subset construction is given as follows.



Design of a Lexical-Analyzer Generator.

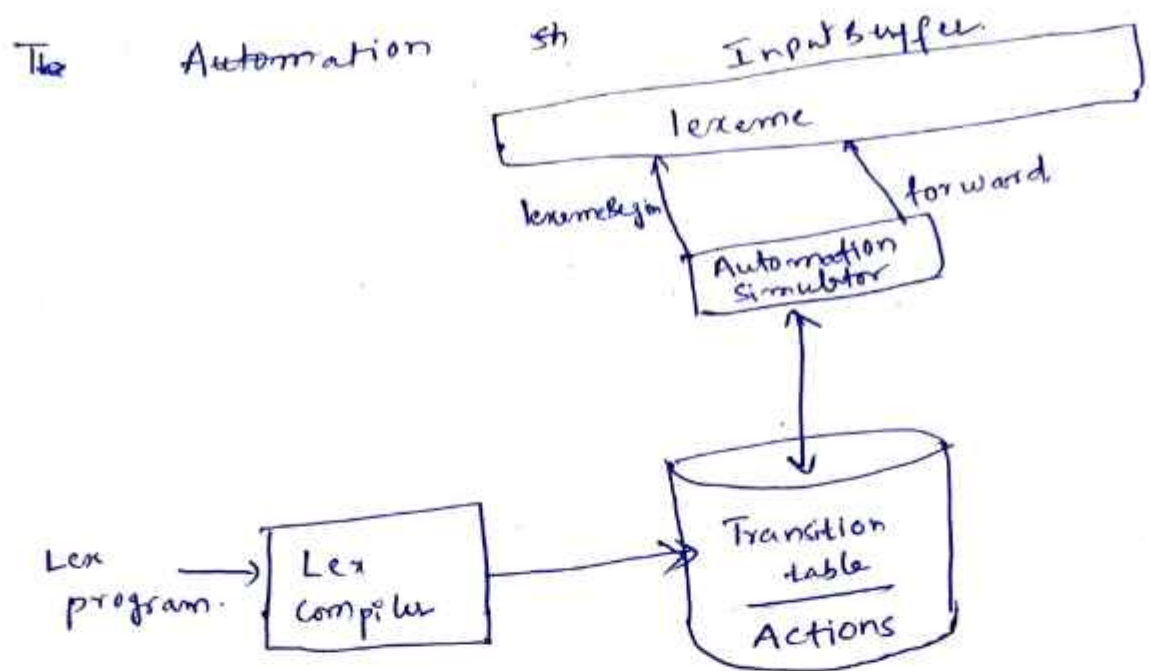
Following figure shows the architecture of lexical analyzer generated by Lex. The program uses either finite automation which can be either deterministic or non deterministic.

Following is a lex program which is turned into a Transition table and actions which will be used by Finite Automata Simulator.

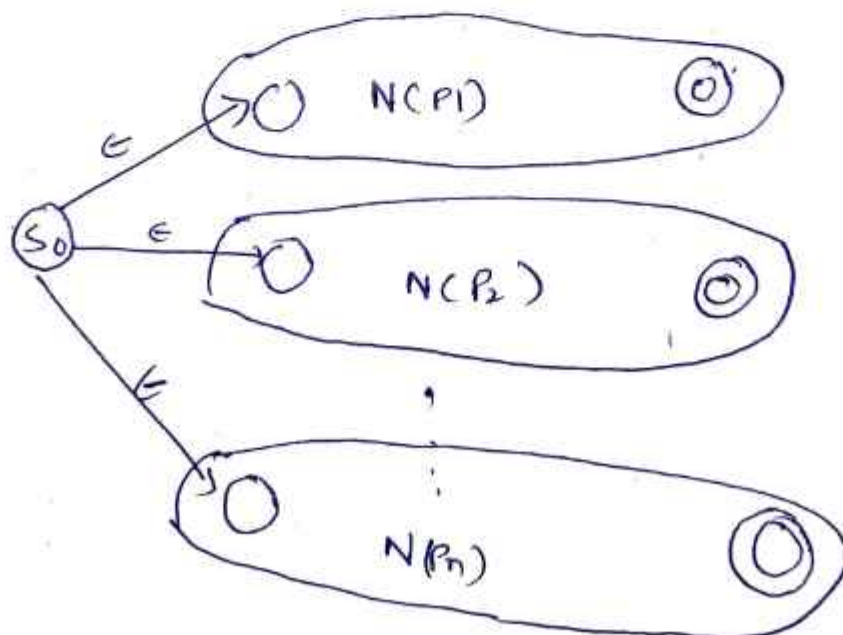
They components include

1. A Transition table for automation.
2. Functions passed directly through Lex to the Output.
3. The actions from the input program, which are to be invoked at the appropriate time by simulator.

To Construct an automation, we begin by taking each regular-expression pattern in Lex program and converting it to NFA.



The Automation will recognize all lexemes matching any of the patterns in source program with ϵ -transitions. to from each start states of the NFA's N_i of pattern P_i as shown in below.



Example.

Consider following patterns.

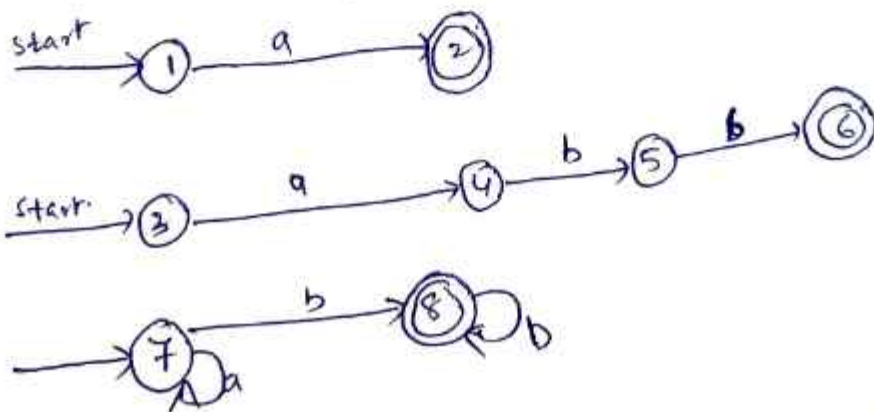
a action A_1 for pattern P_1

abb action A_2 for pattern P_2

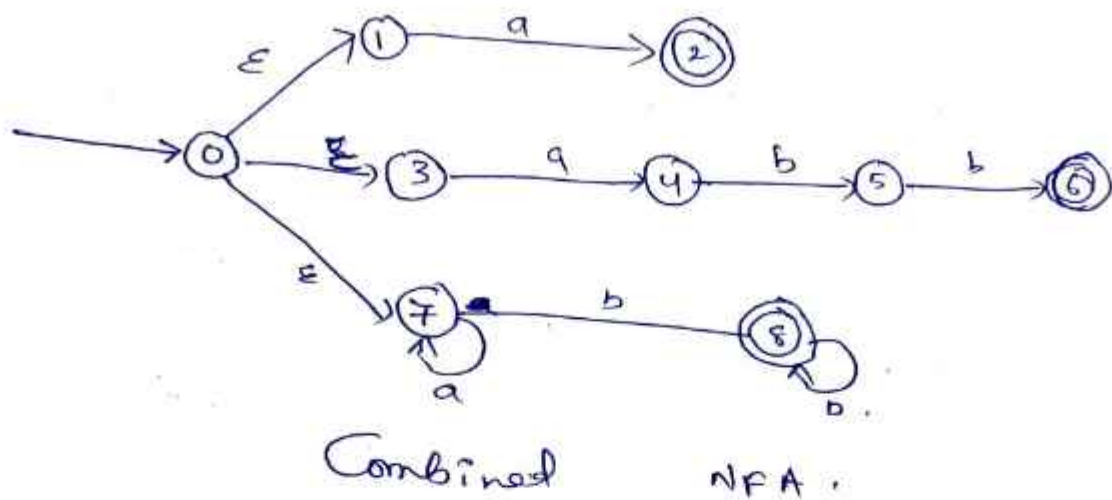
a^*b^+ action A_3 for pattern P_3

There are some conflicts in this example. The pattern abb match P_2 and P_3 , but we consider it as lexeme for P_2 , because it is listed first.

Following figure shows NFA's that recognize three patterns. The last figure shows combination of all above three NFA's.



NFA's for a, abb, and a^*b^+ .



Pattern Matching Based On NFA's.

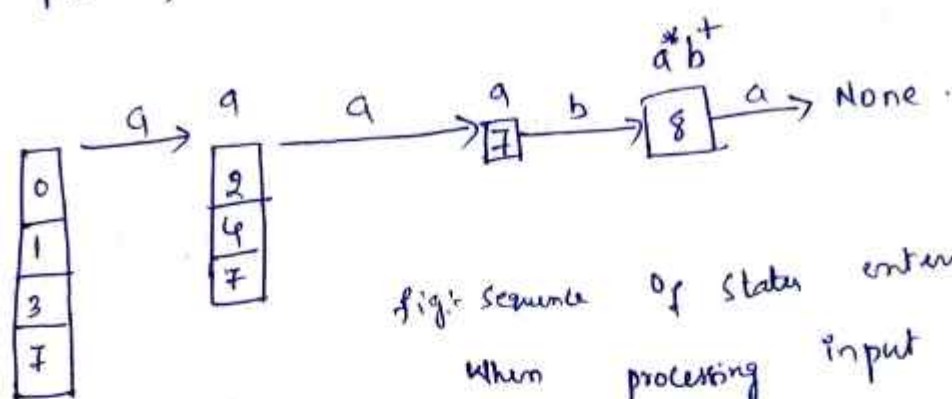
Lexical Analyzer which stimulates NFA, reads the input from the beginning referred to as "lexeme Begin", As it moves forward pointer, it calculates the set of states it is in each point.

When NFA reaches a point on input, where there are no next states, the next set of states will be empty. Then by seeing accepting states in transition we decide the longest prefix that is a lexeme matching the pattern.

We look backward in ^{sequence} set of states, until we find one or more accepting states. If there are more than one accepting state, we take the ^{state with} earliest p_i (pattern) in the list from Lexical Analyzer, move the forward pointer back to the end of lexeme, and perform the action A_i associated with pattern p_i .

Ex 11:- Consider patterns of given example, and input begins with "aaba". The below figure shows set of states of NFA of previous figure, we enter states $\{0, 1, 3, 7\}$ on enclosure ~~a~~ which is final initial state. After reading fourth i/p symbol, T a. We are in empty set of states, since

There are no transitions out of state 8 on input a,



Thus we backup from that position, looking for accepting states, we notice in above figure after reading first input symbol a, we are in state 2, which indicates pattern "a" is matched, but after reading aab, we are in state 8, which indicates a^*b^+ has been matched. Thus prefix aab, is the longest prefix up to an accepting state, we therefore select "aab" as lexeme, and executes Action A_3 , which returns the parser that token whose pattern $p_3 = a^*b^+$ has been found.

Optimization of DFA-Based Pattern Matchers.

There are Three algorithms that can be used to implement and Optimize pattern matchers from Regular Expressions.

1. First algorithm is used in LEX Compiler, it constructs DFA directly from a regular Expression, without constructing intermediate NFA.
2. Second algorithm minimizes the states of DFA, by combining the states that have the same future behavior.
3. The third algorithm produces more compact representations of transition tables than the standard, two dimensional table.

Important states of NFA

An NFA state is known as important if it has non- ϵ out transitions. The Subset Construction uses only the important states in a set T , when it computes ϵ -closure(move(T, a)). The set of states reachable from T on a .

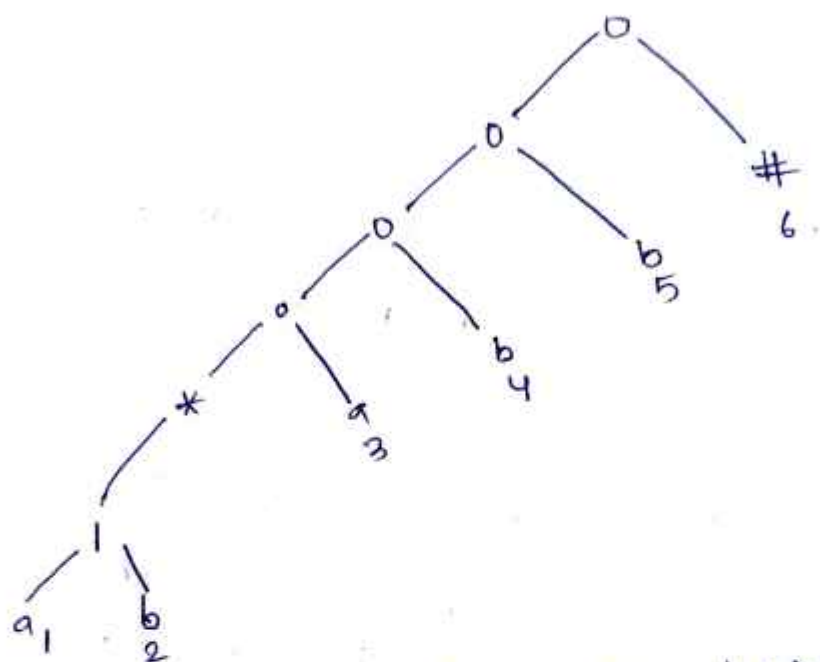
During Subset Construction, two sets of NFA States can be identified

1. Have the same important states; and.
2. Either both have accepting states or neither does.

Each important state corresponds to a particular Operand in regular Expression. The Constructed NFA, has only one Accepting state but does not have out-transitions. By concatenating right end marker $\#$ to a regular expression (r) we give accepting state for r on $\#$ transition of $\#$, making it important state of NFA for $(r)\#$.

The important states of NFA, corresponds to positions in the regular expression that holds symbols of alphabets. Regular Expression can be represented by Syntax Tree, where leaves denotes Operands, interior nodes denotes Operators. Interior nodes are cat-node (dot), or-node ($|$), or star-node ($*$)

Following figure shows Syntax Tree for regular expression $(a|b)^*abb\#$



Leaves in syntax tree are labeled by ϵ , or by alphabet symbol, if the leaf is not labeled by ϵ , we attach an integer, which refers to position of leaf / symbol in regular expression. One symbol can have many positions, for ex, a has position 1 and 3. positions corresponds to important states of NFA.

Functions Computed from the Syntax Tree.

To construct DFA, directly from regular expression, we construct Syntax Tree, and then compute four functions nullable, firstpos, lastpos. and

Followpos, which are defined as follows

1. nullable(n) $\Rightarrow$ This function is true, in a Syntax Tree for node n , if the subexpression represented by n has ϵ in its language.

The subexpression can be made null.

2. Firstpos(n) $\Rightarrow$ is set of positions in subtree rooted at n , that corresponds to first symbol.

3. last Symbol(n) $\Rightarrow$ set of positions in the subtree rooted at n that corresponds to the last symbol of at least one string.

4. Followpos $\Rightarrow$ for a position p , there is a set of positions q in entire Syntax Tree, such that there is a string

$x = a_1 a_2 \dots a_n$ in $L((r)\#)$, where a_i

is matched to position p , and a_{i+1}

to position q .

Example.. Consider the cat-node n , of previous Syntax Tree that corresponds to expression $(a|b)^*a$.

Nullable (n) $\Rightarrow$ Nullable is false, for this node as it generates all the strings of a's and b's ending with a. On the other hand the star node, below it can be nullable, it may generate ϵ along with strings of a's and b's.

firstpos (n) $\Rightarrow \{1, 2, 3\}$, In strings generated like 'aa', The first position corresponds to 1, in string 'ba', The first position corresponds to 2, And string with only 'a', first position corresponds to 3.

lastpos (n) $\Rightarrow \{3\}$, For all strings, generated from expression of node n , The last position is for 'a', and it is 3.

Computing nullable, firstpos, and lastpos.

We can compute nullable, firstpos and lastpos depending on the height of tree. The following table represents the inductive rules for nullable and firstpos. The rules for lastpos are same as for firstpos.

Node (n)	Nullable(n)	firstpos(n)
A leaf node labeled ϵ	true	ϕ
A leaf with position i	false	$\{i\}$.
An or-node $n = c_1 \mid c_2$	nullable(c_1) or nullable(c_2)	firstpos(c_1) $\cup$ firstpos(c_2)
A cat-code $n = c_1 c_2$	nullable(c_1) and nullable(c_2)	if (nullable(c_1)) firstpos(c_1) $\cup$ firstpos(c_2) else firstpos(c_1)
A Star Node $n = c_1^*$	true	firstpos(c_1).

(3)

Example :- From all the nodes of previous Syntax Tree only star-node is nullable, and none of the leaves are nullable because they corresponds to non- ϵ operands.

→ The 'or node' is not nullable, because neither of its children are null.

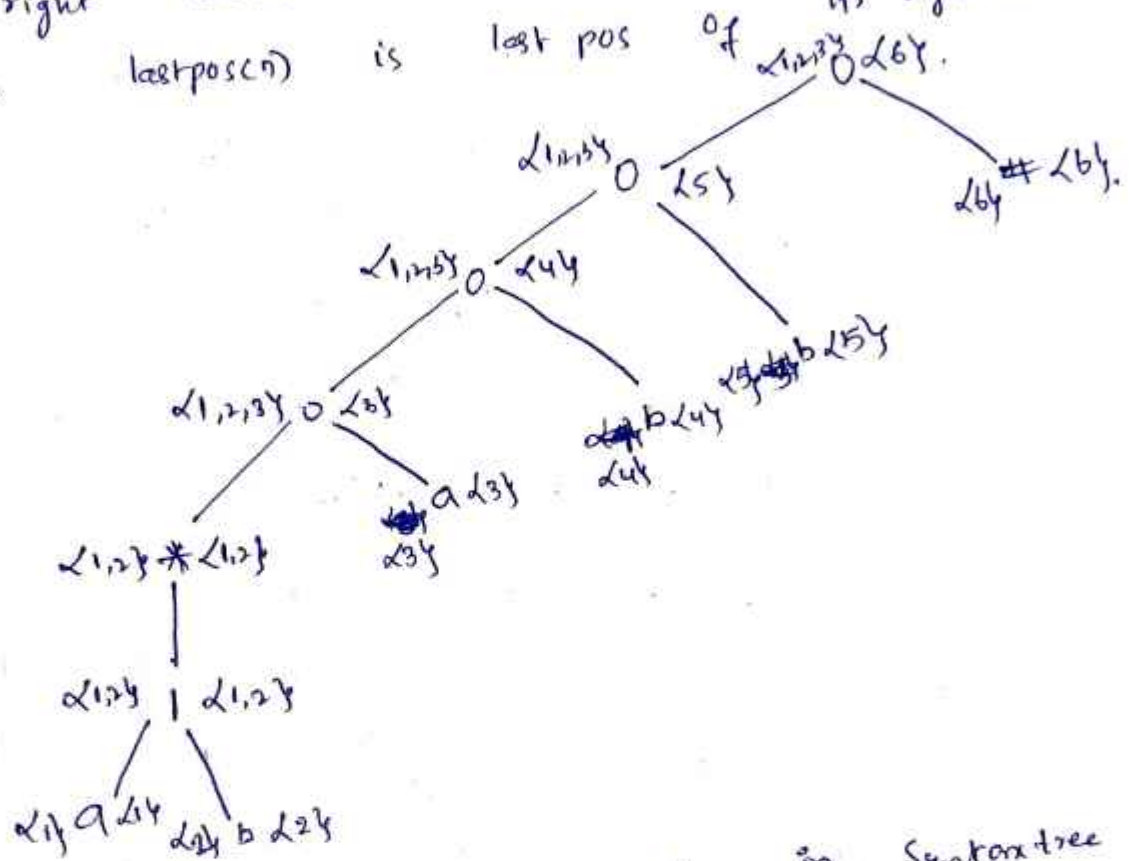
⇒ The star node is nullable because every star node is nullable.

⇒ The cat node is ^{not} nullable if at least one child is non-nullable.

The computation of firstpos and lastpos for each node is shown in below figure, with firstpos(n) to the left ^{of} node n, and lastpos(n) to its right. Each leaf has only itself for firstpos and lastpos, as the rule for non- ϵ leaves. For the or-node, we take the union of firstpos at the children and do the same for lastpos. The rule for the star-node says that we take value of firstpos or lastpos at any one child of that node.

For ex, consider lowest cat-node, consider it as n. Its left child is nullable, and so its first pos is $\{\epsilon\}$, and first pos of right child is $\{3\}$'s union will give i.e. $\{\epsilon, 3\}$ will give

n 's firstpos. The rule for lastpos is
 same like firstpos, with childs interchanged. That is
 to compute lastpos(n), we must check whether
 its right child is nullable, which is not.
 Therefore lastpos(n) is last pos of its right child.



First pos and last pos for nodes in syntax tree
 for $(a|b)^*abb\#$.

Followpos.

Followpos can be applied to only $\text{node}(\bullet)$, and closure/store $\text{node}(\ast)$. The rules for finding followpos are for a node n are.

if n is $\bullet$ then

$$\text{Firstpos}(c_2) \rightarrow \text{lastpos}(c_1)$$

i.e; copy firstpos c_2 to lastpos of c_1

if n is $\ast$ then

$$\text{Firstpos}(n) \rightarrow \text{lastpos}(n).$$

i.e; copy firstpos of n i.e; $\ast$ to lastpos of $\ast$

The final followpos according to above two rules is shown in below table. First we find followpos for $\bullet$ and then for $\ast$.

Position	n	followpos(n)
1		$\{1, 2, 3, 4\}$
2		$\{1, 2, 3, 4\}$
3		$\{4\}$
4		$\{5\}$
5		$\{6\}$
6		ϕ

Converting the Regular Expression to DFA.

We now put all the steps of our example to construct DFA, for regular expression $r = (a|b)^*abb$. We observed that nullable is true only for * node.

The value of firstpos for the root is $\{1, 2, 3\}$, so this is the start node of D, consider this state as A. i.e; $A = \{1, 2, 3\}$, Now we find $D_{trans}[A, a]$, and $D_{trans}[A, b]$. From the tree, among the positions of A, 1 and 3 corresponds to a, and 2 corresponds to b.

So

$$\begin{aligned} D_{trans}[A, a] &= \text{followpos}(1) \cup \text{followpos}(3) \\ &= \{1, 2, 3\} \cup \{4\} \\ &= \{1, 2, 3, 4\} \Rightarrow \text{let this state B.} \end{aligned}$$

$$\begin{aligned} D_{trans}[A, b] &= \text{followpos}(2) \\ &= \{1, 2, 3\} \Rightarrow \text{it is same as A.} \end{aligned}$$

Similarly.

$$\begin{aligned} D_{trans}[B, a] &= \text{followpos}(1) \cup \text{followpos}(3) \\ &= \{1, 2, 3, 4\} = B. \end{aligned}$$

$$\begin{aligned} D_{trans}[B, b] &= \text{followpos}(2) \cup \text{followpos}(4) \\ &= \{1, 2, 3\} \cup \{5\} \\ &= \{1, 2, 3, 5\} = C. \end{aligned}$$

$$D_{trans}[C, a] = followpos(1) \cup followpos(3)$$

$$= \{1, 2, 3, 4\} = B$$

$$D_{trans}[C, b] = followpos(2) \cup followpos(4) \cup followpos(5)$$

$$= \{1, 2, 3\} \cup \{5\} \cup \{6\}$$

$$= \{1, 2, 3, 5, 6\} = D$$

$\Rightarrow$ it is the final state as it has position of end marker.

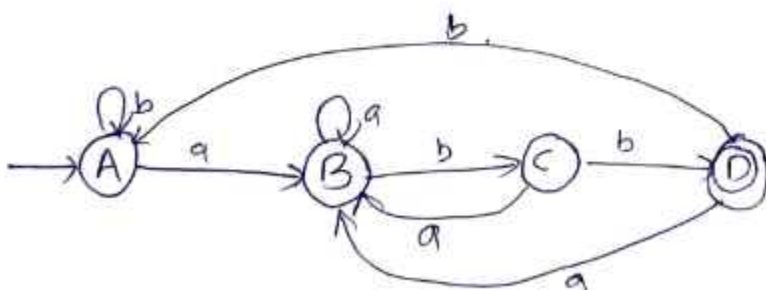
$$D_{trans}[D, a] = followpos(4) \cup followpos(3)$$

$$= \{1, 2, 3, 4\} = B$$

$$D_{trans}[D, b] = followpos(2)$$

$$= \{1, 2, 3\} = A$$

The DFA can be given as.

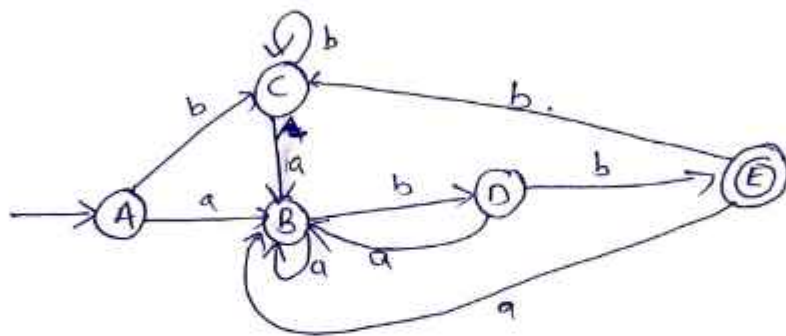


Minimizing The number of States of a DFA:

There can be many DFA's that recognize same language. If we implement a Lexical Analyzer using DFA, we prefer the DFA with few states. The Minimization can be done as follows.

Method:-

Consider below DFA



Step 1:- Separate states into two groups according to accepting states $\Rightarrow \{A, B, C, D\}, \{E\}$.

Step 2:- Consider both groups and inputs a, b. E can not be split, because it has only one state. $\{E\}$ remains in final group.

Step 3:- In group $\{A, B, C, D\}$, on input a, each state goes to B. On input b, A, B, C goes to members of $\{A, B, C, D\}$, D goes to $\{E\}$. So we split $\{A, B, C, D\}$, as $\{A, B, C\}, \{D\}$. The final group contains $\{A, B, C\}, \{D\}, \{E\}$.

Step 4:- Now in Group $\langle A, B, C \rangle$, ~~A and C~~
all states goes to members of $\langle A, B, C \rangle$
On input a, and on b, A and C goes to
members of Group, B. goes to D. So
we split $\langle A, C \rangle$, $\langle B \rangle$.

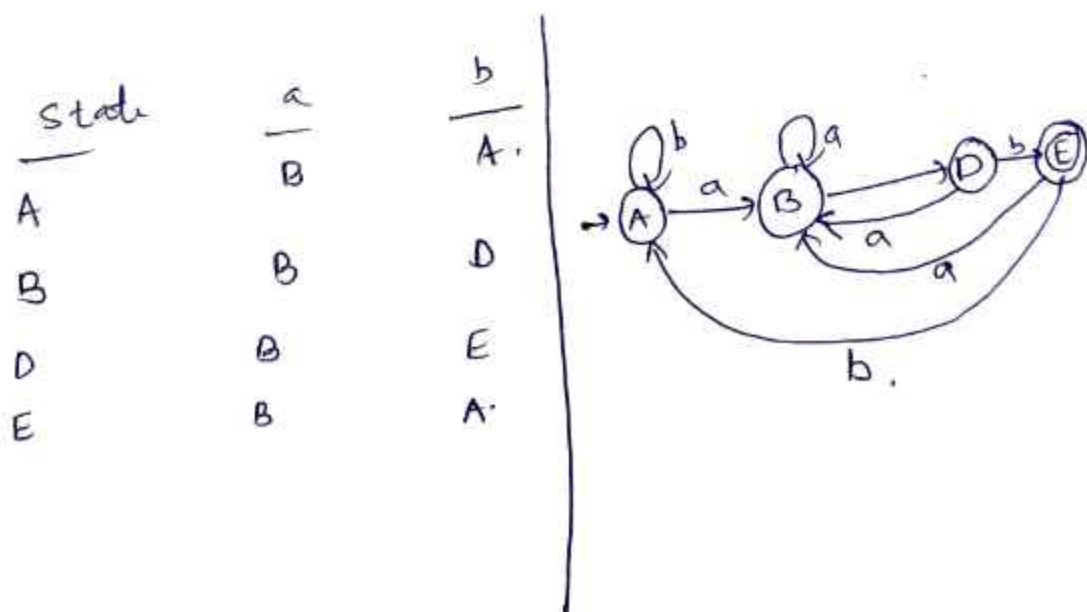
Step 5:- $\langle A, C \rangle$, can't be splitted, because both
goes to same states on input a and b.

The final Group consists of $\langle A, C \rangle$, $\langle B \rangle$, $\langle D \rangle$, $\langle E \rangle$.

The minimum DFA, consists of four states.
And we consider A as representative of $\langle A, C \rangle$.

The initial state is A, and final state is E.

* The Transition Table can be given as
and DFA.
follows.



SYNTAX ANALYSIS: (ROLE OF SYNTAX ANALYZER):

- Syntax Analyzer is also called as the parser, the parser obtains the stream of tokens from the lexical analyzer.
- And, verifies the tokens syntactically, if the tokens are syntactically correct then it will generate a parse tree.

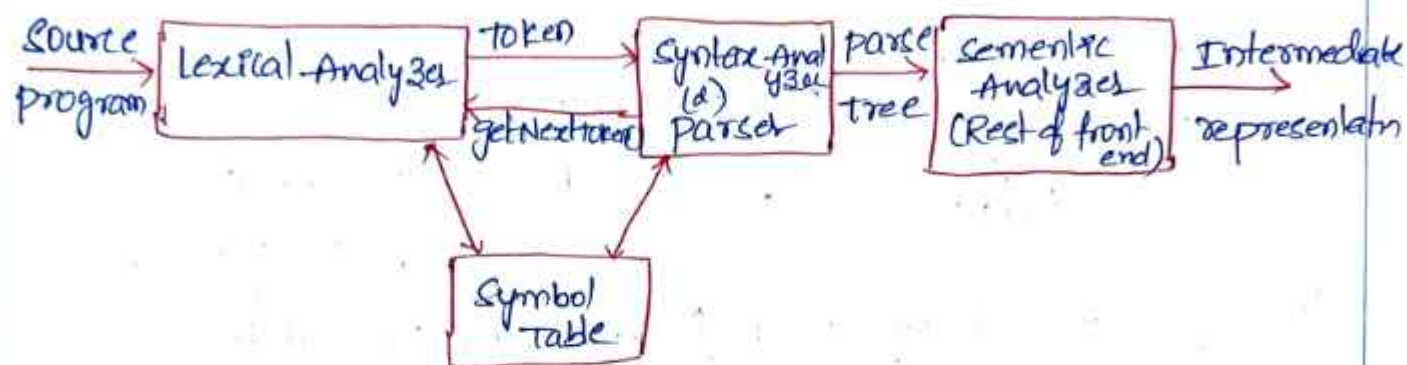
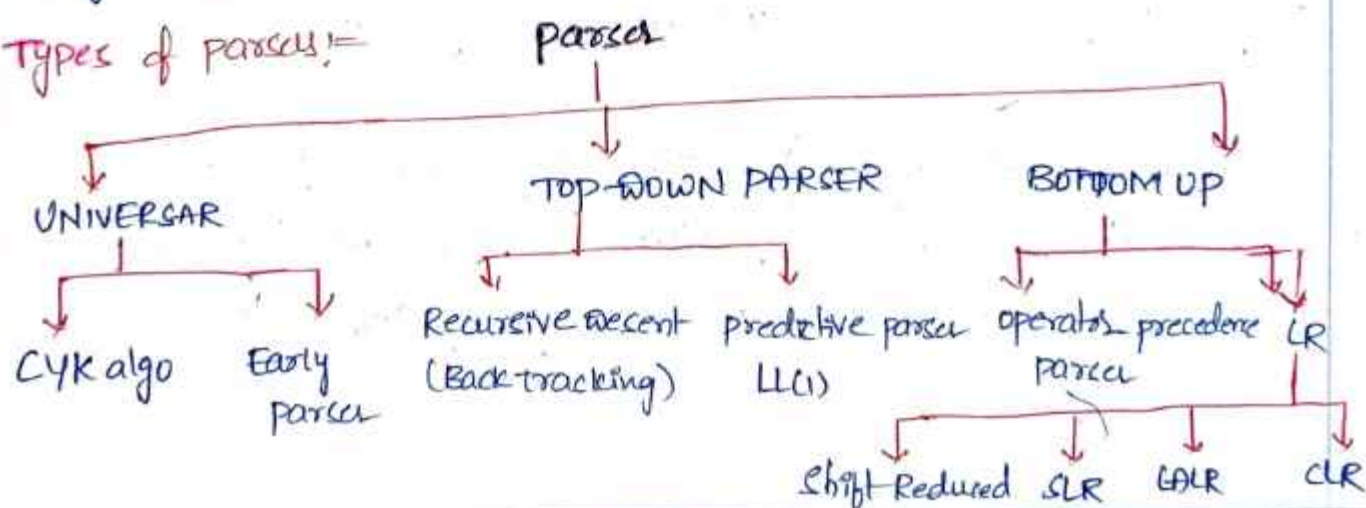


Fig: position of parser in compiler model.

- The main function of parser is to check the syntactic structure of ip code, if the given ip code is out of syntactic errors then it will generate a parse tree.
- if ip code consists of syntax error, then the parser reports an error to the user, now, the user will correct it. It is the responsibility of the user to correct those errors.
- Symbol Table communicates with both Lexical Analyzer & Syntax Analyzer for the token information.

Types of parsers:



→ In top-down parser construct the parse tree from top (root) to the bottom (leaves)

→ Bottom-up parser construct the parse tree from leaves to root (starts from ip symbol and reduced to starting state)

Representative Grammar:-

Associativity and precedence are captured in the following grammar

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

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

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

$$E \rightarrow \text{Expression}$$

$$T \rightarrow \text{Terms}, F \rightarrow \text{Factors that can be}$$

either parenthesized expression

These grammars belong to LR grammars that are suitable for bottom-up parsing. This grammar cannot be used for top-down parsing

The following non-left-recursive grammar will be used for top-down parsing

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

Syntax Error Handling:-

programming errors can occur at many different levels.

lexical errors:-

→ Include misspelling of identifiers, keywords, & operators

eg:- ^{identifier} ellipsesize instead of ellipsize and missing quotes around text intended as a string. (eg: "string");

Syntax Errors := Misplaced semicolon(;), extra & missing braces {}, {}.

Eg := appearance of case statement without an enclosing switch.

Semantic Errors := Eg := $(a + (b * c)) \rightarrow$ missing parenthesis.

→ Type mismatch b/w operator and operands.

Eg. $2 + a[i] \rightarrow$ type mismatch.

② The return of a value in Java method with result type void.

Logical Errors :=

→ errors can be anything from incorrect reasoning on the part of programming.

Eg := In "c" - the assignment operator instead of comparison operator.

if $(C = A)$ X if $(C == A)$ ✓

Errors Handles in parser :=

- should ^{report} the presence of errors clearly and accurately.
- should recovery from each error quickly enough to detect the subsequent errors.
- should not slow down the processing of remaining programs.

Errors Recovery Strategies :=

(i) Panic Mode Recovery :=

Eg :=

```
int a, b;
printf("%d", a);
```

→ The parser discards i/p symbols at a time until one of a designated set of synchronizing tokens is found.

→ The synchronizing tokens are delimiters, ";", "}" & }

→ It is simple to implement & does not goto a loop

(ii) phrase level Recovery :=

→ The parser perform local correction on remaining i/p when the error is discovered.

→ The parser replaces the prefix of the remaining i/p by some strings that allows the parser to carry on its execution

Eg:- replace a (,) comma by semicolon (;), delete an extra semicolon, insert a missing semicolon(;) .

① `print(" ")`, $\rightarrow$ `printf(" ")`;

Disadvantage:- Error correction is difficult when actual errors occur before the point of detection

(ii) Error production.

$\rightarrow$ common errors that can be encountered, we can augment the grammar for the language with productions that generate erroneous constructs. $\Rightarrow$

$\rightarrow$ use a new grammar for the parser.

(iv) Global correction:-

$\rightarrow$ The aim is to make some changes while converting incorrect i/p string to a valid string

Grammar G

- Given an incorrect i/p x , find a parse tree for a related string w (using the given Grammar) such that no. of changes (insertion/deletion) required to transform x to w is minimum
- Too costly to implement

(2) CONTEXT FREE GRAMMAR:-

Context Free Grammars consists of 4-tuples

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

- $V \rightarrow$ Set of variables (or) Non-terminals
- $T \rightarrow$ Set of terminals.
- $P \rightarrow$ Set of production Rules
- $S \rightarrow$ Start symbol.

(i) Non-terminal:-

$\rightarrow$ Denotes set of strings.

- Uppercase letters early in alphabet A, B, C, ...
- $S \rightarrow$ Start symbol
- E, T, F (Expression - E, Term - T, Factor - F)

(ii) Terminal:-

$\rightarrow$ Terminals are the basic symbols from which strings are formed

- Token-name is also called as terminal
- The following are terminal symbols.
 - (a) lowercase letters in the alphabets such as a, b, c.
 - (b) operator symbols such as +, *, ...
 - (c) punctuations (" ", "(", ")", ",", ":", "..." & digits 0, 1, ... 9

(iii) production:- consists of

- (a) A Non terminal called head or left side of the production.

This production defines some of the symbols denoted by the head.

- (b) body or right side $\rightarrow$ consisting of zero or more terminals and Non terminals.

* uppercase letters late in alphabet X, Y, Z represent Grammar symbols either as terminals or Non-terminals.

→ $\alpha, \beta, \gamma \dots$ represent string of grammar symbol

→ A set of productions $A \rightarrow \alpha_1, A \rightarrow \alpha_2 \rightarrow A \rightarrow \alpha_k$ also written as $A \rightarrow \alpha_1 \mid \alpha_2 \mid \dots \mid \alpha_k$

Product $A \rightarrow \alpha$
 $\downarrow \quad \downarrow$
 Non-Terminal either terminal & Nonterminal (VNT)*

Eg:- Consider the grammar

$E \rightarrow E + E \quad E \rightarrow E * E \quad E \rightarrow I \quad I \rightarrow a$

$CFG = (\{E, I\}, \{+, *, a\}, P, E)$

→ derive

Derivations:- Derivation is process of applying a sequence of production rules in order to derive a string.

→ deriving an i/p string from the start symbol of Grammar in one or more steps by replacing the head of the production by its body of the production.

There are two types of derivation.

(i) Left most Derivation (LMD)

(ii) Right most Derivation (RMD)

Left most Derivation:- In each step

we have to expand left-most Non-terminal by one of its production body.

Right most Derivation:-

→ here, we have to expand the Right-most non-terminal

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

Derive a string $id + id * id$

LMD:- $E \xRightarrow{LMD} E + E \quad (E \rightarrow E + E)$

$E \xRightarrow{LMD} id + E \quad (E \rightarrow id)$

$\xRightarrow{LMD} id + E * E \quad (E \rightarrow E * E)$

$\xRightarrow{LMD} id + id * E \quad (E \rightarrow id)$

$\xRightarrow{LMD} id + id * id \quad (E \rightarrow id)$

RMD:-

$E \xRightarrow{RMD} E + E \quad (E \rightarrow E * E)$

$\xRightarrow{RMD} E + E * E \quad (E \rightarrow E * E)$

$\xRightarrow{RMD} E + E * id \quad (E \rightarrow id)$

$\xRightarrow{RMD} E + id * id \quad (E \rightarrow id)$

$\xRightarrow{RMD} id + id * id \quad (E \rightarrow id)$

Parse tree:-

→ A Graphical representation of derivation is called parse tree.

These types of nodes in parse tree

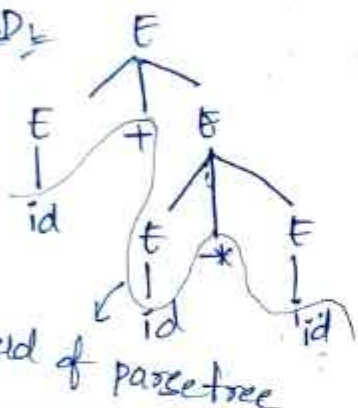
(i) Interior node → are Non terminals

(ii) children node → terminals

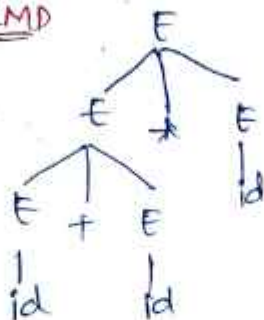
→ In parse tree the root node must be start symbol

construct parse tree for i/p string $W = id + id * id$

LMD:-



RMD



Ambiguity:-

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

→ If a grammar produces more than one parse tree, then that grammar is called as ambiguous.

(i) more than one LMD (different)

(ii) more than one RMD (different).

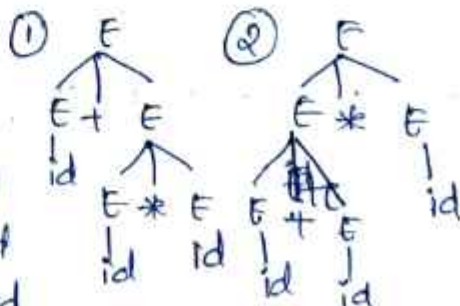
Eg $E \rightarrow E + E \mid E * E \mid (E) \mid id$ string $w = id + id * id$

LMD:-

$E \Rightarrow E + E$
 $\Rightarrow id + E$
 $\Rightarrow id + E * E$
 $\Rightarrow id + id * E$
 $\Rightarrow id + id * id$

RMD:-

$E \Rightarrow E + E$
 $\Rightarrow E + E * E$
 $\Rightarrow E + E * id$
 $\Rightarrow E + id * id$
 $\Rightarrow id + id * id$

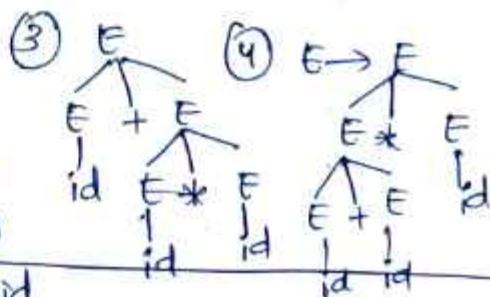


LMD:-

$E \Rightarrow E * E$
 $\Rightarrow E + E * E$
 $\Rightarrow id + E * E$
 $\Rightarrow id + id * E$
 $\Rightarrow id + id * id$

RMD:-

$E \Rightarrow E * E$
 $\Rightarrow E * id$
 $\Rightarrow E + E * id$
 $\Rightarrow E + id * id$
 $\Rightarrow id + id * id$



→ For the above string we got 2 LMD & RMD and 2 parse trees
So, the given grammar is ambiguous

→ Top down parser can't handle ambiguous grammar, so we need to convert this grammar into unambiguous grammar.

* While converting ambiguous to unambiguous:-

→ The grammar should follow Associativity, and precedence rules

$*, /, +, -$ → left Associative (when an operand has operand on both sides, the operand should associate with left side operator)

$1, *, +, -$ → precedence order
1 2 3 4

② WRITING A GRAMMAR:-

→ Grammar is used to describe the syntax of programming language.

Lexical versus Syntactic Analysis:-

- separating syntactic structure into smaller and manageable components of lexical and non-lexical parts.
- lexical rules are simple to understand than grammar.
- We need notations to describe the grammar.
- RE are used to Realtime Examples to make them easy to understand.
- Grammars are used for describing nested structure such as balanced parenthesis, matching begin-end, if-then-else.

Eliminating Ambiguity:-

Properties of CFG

1. Ambiguous
2. Left Recursive
3. Left factoring.

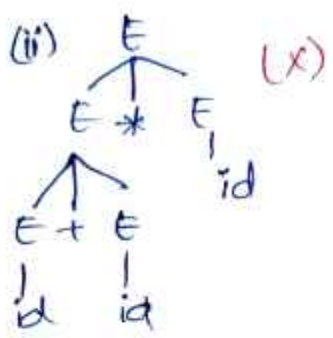
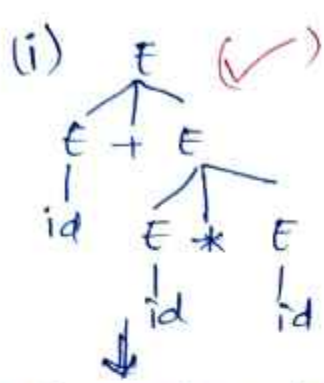
Elimination of ambiguity:-

Consider CFG:-

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

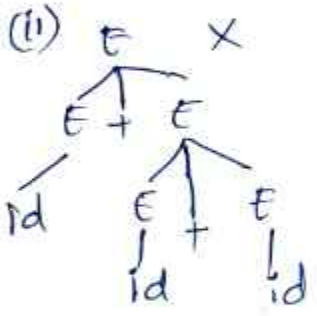
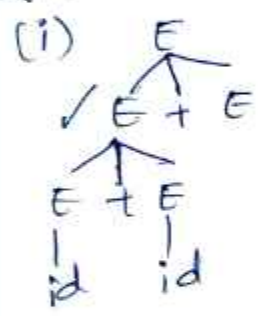
$w = id + id * id$

(There are 2 parse trees for string)



→ This is the valid parse tree, because the "*" operator is appeared at bottom of parse, so it is evaluated first

Eg2:- $w = id + id + id$



→ two factors that are needed to be taken

(1) precedence

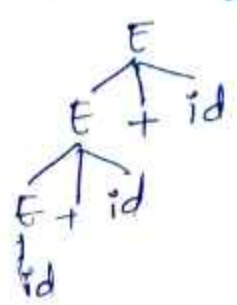
(2) left associativity → To ensure left associativity, we need to convert the grammar into left recursive ^{of RHS}

→ left recursive $E \rightarrow E + E$ (the leftmost symbol is equal to the LHS)

• left recursive, the parse tree can be grown on left side only

$E \rightarrow E + id \mid id$

$id + id + id \Rightarrow (id + id) + id$ (left associative)



left associativity can be achieved

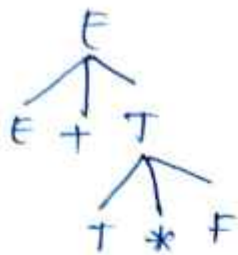
$(id + id) + id$

↓
This expression is evaluated first, so that it is valid parse tree.

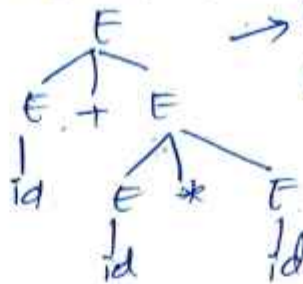
$$E \rightarrow E + T$$

$$T \rightarrow T * F$$

(precedence)

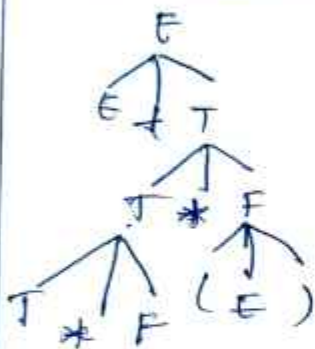
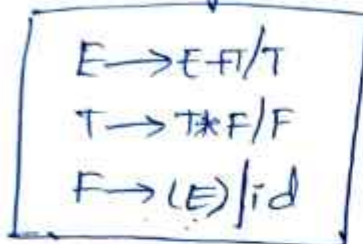


$id + (id * id)$



highest precedence operation is executing first.

The Final ~~good~~ unambiguous Grammars



Dangling else:- In a program if there are more than one if-
stmts then else part is matched with wrong if
stmt, that will lead to wrong results, that is called dangling
else.

Eg $\boxed{\text{if (expr1)} \rightarrow P_1 \rightarrow T}$ (i) if P_1 is True Then it will goto P_2 and
 $\boxed{\begin{array}{l} \text{if (expr2)} \rightarrow T \\ S_1 \\ \text{else} \\ S_2 \end{array}}$ • check the condⁿ 2nd condition, if it also true
 then it will execute S_1
 • if 2nd condition is false it will execute else
 part i.e S_2

other stmt

(ii) if $P_1 \rightarrow$ is false it has to execute the other
statement

(ii) But here if P_1 - False then it is executing else part i.e S_2
it is matching with wrong else.

- considers the following dangling else grammars.

$\text{stmt} \rightarrow \text{if Expr then stmt}$

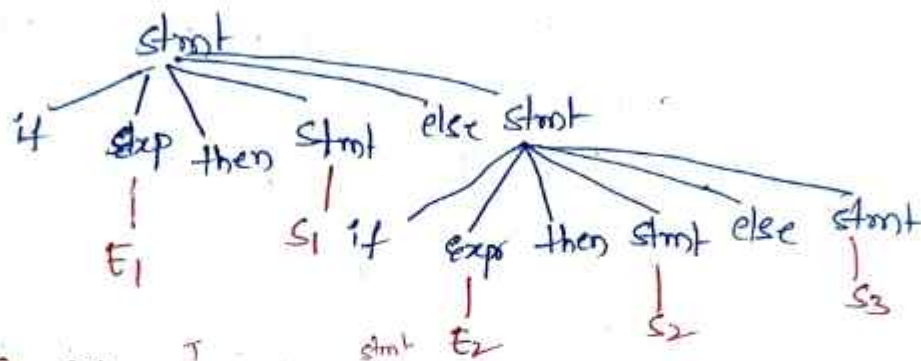
$\text{if Expr then stmt else stmt}$

other

- The compound conditional stmt for the above grammars is

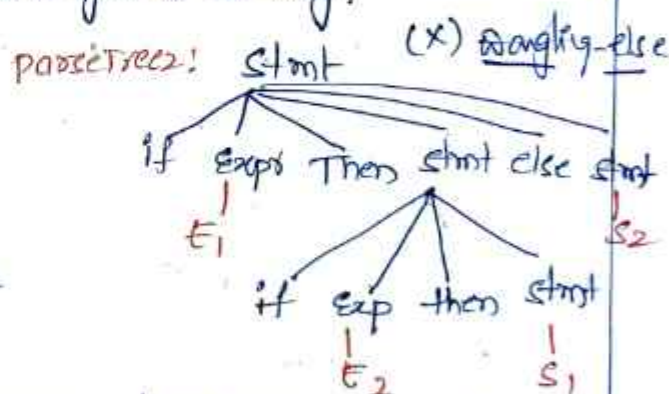
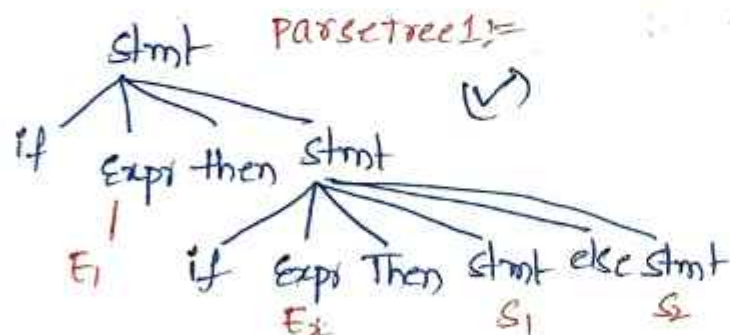
$\text{if } E_1 \text{ then } S_1 [\text{else if } E_2 \text{ then } S_2 \text{ else } S_3]$

- The parse for this stmt



ex i/p string = $\text{if } E_1 \text{ then } [\text{if } E_2 \text{ then } S_1 \text{ else } S_2]$

there exist two parse trees for the given string.



→ In these two parse trees, in programming language, the 1st parse tree is considered because, match the else statement with closest then. if E_1 then [if E_2 then S_1 else S_2] → matched stmt

→ so, The given Grammar is ambiguous Grammar.

The unambiguous grammar will be.

$\text{stmt} \rightarrow \text{matched stmt} \mid \text{open stmt}$

matched stmt → perfect if
open stmt → simple if

matched stmt → $\text{if Expr then matched stmt else matched stmt} \mid \text{other}$

open stmt → if Expr then stmt

$\text{if Expr then matched stmt else open stmt.}$

$$(i) \begin{array}{l} T \rightarrow T * F / F \\ A \quad A \alpha / \beta \end{array}$$

$$T \rightarrow FT'$$

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

$$\text{eg 2: } \begin{array}{l} S \rightarrow \underline{S} \underline{O} \underline{S} \underline{S} / \underline{O} \underline{I} \\ A \quad A \quad \alpha \quad \beta \end{array}$$

$$A \rightarrow A \alpha / \beta$$

$$S \rightarrow \cancel{O} \underline{O} \underline{S} \underline{S}$$

$$S' \rightarrow \underline{O} \underline{S} \underline{S} \underline{S}' / \epsilon$$

$$(ii) F \rightarrow (E) / id \text{ (Non left recursive)}$$

After eliminating the left recursion from the grammar

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

$$F \rightarrow (E) / id / \epsilon$$

$$\text{eg 3: } \begin{array}{l} S \rightarrow (L) / \epsilon \text{ - NO left recursion} \\ L \rightarrow L S / S \end{array}$$

* Elimination of left recursion for multiple production.

$$A \rightarrow A \alpha_1 / A \alpha_2 / \dots / \beta_1 / \beta_2 / \dots$$

$$\boxed{\begin{array}{l} A \rightarrow \beta_1 A' / \beta_2 A' \\ A' \rightarrow \alpha_1 A' / \alpha_2 A' / \epsilon \end{array}}$$

$$\text{eg 1: } \begin{array}{l} \text{Expr} \rightarrow \text{Expr} + \text{Expr} / \text{Expr} * \text{Expr} / id \\ A \quad A \quad \alpha_1 \quad A \quad \alpha_2 \quad \beta_1 \end{array}$$

$$\text{Expr} \rightarrow id \text{Expr}'$$

$$\text{Expr}' \rightarrow + \text{Expr} \text{Expr}' / * \text{Expr} \text{Expr}' / \epsilon$$

$$\text{eg 2: } \begin{array}{l} S \rightarrow Sx / \underline{S} \underline{S} \underline{b} / \underline{x} \underline{S} / a \\ A \quad A \alpha_1, A \alpha_2 \quad \beta_1 \quad \beta_2 \end{array}$$

$$S \rightarrow x \underline{S} \underline{S}' / a \underline{S}'$$

$$S' \rightarrow x \underline{S}' / \underline{b} \underline{S}' / \epsilon$$

$$\text{eg 3: } S \rightarrow Aa / b$$

$$A \rightarrow Ac / Sd / f$$

$$S \rightarrow Aa / b \rightarrow \text{NO left recursion}$$

$$A \rightarrow Ac / Sd / f$$

(1) Elimination of left factoring =

→ A grammar contains production rule in the form

$$A \rightarrow \alpha B_1 / \alpha B_2 / \dots / r_1 / r_2$$

Then that grammar contains left factoring.

→ TOP parsers can't handle left factoring the grammar contains left factoring.

→ we can eliminate the left factoring by replacing with the following production.

$$\begin{aligned} A &\rightarrow \alpha A' / r_1 / r_2 \\ A' &\rightarrow B_1 / B_2 \end{aligned}$$

Eg 1: $S \rightarrow \underline{iets} / \underline{ietses} / a$ (it contains left factoring)

$$E \rightarrow b$$

Here, $A = S$, $\alpha = iets$, $B_1 = E$, $B_2 = es$, $r_1 = a$

$$S \rightarrow ietsd / a$$

$$S' \rightarrow E / es$$

Eg 3: $B \rightarrow bB / b$
 $A \quad \alpha B \quad B_2$
 $B \rightarrow bB$
 $B' \rightarrow B / E$

Eg 2: $A \rightarrow \underline{aAB} / \underline{aA} / a$

$$B \rightarrow bB / b$$

$$A \rightarrow \underline{aAB} / \underline{aA} / a$$

$$A \rightarrow aA'$$

$$A' \rightarrow AB / A / E$$

Eg 4: $X \rightarrow \underline{X+X} / \underline{X * X} / D$
 $A \quad \alpha B_1 \quad \alpha B_2 \quad r_1$

$$X \rightarrow X X' / D$$

$$X' \rightarrow +X / *X$$

Eg 5: $E \rightarrow T + E / T$

$$T \rightarrow \text{int} / \text{int} * T / (E)$$

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

$$A \quad \alpha B_1 \quad \alpha B_2 = E$$

$$E \rightarrow TE'$$

$$E' \rightarrow +E / E$$

$$T \rightarrow \text{int} / (\text{int} * T / (E))$$

$$T \rightarrow \text{int} T' / (E)$$

$$T' \rightarrow E / *T$$

Top Down parsing:-

→ In Top-down parsing, The parse tree is constructed from root node to child node.

→ Top-down parser uses left-most derivation to derive the input string from the grammar.

→ TDP is starting from start symbol of grammar and reaching the i/p string

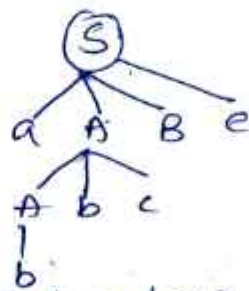
Eg:- Grammar is, derive the i/p string $w = abbcde$.

$$S \rightarrow aABe$$

$$A \rightarrow Abc/b$$

$$B \rightarrow d$$

LMD:- $S \rightarrow aABe$
 $\quad \quad \quad aAbcBe$
 $\quad \quad \quad abbcBe$
 $\quad \quad \quad abbcde$



→ TDP constructed for the grammar if it is free from
 • ambiguity • left recursion.

→ The problem with Top parser is if we have more than one alternative for production, which alternative we should use. $\begin{cases} E \rightarrow E + T \\ T \rightarrow T * F / F \\ F \rightarrow (E) / id \end{cases}$
 ↓ ambiguous, left recursion

Eg:- $E \rightarrow TE'$ → unambiguous grammar

$$E' \rightarrow +TE'$$

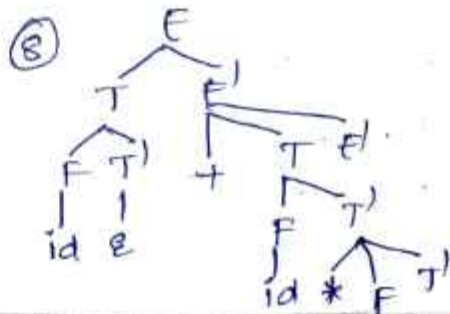
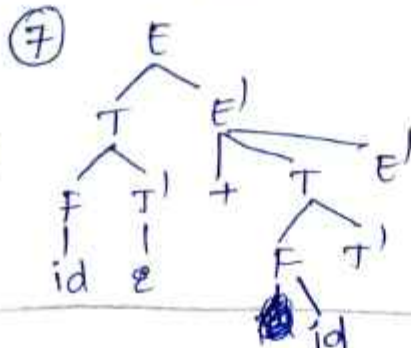
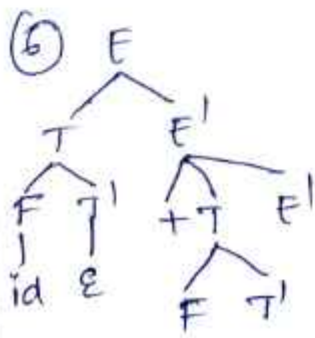
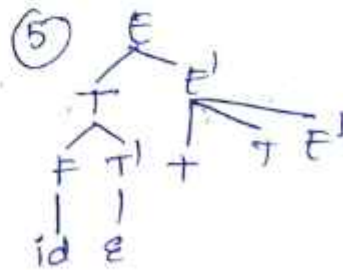
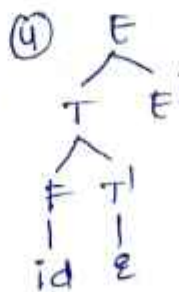
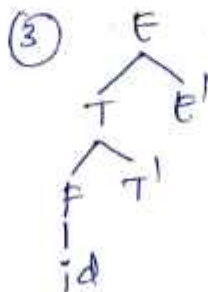
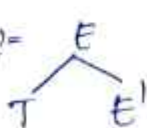
$$T \rightarrow FT'$$

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

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

→ derive an i/p string $w = id + id * id$

step 1:-



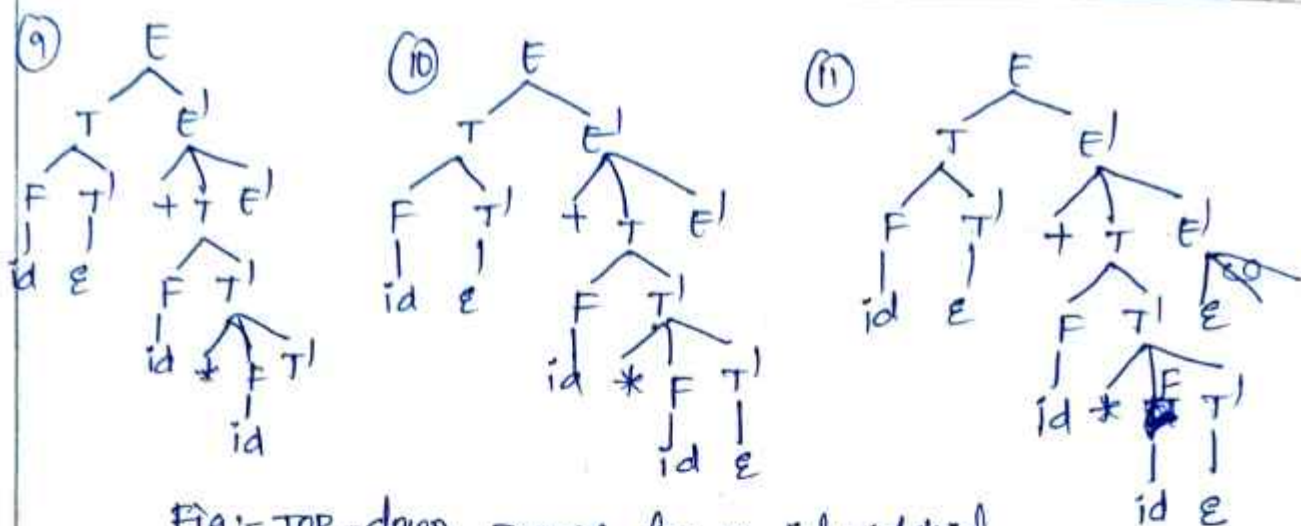


Fig. = TOP-down parser for $w = id + id * id$

(i) Recursive-descent parsing := (with back tracking)

→ The construction of parse tree starts from root & proceed to child node.

→ Recursion := A function which is called by itself.

Steps for constructing Recursive descent parser :=

(i) If the i/p is a non-terminal, then call corresponding function.

(ii) If the i/p is terminal, then compare terminal with i/p symbol.

If both are same, the increment input pointer.

(iii) If a Non-terminal contains more than one production then all the production code should be written in the corresponding function.

Ex 1: $E \rightarrow iE \mid E \mid \epsilon$

(In the given grammar we have 2-Nonterminals E, E' , so define two functions, same as C-lang functions)

$E()$

{

if (input == 'i')

input++;

PRIME();

}

EPRIME();

i/p string

Eg2:- Grammar $E \rightarrow TE' \quad E' \rightarrow +TE' / \epsilon$
 $T \rightarrow FT \quad T' \rightarrow *FT' / \epsilon \quad F \rightarrow (E) / \epsilon$

```

{
  if (input == '+')
  {
    // skip input++
    if (input == '(')
    {
      input++;
      EPRIME();
    }
  }
  else
  {
    return;
  }
}

```

i/p string: i+i

```

E()
{
  T();
  EPRIME();
}
EPRIME();
{
  if (input == '+')
  {
    input++;
    T();
    EPRIME();
  }
  else
  {
    return;
  }
}

```

```

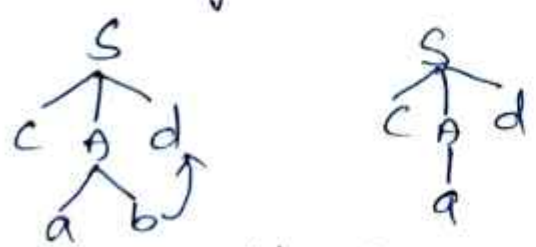
T()
{
  F();
  TPRIME();
}
TPRIME()
{
  if (input == '*')
  {
    // skip
    input++;
    F();
    TPRIME();
  }
  else
  {
    return;
  }
}
F()
{
  if (input == '(')
  {
    input++;
    E();
  }
  if (input == ')')
  {
    input++;
  }
  else if (input == 'id')
  {
    input++;
  }
}

```

Eg3:- $S \rightarrow CAD$

$A \rightarrow ab|a$

w = i/p string w = cad



Backtracking

(If None of the product for is satisfying the condition then we should check on another possibility for "S")


```

input ++
EPRIME();
}

```

```

EPRIME()
{
  if (input == '+')
  {
    input ++;
    if (input == 'i')
    {
      input ++;
      EPRIME();
    }
    else
      return;
  }
}

```

$i/p = id + id$
 $eg: 2 \rightarrow TE$

$E' \rightarrow TE|E$

$T \rightarrow FT$

$TL \rightarrow *FT|E$

$F \rightarrow (E) | id$

$E()$

$T()$

$EPRIME();$

$EPRIME()$

```

{
  if (input == '+')
  {
    input ++;
  }
}

```

```

T()
EPRIME();
}
else
  return;
}

```

$T()$

```

{
  F()
  TPRIME();
}

```

$TPRIME()$

```

{
  if (input == '*')
  {
    input ++;
  }
  F()
  TPRIME();
}
else
  return;
}

```

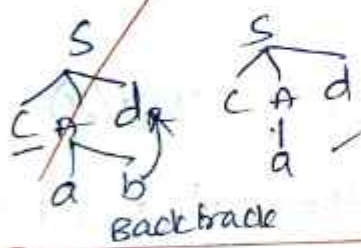
```

F()
{
  if (input == '(')
  {
    input ++;
    E();
  }
  if (input == ')')
  {
    input ++;
  }
  else
  {
    if (input == 'd')
    {
      input ++;
    }
  }
}

```

$i/p: id + id$

$eg \ S \rightarrow CAD$
 $A \rightarrow ab|a$ $w = cad$



($\therefore$ if Non of the product for A is satisfying the condition, then we should check another possibility for S)

(i) predictive parsing: (a) LL(1) & Non recursive parser:-

LL(1) - means

- It scans the input from left to Right by using leftmost derivation.

- LL(1) $\rightarrow$ 1 - no look ahead - at a time we are reading 1 character at a time.

input $\rightarrow$ $\leftarrow$ i/p buffer



LL(1) parser (a) predictive parser

predictive parsing table

$\rightarrow$ o/p buffer

Steps for constructing LL(1) parser:-

- ① Elimination of left recursion
- ② Elimination of left factoring
- ③ calculation of FIRST() & FOLLOW
- ④ construction of parsing table.

LL(1) parser $\rightarrow$ parsing algorithm

LL(1) parsing table - data structure which is constructed from the LL(1) grammars.

Stack := data structure used to store the grammatical symbols

• Always the bottom of the stack is $\$$

$\rightarrow$ The end of ip buffer is $\$$.

$\rightarrow$ after $\$$, the starting symbol of ip is pushed to the stack.

To construct the parsing table we should compute

two functions. (1) FIRST (2) FOLLOW

FIRST(x) :=

• if x is a terminal then $\text{FIRST}(x) = \{x\}$

• if x is a Non terminal & $x \rightarrow y_1 y_2 y_3 \dots y_n$ is a production for x .

$\text{FIRST}(x) = \text{FIRST}(y_1)$ &

if $\text{FIRST}(y_1)$ contains ϵ add $\text{FIRST}(y_2)$ to $\text{FIRST}(x)$

if all $\text{FIRST}(y_1, y_2, \dots, y_n)$ contains ϵ then add ϵ to $\text{FIRST}(x)$

FOLLOW(B) :=

$\rightarrow$ Always the follow of starting symbol is $\$$.

$\rightarrow$ if $A \rightarrow \alpha B \beta$

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

if $\text{FIRST}(\beta)$ contains ϵ or $A \rightarrow \alpha B$ then

add $\text{FOLLOW}(A)$ to $\text{FOLLOW}(B)$

Eg := Given Grammar is

$E \rightarrow E + T / T$ - left recursion

$T \rightarrow T * F / F$ - left recursion

$F \rightarrow (F) / id$

} Eliminate left recursion & left factoring from these 2 productions

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

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

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

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

$$\bullet \text{ FIRST}(T') = \{ *, \epsilon \} \quad \bullet \text{ FIRST}(F) = \{ (, \text{id})$$

$$\rightarrow \text{Follow}(E) = \{ \$, \text{Follow}(F) \}$$

$$\downarrow \quad F \rightarrow (E) \text{ id}$$

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

$$\textcircled{2} \text{ Follow}(E) = \{ \text{Follow}(E) + \text{Follow}(E') \}$$

$$E \rightarrow TE' \quad = \{ \$,) \}$$

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

$$\text{Follow}(T) = \{ \text{Follow}(E) + \text{Follow}(T') \}$$

$$E \rightarrow TE'$$

$$\text{Follow}(T) = \text{FIRST}(E') + \text{Follow}(E)$$

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

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

$$\text{Follow}(T') \Rightarrow \{ \text{Follow}(T) + \text{Follow}(T') \}$$

$$T \rightarrow FT'$$

$$\Rightarrow \{ +, \$,) \}$$

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

$$\text{Follow}(F) = \text{Follow}(T') + \text{FIRST}(T) + \text{Follow}(T)$$

$$T \rightarrow FT'$$

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

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

(1111)
predictive parsing table :-

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

$w = id + id$

Stack	Input	Action
<u>E</u> \$	<u>id</u> + id \$	$E \rightarrow TE'$
<u>TE'</u> \$	<u>id</u> + id \$	$T \rightarrow FT'$
<u>FT'</u> \$	<u>id</u> + id \$	$F \rightarrow id$
<u>id</u> <u>TE'</u> \$	<u>id</u> + id \$	
<u>T'</u> <u>E'</u> \$	+ id \$	$T' \rightarrow \epsilon$
<u>E</u> <u>E'</u> \$	+ id \$	
<u>E'</u> \$	+ id \$	$E' \rightarrow +TE'$
<u>+TE'</u> \$	<u>id</u> \$	$T \rightarrow FT'$
<u>FT'</u> <u>E'</u> \$	<u>id</u> \$	$F \rightarrow id$
<u>id</u> <u>TE'</u> \$	<u>id</u> \$	$T' \rightarrow \epsilon$
<u>E</u> <u>E'</u> \$	\$	$E' \rightarrow \epsilon$
<u>\$</u>	<u>\$</u>	accepted

The i/p is parsed

eg 2: $S \rightarrow (L) / a$
 $L \rightarrow L, S / S$
 $w = (a)$

eg 3: $S \rightarrow aAB / bA / \epsilon$
 $A \rightarrow aAb / \epsilon$
 $B \rightarrow bB / \epsilon$
 $w = "aabb"$

$S \rightarrow iet s / iet ses / a$
 $E \rightarrow b$ (Non-terminating)
 eg 4: $S \rightarrow aAB / cbB / B$
 $A \rightarrow da / B$
 $B \rightarrow g / \epsilon$
 $C \rightarrow h / \epsilon$

Bottom-up parses :-

→ Bottom-up parses construct parse tree from child (bottom) node to root node (top).

→ Bottom up parses use RMD in Reverse order.

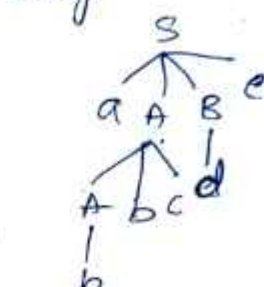
→ In Bottom-up parse the ~~reducer~~ ^{reducer} i/p string w to start symbol of grammar

Eg: $S \rightarrow AABE$

$A \rightarrow abc/b$

$B \rightarrow d$

Input string $w = abcde$



$aAbcde$ ($A \rightarrow b$)

$aAde$ ($A \rightarrow abc$)

$AABe$ ($B \rightarrow d$)

S ($S \rightarrow AABe$)

$S \rightarrow AABE$

$\rightarrow aAde$

$\rightarrow aAbcde$

$\rightarrow abcde$

Eg2: $E \rightarrow E+T/T$

$E \rightarrow T*F/F$

$F \rightarrow (E)/id$

$w = id*id$

$id*id$ ($F \rightarrow id$)

$F*id$ ($T \rightarrow id$)

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

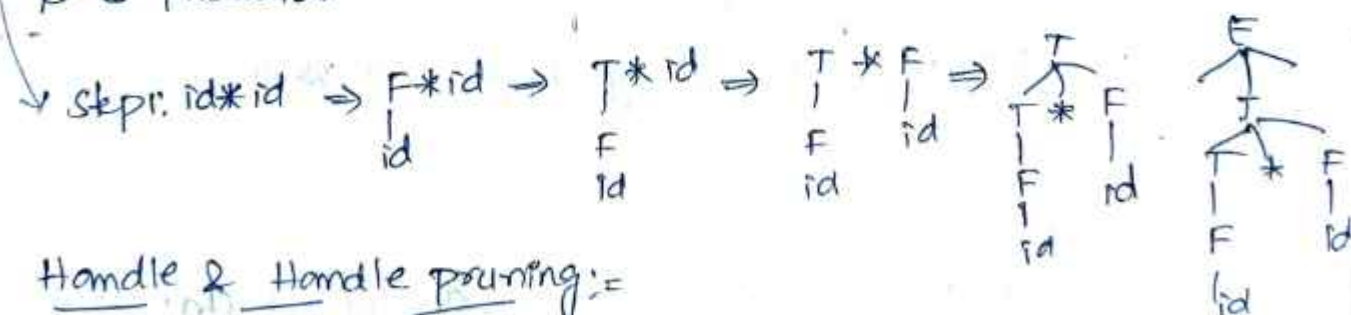
$E*id$ ($F \rightarrow id$)

$T*F$ ($F \rightarrow T*F$)

T ($E \rightarrow T$)

E

→ The problem with Bottom-up parses is when to reduce the ~~pr~~ to production.



Handle & Handle pruning:-

Handle: Handle is a substring, which matches with right side of production.

→ if Handle matches with the right side of production, then it is replaced with left hand side Non-terminal

eg1:- $S \rightarrow aABC$ $w = abbcd$
 $A \rightarrow abc/b$
 $B \rightarrow d$

Right Sentential form	Handle	Reducing production
abbcde	b	$B \rightarrow b$
a ab cde	abc	$A \rightarrow abc$
aAde	d	$B \rightarrow d$
aABe	aABe	$S \rightarrow aABe$
S		

Handles during parse of $id, *, id_2$

eg2:- $E \rightarrow E + T / T$ $w = id, *, id_2$
 $T \rightarrow T * F / F$
 $F \rightarrow (E) id$

Right Sentential form	Handle	Reducing production
<u>id</u> * id	E id	$F \rightarrow id$
F * id	F	$T \rightarrow F$
T * id	id	$F \rightarrow id$
T * F	T * F	$T \rightarrow T * F$
T	T	$E \rightarrow T$

→ Handle pruning is obtained by Right most derivation in reverse order.

(f) Shift-Reduce parser :-

→ Shift-Reduce parser use two data structures ① stack ② i/p buffer

① Stack - is used^{to} store the grammar symbol

② Input buffer - holds the i/p string to be parsed

Stack Input string
 $\$$ $w\$$

Actions of shift-Reduce parser:

1. shift — shift the next I/p symbol onto the top of the stack.
2. Reduce — If the top of stack is matched with Right side of production, then it is reduced to corresponding non-terminal.
3. Accept — ~~a~~ successful completion of parsing
4. Error — Discovers a ^{syntax} error and call error recovery methods

eg: $E \rightarrow E + T / T$ $w = id * id$

$T \rightarrow T * F / F$

$F \rightarrow (E) / id$

Fig: configuration of shift-Reduce parser on $id * id$

Stack	Input Buffer	Action
\$	<u>id</u> * id \$	shift
\$ <u>id</u>	* id \$	Reduce by $F \rightarrow id$
\$ F	* id \$	Reduce $T \rightarrow F$
\$ T	* id \$	shift
\$ T *	id \$	shift
\$ T * <u>id</u>	\$	Reduce $F \rightarrow id$
\$ T * F	\$	Reduce $T \rightarrow T * F$
\$ T	\$	Reduce $E \rightarrow T$
\$ E	\$	Accepted

Conflicts during shift-Reduce parsing

- ① shift-Reduce conflict: cannot decide whether to shift or to reduce
- ② reduce-reduce conflict: if the ^{two} productions have same handle (subtring).

Shift-Reduce

Eg2:- $S \rightarrow (L) \cdot a$ $L \rightarrow L, S / S$ parse the string $(a, (a, a))$ using SR-Parser

Stack	I/P Buffer	Parser Action
\$	$a(a, (a, a))\$$	shift
$(a$	$a, (a, a))\$$	shift
$(a,$	$(a, a))\$$	Reduce $S \rightarrow a$
$(S$	$(a, a))\$$	$L \rightarrow S$ Reduce
$(L$	$(a, a))\$$	shift
$(L,$	$(a, a))\$$	shift
$(L,$	$a, a))\$$	shift
$(L, ($	$a, a))\$$	shift
$(L, (a$	$a, a))\$$	Reduce $S \rightarrow a$
$(L, (S$	$a, a))\$$	shift Reduce $L \rightarrow S$
$(L, (L$	$a, a))\$$	shift
$(L, (L,$	$a, a))\$$	shift
$(L, (L, a$	$a, a))\$$	Reduce $S \rightarrow a$
$(L, (L, S$	$a, a))\$$	Reduce $L \rightarrow L, S$
$(L, (L$	$a, a))\$$	shift
$(L, (L)$	$a, a))\$$	Reduce $S \rightarrow (L)$
$(L, S$	$a, a))\$$	Reduce $L \rightarrow L, S$
$(L$	$a, a))\$$	shift
(L)	$a, a))\$$	Reduce $S \rightarrow (L)$
S	$a, a))\$$	Accept

eg:- $S \rightarrow TL;$
 $T \rightarrow \text{int}/\text{float}$
 $L \rightarrow L \text{ id}/\text{id}$

$w = \text{int id id}$ (4) $S \rightarrow 050/151/2$
 $w = 10201$

Introduction to LR parsing := LR parser is a type of Bottom-up Parser

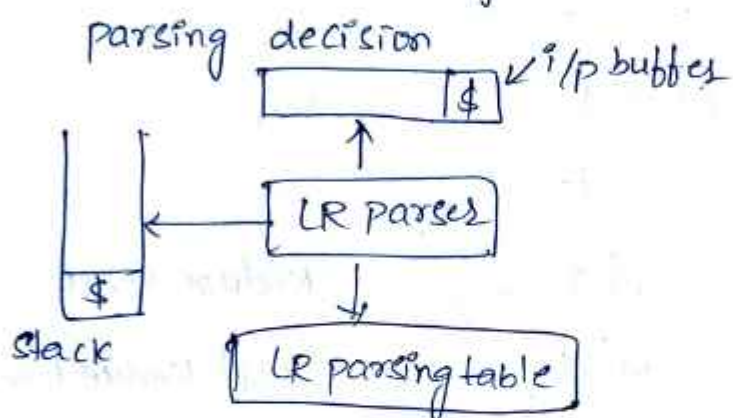
→ LR^(k) parses scans i/p from left to Right order.

→ It use Right most Derivation in Reverse order to

derive the to reduce i/p string to start symbol of grammar

$k \rightarrow$ No. of lookahead symbols that are used to make

parsing decision



Stack - Datastructure used to store grammar symbols.

i/p - " " " " " i/p string to be parsed

LR parser → uses algorithm to make decisions.

LR parsing table - is constructed by using LR(0) items.

- These table uses two functions (i)

(i) closure

(ii) Action.

Benefits of LR(k) parsing :=

→ most generic Non-Backtracking shift Reduced parsing Technique.

→ These parsers can recognize all programming languages for which CFG can be written.

→ They are capable of detecting syntactic errors as soon as possible while scanning of i/p.

Simple LR parsing:

15

Eg: $E \rightarrow E+T$

$E \rightarrow T$

$T \rightarrow T * F$

$T \rightarrow F$

$F \rightarrow (E)$

$F \rightarrow id$

① Augmented Grammars

$E' \rightarrow E$ $E' \rightarrow E$

$E \rightarrow E+T$

$E \rightarrow T$

$T \rightarrow T * F$

$T \rightarrow F$

$F \rightarrow (E)/id$

Steps to construct SLR parser

① Introduce augmented grammars

② Calculate canonical collection of LR(0) items.

③ Construct SLR parsing table by using (i) Goto (ii) Action functions

Item sets:

$E' \rightarrow \cdot E$

$E \rightarrow \cdot E+T$

$E \rightarrow \cdot T$

$T \rightarrow \cdot T * F$

$T \rightarrow \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow \cdot id$

Goto (I_0, E) - I_1

$E' \rightarrow E \cdot$

$E \rightarrow E \cdot +T$

(I_0, T) - I_2

$E \rightarrow T \cdot$

$T \rightarrow T \cdot * F$

(I_0, F) - I_3

$T \rightarrow F \cdot$

($I_0, ($) - I_4

$F \rightarrow (\cdot E)$

$E \rightarrow \cdot E+T$

$E \rightarrow \cdot T$

(I_0, id) - I_5

$F \rightarrow id \cdot$

($I_1, +$) - I_6

$E \rightarrow E+ \cdot T$

$T \rightarrow \cdot F$

$T \rightarrow \cdot T * F$

$F \rightarrow \cdot (E)$

$F \rightarrow \cdot id$

($I_2, *$) - I_7

$T \rightarrow T * \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow \cdot id$

(I_4, E) - I_8

$F \rightarrow (E) \cdot$

$E \rightarrow E \cdot +T$

(I_4, T) - I_9

$E \rightarrow T \cdot$

$E \rightarrow E \cdot +T$

$T \rightarrow T \cdot * F$

(I_4, F) - I_{10}

$T \rightarrow F \cdot$

($I_4, ($) - I_{11}

$F \rightarrow (\cdot E)$

$E \rightarrow \cdot E+T$

$E \rightarrow \cdot T$

$T \rightarrow \cdot T * F$

$T \rightarrow \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow \cdot id$

(I_4, id) - I_{12}

$F \rightarrow id \cdot$

(I_6, T) - I_{13}

$E \rightarrow E+ \cdot T$

$T \rightarrow T \cdot * F$

(I_6, F) - I_{14}

$T \rightarrow F \cdot$

($I_6, ($) - I_{15}

$F \rightarrow (\cdot E)$

$E \rightarrow \cdot E+T$

$E \rightarrow \cdot T$

$T \rightarrow \cdot T * F$

$T \rightarrow \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow id$

(I_6, id) - I_{16}

$F \rightarrow id \cdot$

(I_7, F) - I_{17}

$T \rightarrow T * \cdot F$

($I_7, ($) - I_{18}

$F \rightarrow (\cdot E)$

$E \rightarrow \cdot E+T$

$E \rightarrow \cdot T$

$T \rightarrow \cdot T * F$

$T \rightarrow \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow \cdot id$

(I_7, id) - I_{19}

$F \rightarrow id \cdot$

($I_8,)$ - I_{20}

$F \rightarrow (E) \cdot$

($I_8, +$) - I_{21}

$E \rightarrow E+ \cdot T$

$T \rightarrow \cdot T * F$

$T \rightarrow \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow id$

$(I_7, *) - (I_7)$

$T \rightarrow T * \cdot F$

$F \rightarrow \cdot (E)$

$F \rightarrow \cdot id$

STATE	ACTION						GOTO		
	id	+	*	(	)	\$	E	T	F
0	S5			S4			1	2	3
1		S6				Acceptance			
2		r2	S7		r2	r2			
3		r4	r4		r4	r4			
4	S5			S4			8	2	3
5	S5	r6	r6		r6	r6			
6	S5			S4				9	3
7	S5			S4					10
8		S6			S11				
9		r1	S7		r1	r1			
10		r3	r3		r3	r3			
11		r5	r5		r5	r5			

1. $I_1: E \rightarrow E.$

Follow(E) = $\emptyset$

Follow(E) = $\{ \$ \}$

2. $I_2: E \rightarrow T.$

$E \rightarrow \cdot E$

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

3. $I_3: T \rightarrow F. (r_4)$

1 $E \rightarrow E \cdot T$

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

4. $I_5: F \rightarrow id.$

2 $E \rightarrow T (r_2)$

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

5. $I_9: E \rightarrow E \cdot T.$

3 $T \rightarrow T * F$

(2, \$) (2, +) (2,) = r_2

$I_{10}: T \rightarrow T * F.$

4 $T \rightarrow F (r_3)$

(3, \$) (3, +) (3,) (2, +) = r_3

$I_{11}: F \rightarrow (E).$

5 $F \rightarrow (E)$

(5, \$) (5, +) (5,) (5, *) = r_5

6 $F \rightarrow id$

$I_1: S \xrightarrow{1} S.$

$$w = id * id + id$$

Stack	Input	Action
\$0	id*id+id\$	shift 5
\$0id5	*id+id\$	reduce 6 $F \rightarrow id$
\$0F3	*id+id\$	reduce 4 $T \rightarrow F$
\$0T2	*id+id\$	S7
\$0T2*7	id+id\$	S5
\$0T2*7id5	+id\$	r6 $F \rightarrow id$
\$0T2*7F10	+id\$	r3 $T \rightarrow T * F$
\$0T2	+id\$	r2 $E \rightarrow T$
\$0E1	+id\$	S6
\$0E1+6	id\$	S5 $F \rightarrow id$
\$0E1+6id5	\$	r6 $F \rightarrow id$
\$0E1+6F3	\$	r4 $T \rightarrow F$
\$0E1+6T9	\$	r1 $E \rightarrow E + T$
\$0E1	\$	Acceptance

eg: 2 $S \rightarrow AS/b$
 $A \rightarrow SA/a$ $w = 'abab'$

→ shift → while performing shifting first we have to shift
 1/p symbol on to the top of the stack & then shift the state
 Number.

→ reduce → while performing reduce operation. (if $F \rightarrow id$ we
 reducing), if Right hand side contains one symbol we need to
 pop two symbols ^{from} top of the stack.

(ii) LR(0) parser:-

steps for construct LR(0) parser

1) Introduce Augmented grammars

2) calculate canonical collection of LR(0) items.

3) construct LR(0) parsing table by using (i) Goto (ii) Action functions.

eg:- $S \rightarrow AA$
 $A \rightarrow aA/b$

w = string = aabbb

Closure:-

LR(0) items - (I₀)

$S' \rightarrow \cdot S$

$S \rightarrow \cdot AA$

$A \rightarrow \cdot aA/b$

Goto
 (I₀, S) - (I₁)

$S' \rightarrow S \cdot$

Augmented grammars

$S' \rightarrow S$

$S \rightarrow AA$

Goto (I₀, A) - (I₂)

$S \rightarrow A \cdot A$

$A \rightarrow \cdot aA/b$

Goto (I₀, a) - (I₃) (I₄)

$A \rightarrow a \cdot A$

$A \rightarrow \cdot aA/b$
 $A \rightarrow b \cdot$

(I₂, A) - (I₅)

$S \rightarrow AA \cdot$

(I₂, a) - (I₃)

$A \rightarrow a \cdot A$

$A \rightarrow \cdot aA/b$

(I₂, b) - (I₄) I₃

~~$A \rightarrow aA$~~

$A \rightarrow b \cdot$

(I₃, a) - I₃

$A \rightarrow a \cdot A$

$A \rightarrow \cdot aA/b$

(I₃, A) - (I₆)

$A \rightarrow aA \cdot$

(I₃, b) - (I₄)

$A \rightarrow b \cdot$

LR(0) parsing Table

	Action			Goto	
	a	b	\$	A	S
0	S ₃	S ₄		2	1
1			Accept		
2	S ₃	S ₄		5	
3	S ₃	S ₄		6	
4	r ₃	r ₃	r ₃		
5	r ₁	r ₁	r ₁		
6	r ₂	r ₂	r ₂		

wrote down the states that contains

"." at the end of the production.

I₁ → S' → S.
 I₄ → A → b.
 I₆ → A → aA.
 I₅ → S → AA.

Given Grammar

$S \rightarrow AA$ - (1)
 $A \rightarrow aA$ - (2)
 $A \rightarrow b$ - (3)

→ The given string $w = aabb$ $\Rightarrow w = aabb\$$

Stack	Input	Action
\$0	aabb\$	shift 3
\$0a3	abb\$	shift 3
\$0a3a3	bb\$	shift 4
\$0a3a3b4	b\$	reduce r3
\$0a3a3	b\$	A → b
\$0a3a3b4	\$	shift 4

while performing shifting if we have to shift i/p symbol on to the top of the stack & then state Number

Stack	Input	Action
\$0	aabb\$	S3
\$0a3	abb\$	S3
\$0a3a3	bb\$	S4
\$0a3a3b4	b\$	reduce r3 A → b
\$0a3a3A6	b\$	r2 A → aA
\$0a3A6	b\$	r2 A → aA
\$0A2	b\$	S4
\$0A2b4	\$	r3 A → b
\$0A2A5	\$	r1 S → AA
\$0S1	\$	Accepted.

String $w = aabb$ is parsed through the grammar by using LR(0) parser.

Canonical LR parsing:-

Step 1: Augmented grammars

Eg: $S \rightarrow cc$
 $C \rightarrow ac$
 $C \rightarrow d$

$S \rightarrow s$
 $S \rightarrow cc$
 $C \rightarrow ac$
 $C \rightarrow d$

$w = add$

Step 2: LR(1) items = LR(0) + lookahead

$(I_0, \$) - (I_1)$ $S \rightarrow \cdot s, \$$ $S \rightarrow \cdot cc, \$$ $C \rightarrow \cdot ac, ald$ $C \rightarrow \cdot d, ald$	$(I_2, d) - (I_3)$ $c \rightarrow d \cdot, \$$
$(I_0, s) - (I_1)$ $s \rightarrow s \cdot, \$$	$(I_3, c) - (I_8)$ $C \rightarrow ac \cdot, ald$
$(I_0, c) - (I_2)$ $S \rightarrow c \cdot c, \$$ $C \rightarrow \cdot ac, \$$ $C \rightarrow \cdot d, \$$	$(I_3, a) - (I_3)$ $c \rightarrow a \cdot c, ald$ $C \rightarrow \cdot ac, ald$ $C \rightarrow \cdot d, ald$
$(I_0, a) - (I_3)$ $C \rightarrow a \cdot c, ald$ $C \rightarrow \cdot ac, ald$ $C \rightarrow \cdot d, ald$	$(I_3, d) - (I_4)$ $C \rightarrow d \cdot, ald$
$(I_0, d) - (I_4)$ $C \rightarrow d \cdot, ald$	$(I_6, c) - (I_9)$ $S \rightarrow ac \cdot, \$$
$(I_2, c) - (I_5)$ $S \rightarrow cc \cdot, \$$	$(I_6, a) - (I_6)$ $C \rightarrow a \cdot c, \$$ $C \rightarrow \cdot ac, \$$ $C \rightarrow \cdot d, \$$
$(I_2, a) - (I_6)$ $S \rightarrow a \cdot c, \$$ $C \rightarrow \cdot ac, \$$ $C \rightarrow \cdot d, \$$	$(I_6, d) - (I_7)$ $C \rightarrow d \cdot, \$$

Step 3) = construction of CLR parsing Table.

	Action			Goto		
	a	d	\$	S	C	
0	S ₃	S ₄		1	2	
1			Accepted			
2	S ₆	S ₇			5	1: S → CC
3	S ₃	S ₄			8	2: C → ac
4	C → d r ₃	C → d r ₃				3: C → d
5			S → CC r ₁			I ₁ : S' → S. \$
6	S ₆	S ₇				I ₄ : C → d, a, d
7			C → d r ₃		9	I ₅ : S → CC, \$
8	C → ac r ₂	C → ac r ₂				I ₇ : C → d, \$
9			C → ac r ₂			I ₈ : C → ac, a, d
						I ₉ : S → ac, \$

LALR Table: = I₃ = I₆
I₄ = I₇
I₈ = I₉

	Action			Goto	
	a	d	\$	S	C
0	S ₃₆	S ₄₇		1	2
1			ACC ACC		
2	S ₃₆	S ₄₇			5
36	S ₃₆	S ₄₇			89
47	r ₃	r ₃	r ₃		
5			r ₁		
89	r ₂	r ₂	r ₂		

LALR parsing table



→ On LALR parser combine the same states that are having different lookahead symbols.

w = add

Stack	Input buffer	Action
\$0	add \$	S ₃₆
\$0a36	ad \$	S ₄₇
\$0a36d47	d \$	r ₃ C → d
\$0a36c89	d \$	r ₂ C → ac
\$0c2	d \$	S ₇
\$0c2d47	\$	r ₃ C → d
\$0c2d47	\$	
\$0c2c5	\$	r ₁ S → CC
\$0s1	\$	Accepted

Canonical LR parser = (CLR parser)

15 16

Step 1: Augmented grammar

$S \rightarrow CC$

$C \rightarrow cC$

$C \rightarrow d$

$S \rightarrow S$

$S \rightarrow CC$

$C \rightarrow cC$

$C \rightarrow d$

Step 2: Calculating LR(0) items

$I_0 \rightarrow S \rightarrow \cdot S, \$$

$S \rightarrow \cdot CC, \$$

$C \rightarrow \cdot cC, c/d$

$C \rightarrow \cdot d, c/d$

(I_0, S)

$(I_0, S) - I_1$

$S \rightarrow S, \$$

$(I_0, C) - I_2$

$S \rightarrow C.C, \$$

$C \rightarrow c.C, \$$

$C \rightarrow \cdot d, \$$

$(I_0, e) - I_3$

$C \rightarrow c.C, c/d$

$C \rightarrow \cdot cC, c/d$

$C \rightarrow \cdot d, c/d$

$(I_0, d) - I_4$

$C \rightarrow d, c/d$

$(I_2, C) - I_5$

$S \rightarrow CC, \$$

$(I_2, C) - I_6$

$C \rightarrow c.C, \$$

$C \rightarrow \cdot cC, \$$

$C \rightarrow \cdot d, \$$

$(I_2, d) - I_7$

$C \rightarrow d, \$$

$(I_3, C) - I_8$

$C \rightarrow CC, c/d$

$(I_3, c) - I_9$

$C \rightarrow c.C, c/d$

$C \rightarrow \cdot cC, c/d$

$C \rightarrow \cdot d, c/d$

$(I_3, d) - I_{10}$

$C \rightarrow d, c/d$

$(I_6, C) - I_{11}$

$C \rightarrow CC, \$$

$(I_6, c) - I_{12}$

$C \rightarrow c.C, \$$

$C \rightarrow \cdot cC, \$$

$C \rightarrow \cdot d, \$$

$(I_6, d) - I_{13}$

$C \rightarrow d, \$$

construction of CLR parsing table =

Action

Goto

	c	d	\$	S	C
0	S ₃	S ₄		1	2
1			Accepted		
2	S ₆	S ₇		5	
3	S ₃	S ₄		8	
4	r ₃	r ₃			
5			r ₁		
6	S ₆	S ₇		9	
7			r ₃		
8	r ₁	r ₁			

$S \rightarrow S$
 $S \rightarrow CC \rightarrow 1$
 $C \rightarrow cC \rightarrow 2$
 $C \rightarrow d \rightarrow 3$
 $I_1: S \rightarrow S, \$$
 $I_4: C \rightarrow d, c/d$
 $I_5: CC, \$$
 $I_7: C \rightarrow d, \$$
 $I_8: C \rightarrow CC, c/d$
 $I_9: C \rightarrow cC, \$$

parse the input string $w = add\$$

Stack	Input	String
\$0	dd\$	shift 4
\$0d1	d\$	reduce 3 $C \rightarrow d$
\$0C2	d\$	shift 7
\$0C2d7	\$	reduce 3 $C \rightarrow d$
\$0C2C5	\$	reduce 8, $S \rightarrow CC$
\$0S1	\$	Accepted

↳ So, the given string is accepted by the grammar.

↳

LALR parser = (Look Ahead LR parser)

eg 1 = steps to construct LALR parser

- (1) Introduce Augmented grammar
- (2) calculate the LR(1) items
- (3) construction of LALR parse table.
- (4) parsing of given string.

eg 1 = Construct LALR parser for

$S \rightarrow CC$

$C \rightarrow ac$

$C \rightarrow d$

and also parse the string $w = add$.

Parser generator:-

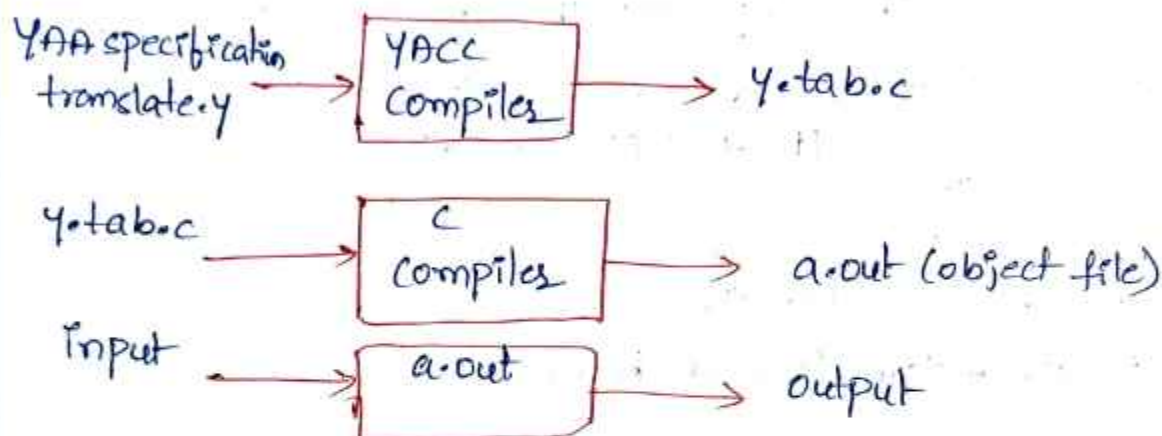
50 (18)

→ we use LALR parser generator YACC.

→ YACC → "yet another compiler compiler"

→ YACC is a tool to generate parser, YACC accepts an tokens as i/p and produces parse tree as o/p.

The parser generator YACC:-



→ translate.y consists of YACC specification

Structure of YACC program:-

it has three parts.

declarations

/*

translation Rules

*/

Auxiliary functions

Declarations:-

→ used to declare the C variables and constants, & header files are also specified here.

Syntax:-
%d
%}

eg:- %d

int a, b;

const int a=20;

#include <stdio.h>

%}

(i) Translation Rules:-

→ Translation Rules are enclosed between `"%{"` & `"%}"`

Syntax:-

$\text{head} \rightarrow \text{body}_1 / \text{body}_2 / \text{body}_3 \dots$ eg: $C \rightarrow aa / bb$

$\text{head} : \text{body}_1 \{ \text{semantic action} \}$

$\quad \quad \quad \{ \text{body}_2 \{ \text{semantic action} \}$

$\quad \quad \quad \quad \quad \quad \quad \{ \text{body}_3 \{ \text{semantic action} \}$

→ In translation Rules we use two symbols $\boxed{\$i}$, $\$i \rightarrow$

$\$i \rightarrow$ represent LHS attribute value, to access LHS non-terminal value we use $\$i$

(ii) Auxiliary function:- $\$i \rightarrow$ represent the i th symbol of body.

eg: $E \rightarrow E + T$ (if we want to access E , we have to use $\$1$, $+ \rightarrow \$2$, $T \rightarrow \$3$)

(iii) Auxiliary function:-

→ used to define "C" function.

→ yacc is ~~use~~ function is used here.

eg: YACC specification (program) of simple desk calculator.

EX:- $E \rightarrow E + E / T$

$T \rightarrow T * F / F$

$F \rightarrow (E) / id.$

Declaration part:-

`%{`

`#include <ctype.h>`

`%}`

`% token digit` → specification of RE

`%{"`

→ it specifies r/p to desk calculator is an expression by in

`line: expr '\n' { printf("%d\n", $1); }`

`expr: expr '+' term { $$ = $1 + $3; }`

`term`

`term: term '*' factor { $$ = $1 * $3; }`

`factor`

(2)

D

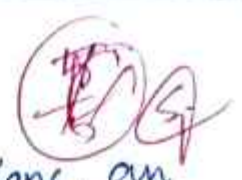
factos : '(' 'expr')' { \$\$ = \$2; }
| DIGIT;

*/

yylex()

```
{  
    int c;  
    c = getchar();  
    if (isdigit(c))  
    {  
        yylval = c - '0';  
        return DIGIT;  
    }  
    return c;  
}
```

yylex → is a lexical analyzer function
which takes our source program
as input and produces tokens as o/p



1) Using Ambiguous Grammars

- For language constructs like expressions, an ambiguous grammar provides a shorter, more natural specification than any equivalent unambiguous grammar.
- Another use of ambiguous grammar is in isolating commonly occurring syntactic constructs for special-case optimization, we can add new productions to grammar.
- Sometimes ambiguity rules allow only one parse tree, in such case it is unambiguous, it is possible to design an LR parser that allows same ambiguity resolving choices.

Precedence & Associativity to Resolve Conflicts:-

- Consider ambiguous grammar with operators '+' & '*'

$$E \rightarrow E + E \mid E * E \mid (E) \mid id \quad (1)$$
- $E \rightarrow E + T$, $T \rightarrow T * F$, generates same language gives lower precedence to '+' than *, makes left associative. [use ambiguous grammar]
- ~~But~~ First we change associativity and precedence of operators + & * without disturbing above grammar.
- Second, the parser for the unambiguous grammar will spend a substantial fraction of its time reducing by the productions $E \rightarrow T$ & $T \rightarrow F$
- $E' \rightarrow E$, (1) is ambiguous there will be parsing-action conflicts when we try to produce LR parsing table.

$I_0: E' \rightarrow \cdot E$
 $E \rightarrow \cdot E + E$
 $E \rightarrow \cdot E * E$
 $E \rightarrow \cdot (E)$
 $E \rightarrow \cdot id$

$I_1: E' \rightarrow E \cdot$
 $E \rightarrow E \cdot + E$
 $E \rightarrow E \cdot * E$

$I_2: E \rightarrow (\cdot E)$
 $E \rightarrow \cdot E + E$
 $E \rightarrow \cdot E * E$
 $E \rightarrow \cdot (E)$
 $E \rightarrow \cdot id$

$I_3: E \rightarrow E \cdot id$

$I_4: E \rightarrow E + \cdot E$
 $E \rightarrow \cdot E + E$
 $E \rightarrow \cdot E * E$
 $E \rightarrow \cdot (E)$
 $E \rightarrow \cdot id$

$I_5: E \rightarrow E * \cdot E$
 $E \rightarrow \cdot E + E$
 $E \rightarrow \cdot E * E$
 $E \rightarrow \cdot (E)$
 $E \rightarrow \cdot id$

$I_6: E \rightarrow (E \cdot)$
 $E \rightarrow E \cdot + E$
 $E \rightarrow E \cdot * E$

$I_7: E \rightarrow E + E \cdot$
 $E \rightarrow E \cdot + E$
 $E \rightarrow E \cdot * E$

$I_8: E \rightarrow E * E \cdot$
 $E \rightarrow E \cdot + E$
 $E \rightarrow E \cdot * E$

$I_9: E \rightarrow (E) \cdot$

LR(0) items of augmented exp grammar

Conflict occur in I_7 & I_8 $E \rightarrow E + E$, $E \rightarrow E * E$

PREFIX

STACK

INPUT

$E + E$

0147

$* id \$$

If ilp is $id + id$.

State	ACTION						GOTO
	id	+	*	(	)	\$	
0	s ₃			s ₂			1
1		s ₄	s ₅			accept	
2	s ₃			s ₂			6
3		r ₄	r ₄		r ₄	r ₄	
4	s ₃			s ₂			7
5	s ₃			s ₂			8
6		s ₄	s ₅		s ₉		
7		r ₁	s ₅		r ₁	r ₁	
8		r ₂	r ₂		r ₂	r ₂	
9		r ₃	r ₃		r ₃	r ₃	

Fig - parsing table for grammar

"Dangling-Else" Ambiguity:-

stmt $\rightarrow$ if expr then stmt else stmt
if expr then stmt
other

Consider

The grammar, an abstraction of this grammar where 'i' stands for if expr then, e stands for else and 'a' stands for "all other production".

$S' \rightarrow S$ (augmented grammar)

$S \rightarrow ises / is / a$

$I_0: S' \rightarrow \cdot S$

$S \rightarrow \cdot ises$

$S \rightarrow \cdot is$

$S \rightarrow \cdot a$

$I_1: S' \rightarrow S \cdot$

$I_2: S \rightarrow i \cdot ses$

$S \rightarrow i \cdot s$

$S \rightarrow \cdot ises$

$S \rightarrow \cdot is$

$S \rightarrow \cdot a$

$I_3: S \rightarrow a \cdot$

$I_4: S \rightarrow is \cdot es$

$I_5: S \rightarrow ise \cdot s$

$S \rightarrow \cdot ises$

$S \rightarrow \cdot is$

$S \rightarrow \cdot a$

$I_6: S \rightarrow ises \cdot$

Fig:- LR(0) states for augmented grammar

if expr then stmt — (2)

Ambiguity arises in I_4 shift/reduce conflict

$S \rightarrow is \cdot es$ calls for a shift of e

$FOLLOW(S) = \{e, \$ \}$, $S \rightarrow is \cdot$ call for reduction by ~~sexxis~~

$S \rightarrow is$ on ifp 'e'

In (2) should we shift else onto stack or reduce if expr then ~~stmt~~, we should shift else because it is associated with previous then.

SLR parsing table is Constructed.

STATE	ACTION				GOTO
	i	e	a	\$	
0	S2		S3	accepted	1
1					
2	S2		S3		4
3		r3			
4		S5		r2	
5	S2		S3		6
6		r1		r1	

input is $iiaea$, At line (5) state 4 selects the shift action on input e , whereas at line (9) state 4 calls for reduction by $s \rightarrow is$ on input $\$$.

STACK	SYMBOLS	INPUT	ACTION
(1) 0		$iiaea \$$	shift
(2) 02	i	$iaea \$$	shift
(3) 022	ii	$aea \$$	shift
(4) 0223	iia	$ea \$$	shift
(5) 0224	ii's	$ea \$$	reduce by $s \rightarrow a$
(6) 02245	iiSe	$a \$$	shift
(7) 022453	iisea	$\$$	reduce by $s \rightarrow a$
(8) 022456	iises	$\$$	reduce by $s \rightarrow ises$
(9) 024	is	$\$$	reduce by $s \rightarrow is$
(10) 01	S	$\$$	accept

Fig:- parsing actions on input $iiaea$

Syntax Directed Definition (SDD):

- Semantic Analysis phase & ICG phase uses SDD's.
- SDD is Context Free Grammar with Semantic Rules and attributes

$$\boxed{\text{SDD} = \text{CFG} + \text{Semantic Rules}}$$

- SDD is also called as "attribute grammar".
- Attributes are associated with grammar symbols and semantic Rules are associated with productions.
- If 'x' is a symbol and 'a' is one of attribute then x.a denotes value at node 'x'.
- Attributes may be numbers, string, data types & references etc

eg

production

Semantic Rules.

$$E \rightarrow E + T$$

$$E.val = E.val + T.val \quad (\text{Here } E \rightarrow \text{ is grammar symbol})$$

$$E \rightarrow T$$

$$E.val = T.val \quad \text{val} \rightarrow \text{is attribute of } E$$

$$D \rightarrow TL$$

$$L.inh = T.type$$

→ Semantic Rules have two Notations. (i) SDD (ii) SDT

Types of attribute:-

(i) Synthesized Attribute:- If a node takes value from its children node then it is called as synthesized attribute.

Eg:- $A \rightarrow BCD$ $A.s = B.s$ $A.s = C.s$ $A.s = D.s$ } $A \rightarrow$ is synthesized attribute.
 parent node takes value from child nodes
 $A \rightarrow$ be a parent node
 $B, C, D \rightarrow$ are children node.

(ii) Inherited Attributes:- If a node takes value from its sibling or Parent then it is called as inherited attributes.

$$D \rightarrow TL$$

production

Semantic Rules

Eg

Eg:- $A \rightarrow BCD$ $\rightarrow C$ is inherited attribute

$$C.i = A.i \rightarrow \text{Here } C \text{ is taking value from its parent } A.$$

$$C.i = B.i \rightarrow C \text{ is taking value from its sibling } B.$$

$$C.i = D.i \rightarrow C \text{ is taking value from its sibling } D.$$

Eqn =	production	Semantic Rules
	$L \rightarrow E_n$	$E.val = E.val$
	$E \rightarrow E_1 + T$	$E.val = E_1.val + T.val$
	$E \rightarrow T$	$E.val = T.val$
	$T \rightarrow T_1 * F$	$T.val = T_1.val * F.val$
	$T \rightarrow F$	$T.val = F.val$
	$F \rightarrow (E)$	$F.val = E.val$
	$F \rightarrow \text{digit}$	$F.val = \text{digit.lexval}$

Fig: SPP for a simple desk calculator

Evaluating an SPP at Nodes of parse tree:-

Annotated parse tree:-

- A parse tree which contains values at each node is known as annotated parse tree.
- A parse tree is constructed in order to evaluate the attribute value at each node of parse tree.
- If an attribute is synthesized.
 - 1st evaluate the val attribute at all the children node.
 - evaluate the value attribute at parent node.
 - synthesized attributes, attributes are evaluated in bottom-up manner.

production

$A \rightarrow B$

Semantic Rules

$A.s = B.i$

$B.i = A.s + 1$

These rules are circular, it is impossible to evaluate A.s without first evaluating B.i at some node.

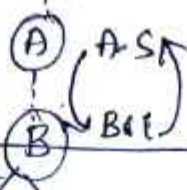
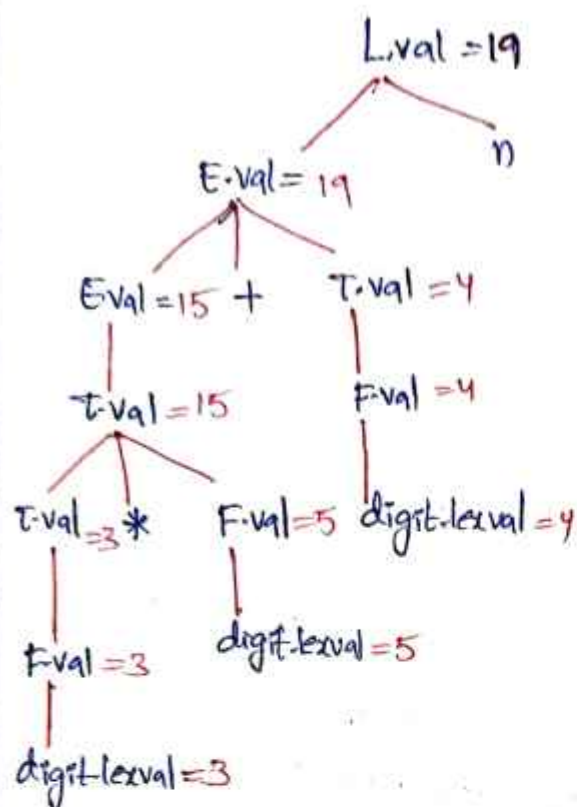


Fig: The circular dependency of A.s and B.i on one another

Eg1: Construct an SDD for simple arithmetic calculator grammars and construct parse tree for $3*5+4n$.



production	Semantic Rule
$L \rightarrow En$	$L.val = E.val$
$E \rightarrow E+T$	$E.val = E.val + T.val$
$E \rightarrow T$	$E.val = T.val$
$T \rightarrow T_1 * F$	$T.val = T_1.val * F.val$
$T \rightarrow F$	$T.val = F.val$
$F \rightarrow (E)$	$F.val = E.val$
$F \rightarrow digit$	$F.val = digit.lexval$

Fig: SDD for simple arithmetic calculator

Eg2: construct Annotated parse tree for $2+3*4$.

Fig: Annotated parse tree for $3*5+4n$

→ Annotated parse tree contains values at each node.

→ To construct Annotated parse tree, we have to perform top-down left to Right traversing, if there is reduction execute corresponding action.

Eg2: production Semantic Rules

$D \rightarrow TL$

$L.inh = T.type \rightarrow$ list in herit is type T.

parse tree =

$T \rightarrow int$

$T.type = "integer"$

$T \rightarrow real$

$T.type = "real"$

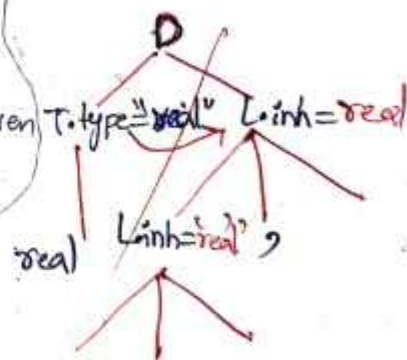
$L \rightarrow L, id$

$L.inh = L.inh \rightarrow$ here id needs to be given $T.type$.

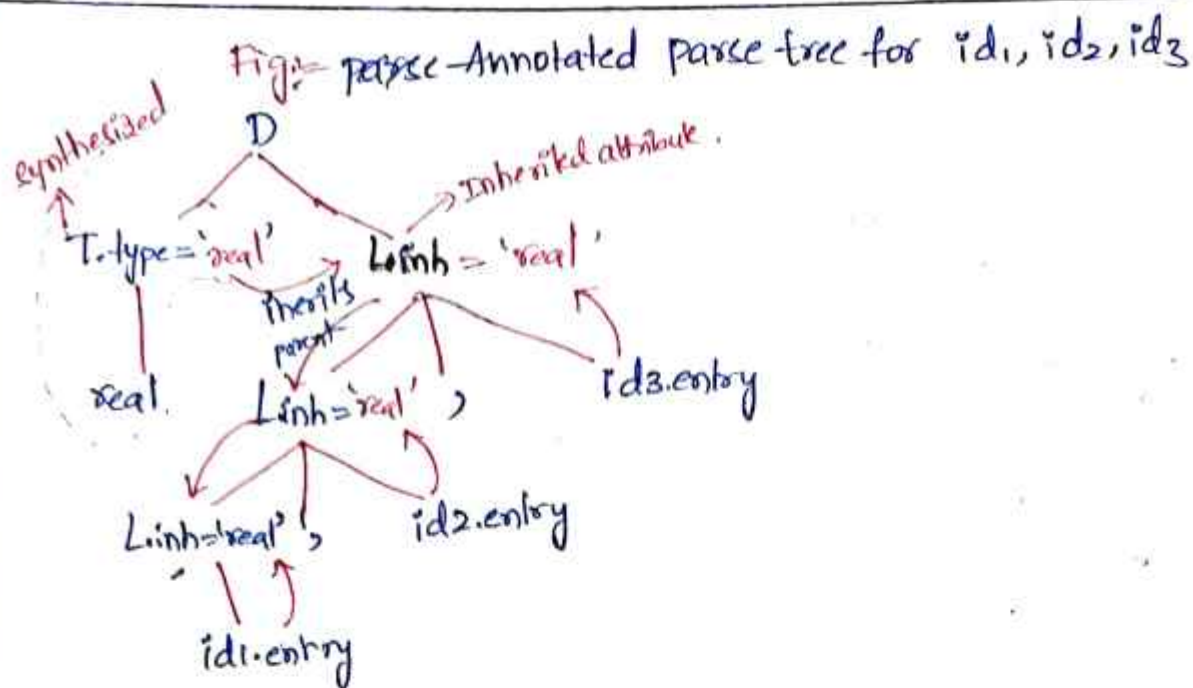
$addtype(id.entry, L.inh)$

$L \rightarrow id$

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



→ In dependency graph we have an edges directed edges



② Evaluating order SDD's :-

→ Dependency graph is used to determine an evaluation order for the attribute given

Dependency Graph :- $\rightarrow$ A directed graph that represents the interdependency b/w synthesized and inherited attribute at node in the parse tree

→ Dependency Graph represents the flow of information among the attributes in parse tree.

→ used to determine the evaluation order for attribute in a parse tree. (which semantic action should execute first)

→ An Annotated parse tree shows the value of attributes, a dependency graph determines how those values can be computed.

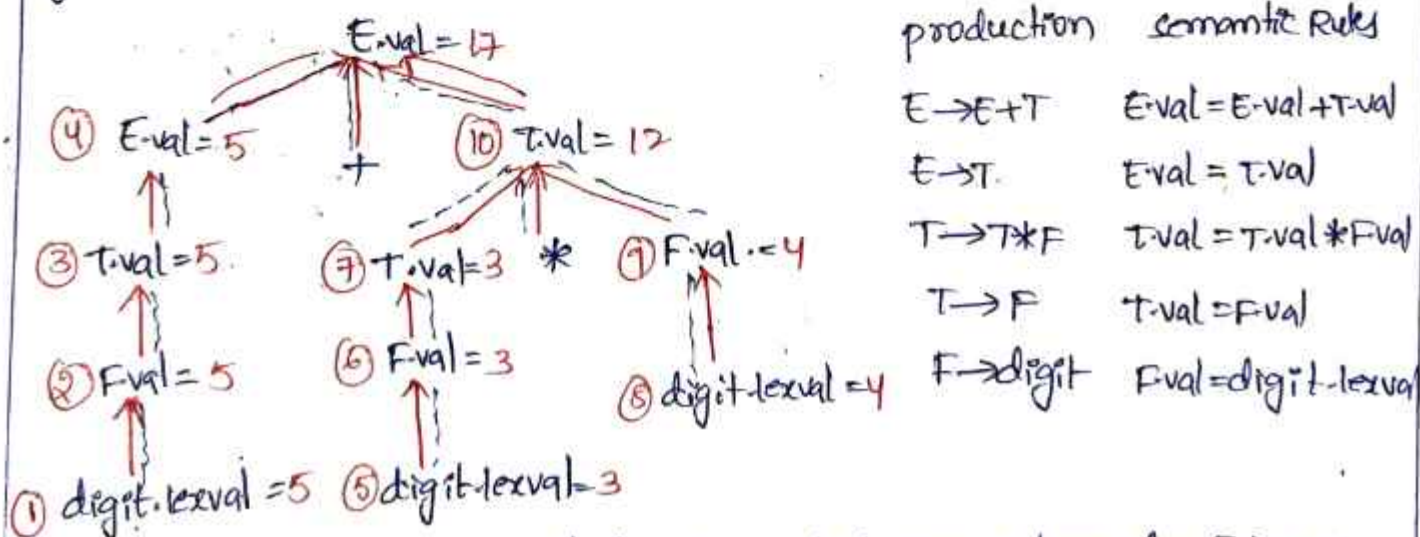


Fig. = Dependency graph for Annotated parse tree for $5 + 3 * 4$

→ It is also called as postfix "SDD"

→ Attributes are evaluated with Bottom-up parsing



L-Attributed SDD:

* A SDD that use both synthesis attributes & Inherited attributes is called as L-Attributed SDD.

* In L-Attributed SDD, Inherited attribute is restricted to inherit from parent & left sibling only.

Eg: $A \rightarrow XYZ \{ y.s = A.s, y.s = x.s, y.s = z.x \}$ production semantic action

* semantic actions are placed anywhere on R.H.S. Eg: $E \rightarrow E+T \{ E \rightarrow E+T \}$

* evaluated attributes are evaluated by traversing parse tree depthfirst, left to right order

Production

$D \rightarrow TL$

$T \rightarrow \text{int}$

$T \rightarrow \text{float}$

$L \rightarrow L, id$

$L \rightarrow id$

Semantic Rules

$L.inh = T.type$

$T.type = \text{integer}$

$T.type = \text{float}$

$L.inh = L.inh.addtype(id.entry, L.inh)$

$L.inh = addtype(id.entry, L.inh)$

Eg: SDD for simple type declaration for L-Attributed SDD

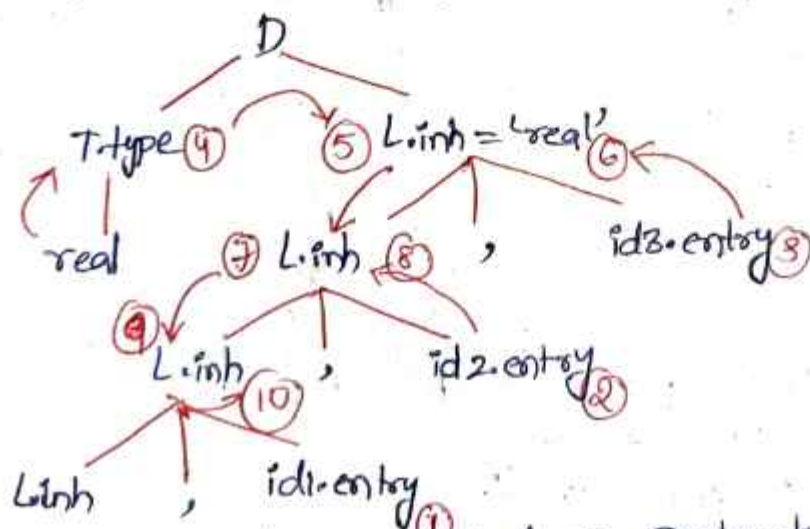


Fig: Dependency graph for a Declaration Ahead id_1, id_2, id_3

(3) Application of syntax directed Translation:-

Syntax Directed Translation (SDT):-

↳ SDT is a CFG together with semantic actions.

↳ Semantic actions are enclosed in '{', '}' eg: $A \rightarrow ABC \rightarrow ABC$
 $\rightarrow A \{ \} BC$
 $\rightarrow \{ \} ABC$

↳ Semantic actions are placed at any where on RHS of productions

↳ Semantic actions specifies in which order the expression is executed.

eg 1:- production

$A \rightarrow B + C$

semantic action

{ printf ('+') ; }

eg 2:-

production semantic action

$E \rightarrow E + T$ { printf ('+') ; } ①

$E \rightarrow T$ { } ②

$E \rightarrow T * F$ { printf ('*') ; } ③

$T \rightarrow F$ { } ④

$F \rightarrow \text{num}$ { printf (num.val) ; } ⑤

eg 2:-

production

semantic action

$E \rightarrow E + T$ { E.val = E.val + T.val ; } ①
 T { T.val = T.val ; }

$E \rightarrow T * F$ { T.val = T.val * F.val ; }
 F { F.val = F.val ; }

$T \rightarrow F$ { F.val = F.val ; }

2+3*4

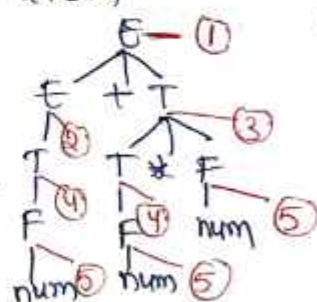


Fig:- SDT

Fig:- SDT for evaluation of Expression.

Applications of SDT:-

Construction of syntax Tree:-

→ syntax tree is an intermediate representation

→ The nodes in syntax tree is implemented by objects with suitable no. of fields

→ Each field will have an op-field that is the label of the node.

We can construct the syntax tree by using the following functions.

- (1) $\text{Mknode}(\text{op}, \text{left}, \text{right})$
- (2) $\text{mkleaf}(\text{id}, \text{entry to symbol table})$
- (3) $\text{mkleaf}(\text{num}, \text{value})$

→ If the node is leaf, an additional field hold the lexical value for the leaf

$\text{leaf}(\text{op}, \text{val})$ — create leaf object

→ If Node is an operator, create an object with first field: op and k additional for the k children $C_1 \dots C_k$

Ex = Construct the syntax tree for the following grammar for the expression $x * y - 5 + 3$.

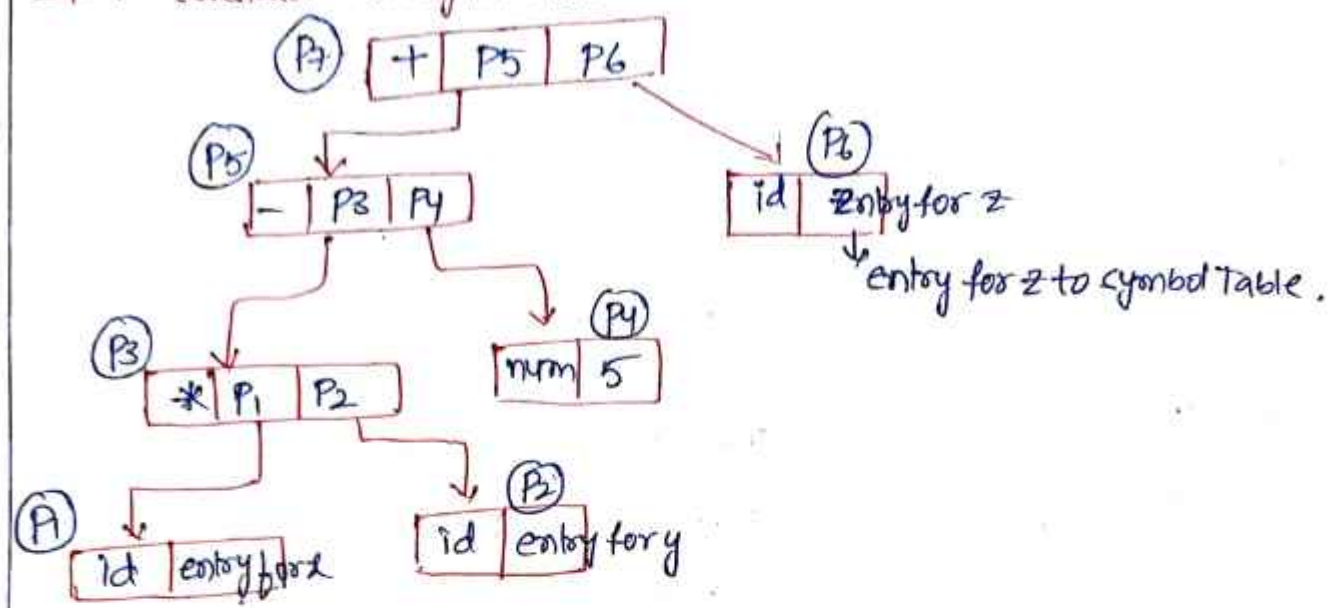
Step 1 = Construct SDD for the given grammar

production	Semantic Rule
$E \rightarrow E_1 + T$	$E.\text{node} = \text{mknode}('+', E_1.\text{node}, T.\text{node})$ (a) $E.\text{node} = \text{newleaf}('+', E.\text{node}, T.\text{node})$
$E \rightarrow E_1 - T$	$E.\text{node} = \text{mknode}('-', E_1.\text{node}, T.\text{node})$
$E \rightarrow T$	$E.\text{node} = T.\text{node}$
$T \rightarrow (E)$	$T.\text{node} = E.\text{node}$
$T \rightarrow \text{id}$	$T.\text{node} = \text{mkleaf}(\text{id}, \text{id}.\text{entry})$ (b) $T.\text{node} = \text{newleaf}(\text{id}, \text{id}.\text{entry})$
$T \rightarrow \text{num}$	$T.\text{node} = \text{mkleaf}(\text{num}, \text{num}.\text{val})$

<u>Step 2</u> :-	Symbol	operation
	x	$P_1 = \text{mkleaf}(\text{id}, \text{entry} - x)$
	y	$P_2 = \text{mkleaf}(\text{id}, \text{entry} - y)$
	*	$P_3 = \text{mknode}('*', P_1, P_2)$
	5	$P_4 = \text{mkleaf}(\text{num}, 5)$
	-	$P_5 = \text{mknode}('-', P_3, P_4)$

3 $P_6 = \text{mkleaf}(\text{id}, \text{entry}-3)$
 + $P_7 = \text{mnnode}('-', P_5, P_6)$

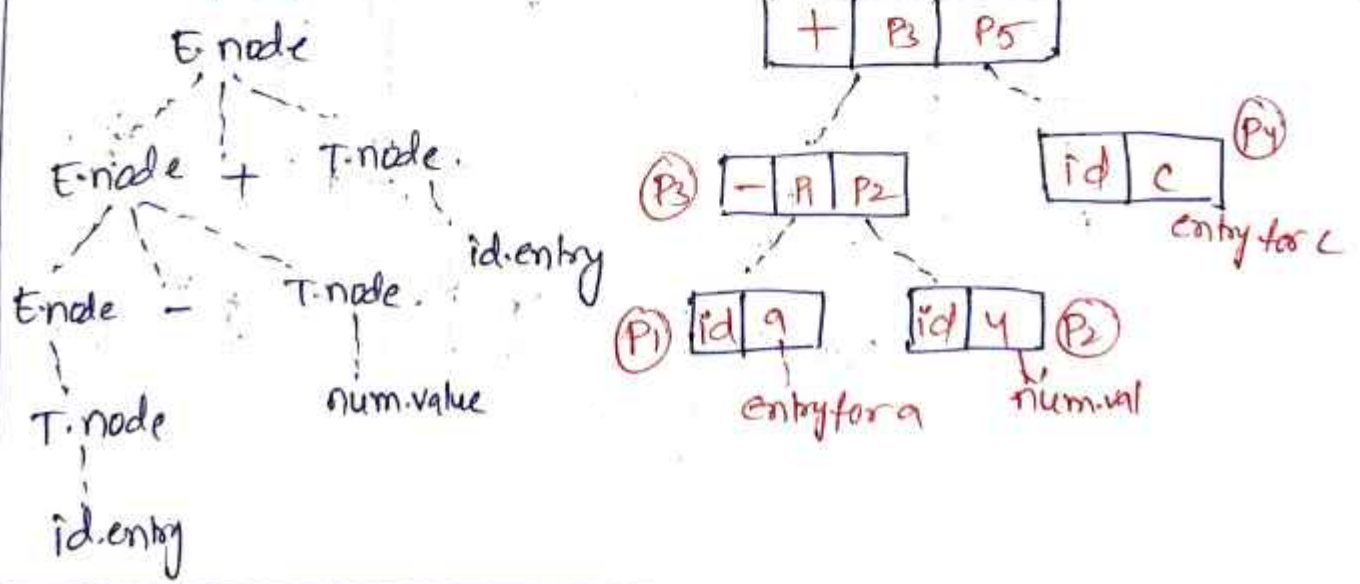
Step 3:- construct the syntax tree.



Eg 2:- $a-4+c$

symbol operation
 a $P_1 = \text{mkleaf}(\text{id}, \text{entry}-a)$
 4 $P_2 = \text{mkleaf}(\text{num}, 4)$
 - $P_3 = \text{mnnode}('-', P_1, P_2)$
 c $P_4 = \text{mkleaf}(\text{id}, \text{entry}-c)$
 + $P_5 = \text{mnnode}('+', P_3, P_4)$

Constructing syntax tree:-



Eg3 := production

Semantic Rules

$E \rightarrow TE'$

$E \cdot node = E \cdot syn$

$E \rightarrow +TE'$

$E \cdot inh = T \cdot node$

$E' \cdot inh = newnode(+', E' \cdot inh, T \cdot node)$

$E' \rightarrow -TE'$

$E' \cdot syn = E' \cdot syn$

$E' \cdot inh = newnode('-', E' \cdot inh, T \cdot node)$

$E' \rightarrow \epsilon$

$E' \cdot syn = E' \cdot syn$

$T \rightarrow (E)$

$E' \cdot syn = E' \cdot inh$

$T \rightarrow id$

$T \cdot node = E \cdot node$

$T \rightarrow num$

$T \cdot node = newleaf(id, id \cdot entry)$

$T \cdot node = newleaf(num, num \cdot val)$

Eg1 := $a - 4 + c$

symbol operation

a $P_1 = mkleaf(id, entry \cdot a)$

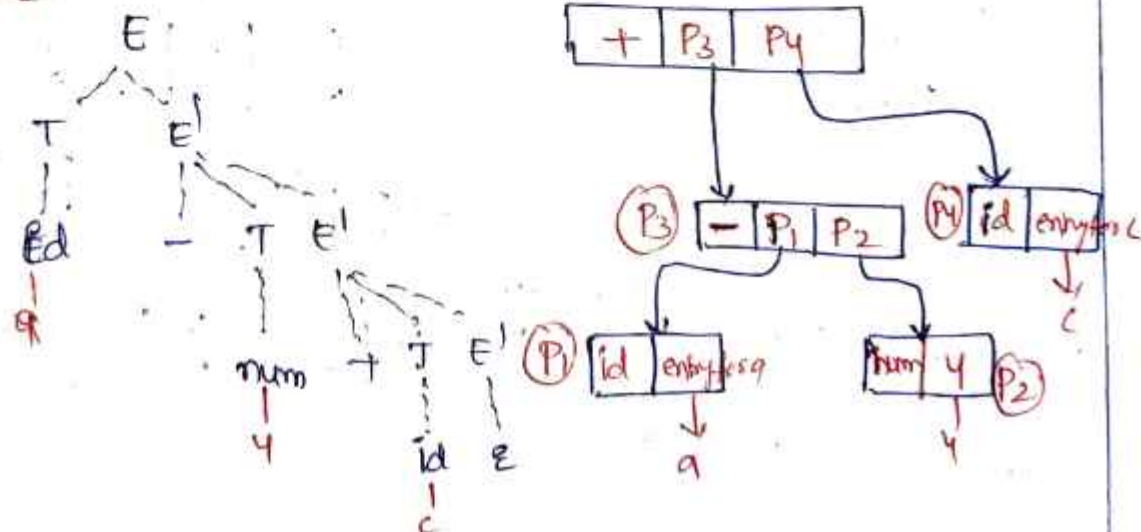
4 $P_2 = mkleaf(num, 4)$

$-$ $P_3 = mnode('-', P_1, P_2)$

c $P_4 = mkleaf(id, entry \cdot a)$

$+$ $P_5 = mnode(+', P_3, P_4)$

Syntax tree:



5) Syntax-Directed Translation Schemes

→ A SDT is a context free Grammar combined with semantic actions.

→ Semantic actions are also called as program-fragments.

→ Semantic actions are embedded with production body.

→ Any SDT can be implemented by constructing parse-tree and then performing the actions in a left-to-right depth-first order.

SDTs are used to implement two important classes of SDDs

1) The underlying grammar is LR-parable, & the SDD is S-attributed

2) The underlying grammar is LL-parable, and the SDD is L-attributed.

Postfix Translation Schemes

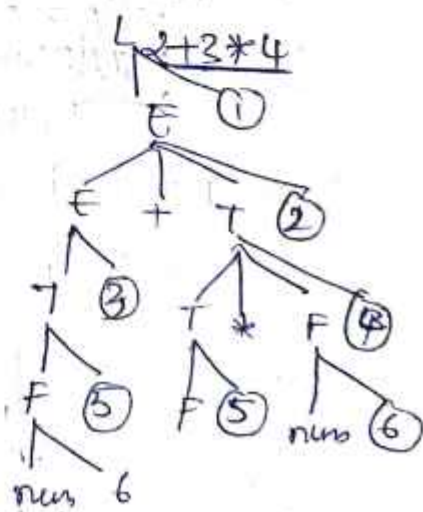
→ It is used to convert infix Expression to postfix Expression

→ Infix Expression - operator appears b/w operands

postfix Expression - operator appears after the operands.

Production	Semantic actions
$E \rightarrow E_n$	{ print (E.val); } ①
$E \rightarrow E + T$	{ print ('+'); } ②
$E \rightarrow T$	{ } ③
$E \rightarrow T * F$	{ print ('*'); } ④
$T \rightarrow F$	{ } ⑤
$F \rightarrow \text{num}$	{ print (num.val); } ⑥

Fig:- SDT for desc calculator



2nd method to convert the postfix infix expression to postfix expression

Eg:- $E \rightarrow E+T \mid T$ for $1+2+3$
 $T \rightarrow \text{num}$

=(here we are having only operation, so that we need check the left recursion in the grammar)

→ If Grammar contains left recursion we need to convert it to eliminate left recursion from the grammar

production	SAT
$E \rightarrow E+T$	$\{ \text{print}+('+'); \}$
$E \rightarrow T$	$\{ \}$
$T \rightarrow \text{num}$	$\{ \text{print num-value} \}$

SAT for given grammar

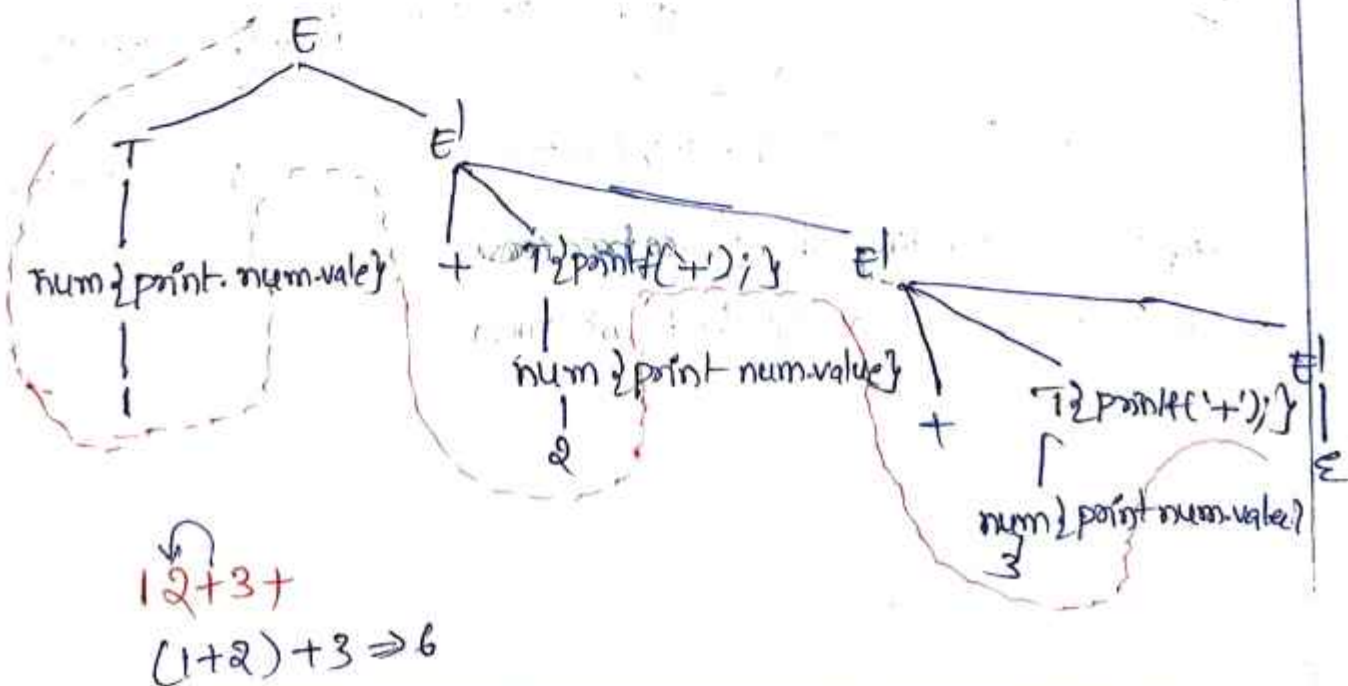
So, the given Grammar contains left recursion

$E \rightarrow E+T \mid T$ $E \rightarrow E+T \mid \{ \text{print}+('+'); \} \mid T$
 $T \rightarrow \text{num}$ $\alpha \quad \beta$

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

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

The grammar after eliminating the left recursion

$$\begin{cases} E \rightarrow TE' \\ E' \rightarrow +T \{ \text{print}+('+'); \} E' \mid \epsilon \\ T \rightarrow \text{num} \{ \text{print num-value} \} \end{cases}$$


Ex: = Give translation scheme that convert infix expr to postfix expression for the following grammar and also generate annotated parse for input string 2+6+1

Grammar $\rightarrow E \rightarrow E+T \{ \text{print}('+') \} / T$
 $T \rightarrow 0 | 1 | 2 | \dots | 9$ (it contains left recursion)

$A \rightarrow A\alpha | B$

$A \rightarrow BA$
 $A \rightarrow \alpha A | \epsilon$

$E \rightarrow TE'$
 $E' \rightarrow +T \{ \text{print}('+') \} E'$
 $E' \rightarrow \epsilon$

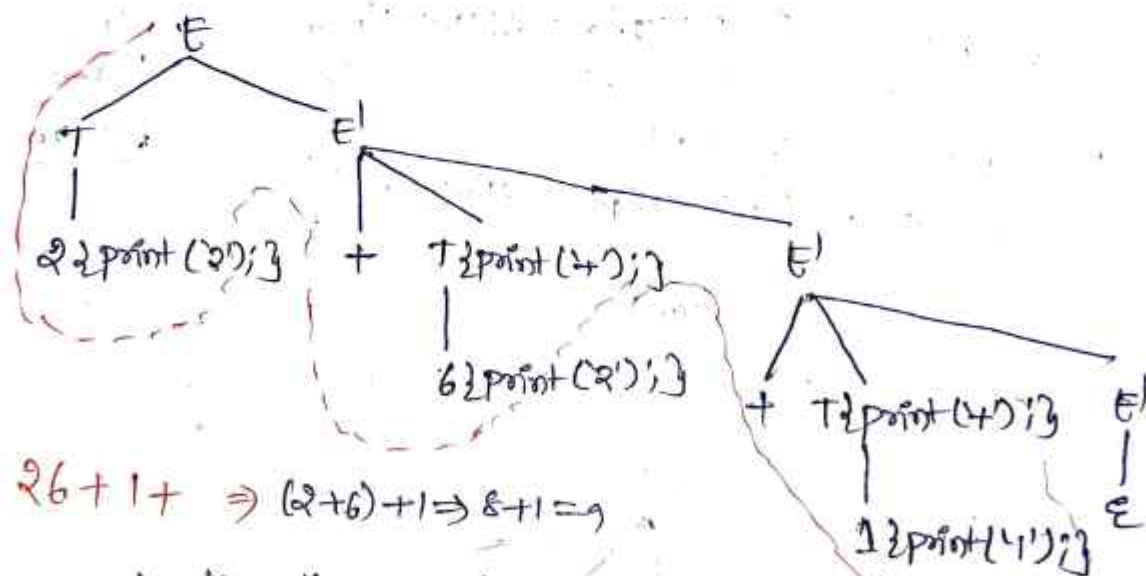
The Grammar after eliminating left recursion

$T \rightarrow 0 \{ \text{print}('0') \} \epsilon$

$T \rightarrow 1 \{ \text{print}('1') \} \epsilon$

$T \rightarrow 9 \{ \text{print}('9') \} \epsilon$

Fig: SDT to convert infix to postfix
 Annotated Parse tree: The given i/p string 2+6+1 $\Rightarrow$ 26+1+



$\rightarrow$ After constructing the parse tree, traverse the parse from top-down and left to right manner

$L \rightarrow E_n \quad \{ \text{print}(E.\text{val}); \}$

$E \rightarrow E + T \quad \{ E.\text{val} = E_1.\text{val} + T.\text{val}; \}$

$E \rightarrow T \quad \{ E.\text{val} = T.\text{val}; \}$

$T \rightarrow T_1 * F \quad \{ T.\text{val} = T_1.\text{val} * F.\text{val}; \}$

$T \rightarrow F \quad \{ T.\text{val} = F.\text{val}; \}$

$F \rightarrow (E) \quad \{ F.\text{val} = E.\text{val}; \}$

$F \rightarrow \text{digit} \quad \{ F.\text{val} = \text{digit}.\text{lexval}; \}$

Fig = postfix SDT implementing the desc calculator

parser - stack implementation of postfix SDT's =

→ postfix SDT's can be implemented during LR parsing by executing the actions when the reductions occur.

→ The grammar symbols of each grammar =

→ The attributes of each grammar symbols can be placed in stack during the parsing.

→ The parser stack contain records with fields for a grammar symbol.

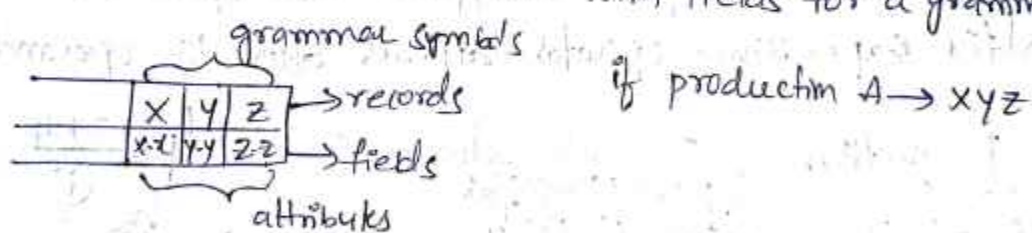


Fig = parser stack with a field for synthesized attributes

q1 $E \rightarrow E + T \quad \{ \text{print}(+); \}$ 2 $E \rightarrow E * T$ 2+3*4

T
+
E

↓
7

+
T
E

814

*
F
T

12

production	Actions.
$L \rightarrow E_n$	{ print (stack[top-1], val); top = top - 1
$E \rightarrow E + T$	{ stack[top+2].val = stack[top-2].val + stack[top].val; top = top - 2; }
$E \rightarrow T$	
$T \rightarrow T_1 * F$	{ stack[top-2].val = stack[top-2].val * stack[top].val; top = top - 2; }
$T \rightarrow F$	
$F \rightarrow (E)$	{ stack[top-2].val = stack[top-1].val; top = top - 2; }
$F \rightarrow \text{digit}$	

Fig:- Implementing the decal calculator on bottomup-parsing stack.

SOL's with Actions Inside productions:-

→ an actions may be placed at any position ~~on~~ right within the body of production.

eg:- $B \rightarrow x \{a\} y.$



The action "a" is executes after recognizing the ~~the~~ "x"

- If the parse is bottom-up then we perform action "a" when "x" is appears on the top of the stack.
- If the parse is top-down, we perform a just before expanding the "y"

Any SDT can be implemented as follows.

- 1) Ignoring the actions, parse the i/p & produce a parse tree as a result.
- 2) examine each node and add additional node for corresponding action
- 3) perform the preorder traversal of the tree, and as soon as a node is labeled by action is visited, perform that action.

Ex

Ex: parse tree for expression $3*5+4$ with actions inserted
 we get if we visit the nodes in preorder, we get the prefix form of expression $+*354$

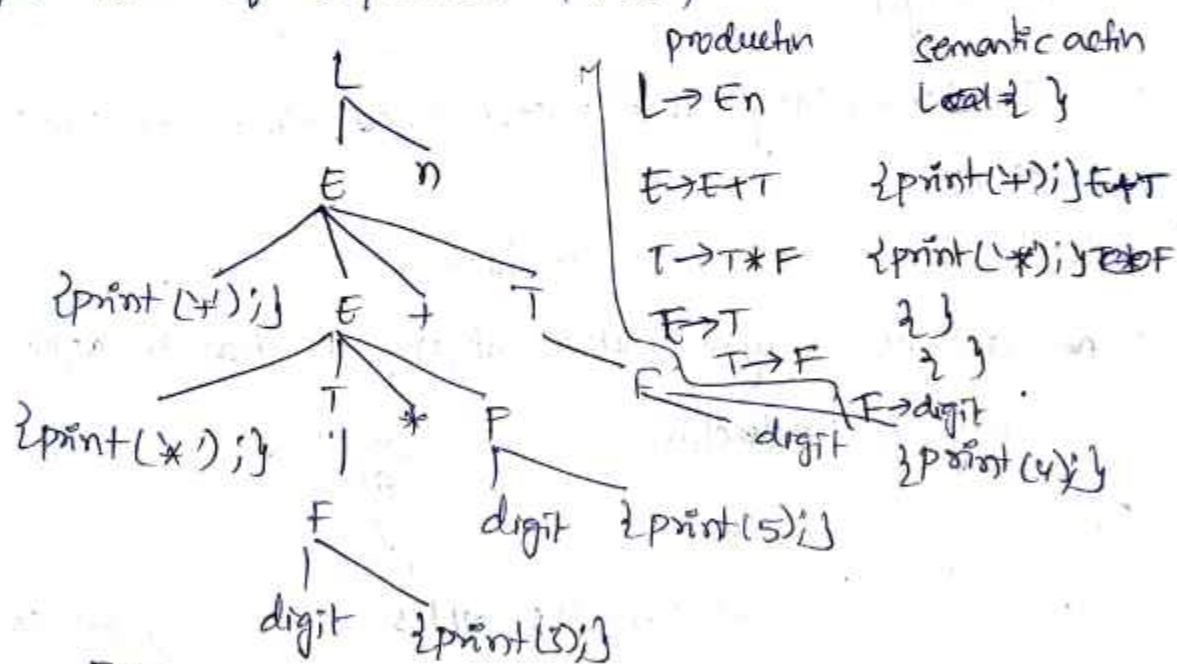


Fig: parse tree with actions embedded.

Intermediate code generation:-

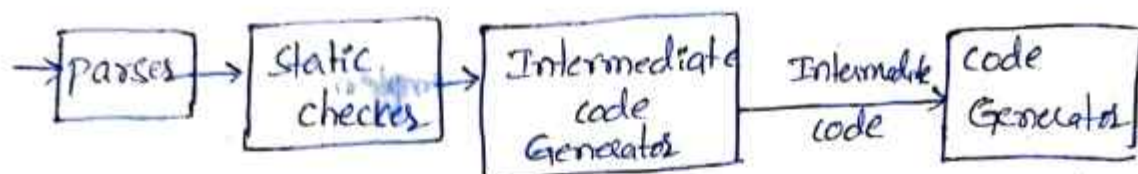


Fig:- Logical structure of a compiler's Frontend.

- Intermediate code is used to translate the source code into the machine-code.
- In the above figure parsing, static checking and Intermediate-code generation are done sequentially.
- Static type checking includes type checking, which ensures that operators are applied to compatible operand.
- ICG receives from its predecessor phase & semantic analysis phase.
- It takes i/p in the form of an annotated syntax tree.
- In process of translating a source program into target code compiler may construct a sequence of intermediate representation.



- Syntax trees are high level representation.
- A low level representation is suitable for machine-dependent tasks like register allocation and instruction selection.

⑦ Variants of syntax tree:-

Directed Acyclic Graph (DAG) for expressions:-

→ DAG is a data structure used for implementing transformations on basic blocks.

→ DAG nodes in:

→ DAG represent the structure of a basic block.

- In DAG internal nodes represent ^{operators} and leaf nodes represent identifiers, constants.



- Internal nodes represent the result of expression

→ The only difference b/w syntax tree and DAG is, In DAG a node has more than one parent.

Applications of DAG:-

* Determining the common subexpression.

* Determining which names are used inside the block and computed outside, the block.

• Determining which statement of the block could have their computed value outside the block.

• ^{eg} Eliminating common subexpressions it simplify the code.

Eg:- $a + a * (b - c) + (b - c) * d$

Syntax tree

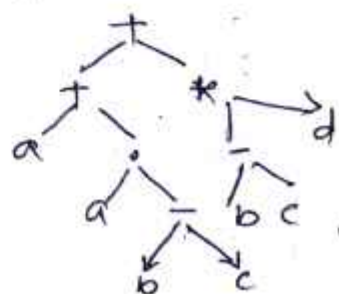
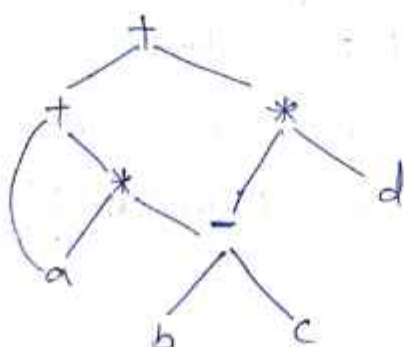


Fig:- DAG for expression $a + a * (b - c) + (b - c) * d$

production	Semantic Rules
$E \rightarrow E + T$	$E\text{-node} = \text{new Node}('+', E_1.\text{node}, T.\text{node})$
$E \rightarrow E_1 - T$	$E.\text{node} = \text{new Node}('-', E_1.\text{node}, T.\text{node})$
$E \rightarrow T$	$E.\text{node} = T.\text{node}$
$T \rightarrow (E)$	$T.\text{node} = E.\text{node}$
$T \rightarrow \text{id}$	$T.\text{node} = \text{new leaf}(\text{id}, \text{id}, \text{entry})$
$T \rightarrow \text{num}$	$T.\text{node} = \text{new leaf}(\text{num}, \text{num.val})$

Fig 1:- SDD ~~for~~ to produce syntax tree & DAG's

1) $P_1 = \text{leaf}(\text{id}, \text{entry}-a)$

$P_2 = \text{leaf}(\text{id}, \text{entry}-a) = P_1$ Ex 2:- construct DAG for

$P_3 = \text{leaf}(\text{id}, \text{entry}-b)$

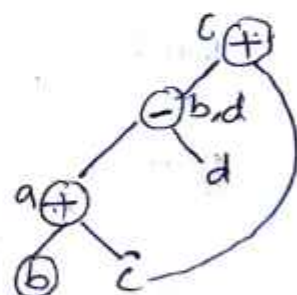
1 $a = b + c$ 3 $c = b + c$
2 $b = a - d$ 4 $d = a - b$

$P_4 = \text{leaf}(\text{id}, \text{entry}-c)$

$P_5 = \text{Node}('-', P_3, P_4)$

$P_6 = \text{Node}('*', P_1, P_5)$

$P_7 = \text{Node}('+', P_1, P_6)$



$P_8 = \text{leaf}(\text{id}, \text{entry}-b) = P_3$

$P_9 = \text{leaf}(\text{id}, \text{entry}-c) = P_4$

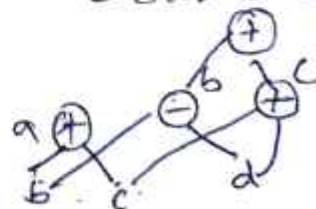
$P_{10} = \text{Node}('-', P_3, P_4) = P_5$

$P_{11} = \text{leaf}(\text{id}, \text{entry}-d)$

Ex 3:- 1 $a = b + c$ 2 $b = b - d$
3 $c = c + d$ 4 $e = b + c$

$P_{12} = \text{Node}('*', P_5, P_{11})$

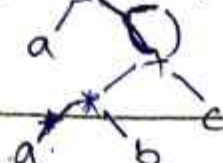
$P_{13} = \text{Node}('+', P_7, P_{12})$



The value number Fig:- steps for constructing the DAG.

Ex 2:- $a = b * -c + b * -c$

Ex 3:- $a = (a * b + c) - (a * b + c)$



The Value Number method for constructing DAG's

- Nodes of syntax tree & DAG are stored in array of records.
- Each row of array represent one record (node)
- In each record first field is operation code, indicating the label of the node.
- leaves has ~~the leaves~~ one additional field which holds the lexical value.
- Interior nodes have two additional field indicating left and right children.

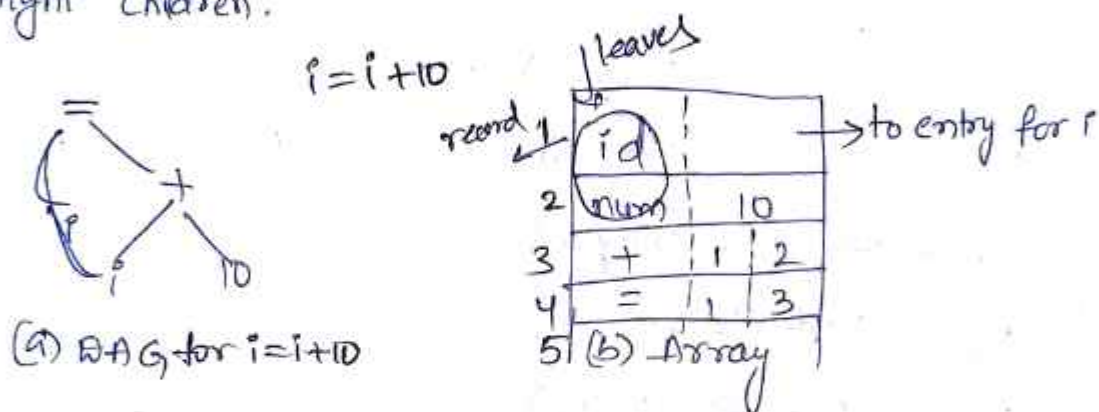


Fig 1 = Nodes of a DAG for $i = i + 10$ allocated in an array.

- The array index is used for reference a node rather than a pointer
- Initially the array is empty.
- First it searches for $\langle id, lexical \rangle$, if it is not there, we will make new record and so on.
- if already, record is present, it is just used for further records
- In the array, we refer to node by giving integer index of the record, for that node within the array this integer is called 'Value Number'. $\langle op, \text{value-num of left}, \text{value-num of right} \rangle$ is also called as signature of node

(23)

Drawbacks =

- Search options: It takes time for every New/Old record to search.
- To overcome this we are using hash-functions, in which the nodes are put into "buckets" (hash table)
 - These buckets have only few nodes.
 - hashtable is a data structure that support the dictionaries.
 - Dictionaries are used to insert & delete elements of a set.
 - Dictionaries are used to determine whether a given element is currently in the set. this is done
 - It search the elements in less time and independent of the size of set.

TO construct a has table for node.. of a ~~node~~, we hashfunction "h" is used that computes the index of bucket.

- ~~has~~ ~~copy~~ ~~has~~ The bucket index $h \in [0, L-1]$
- bucket can be implemented as linked list.
- An array indexed by hash value, hold the bucket headers, each of which point to first cell of list

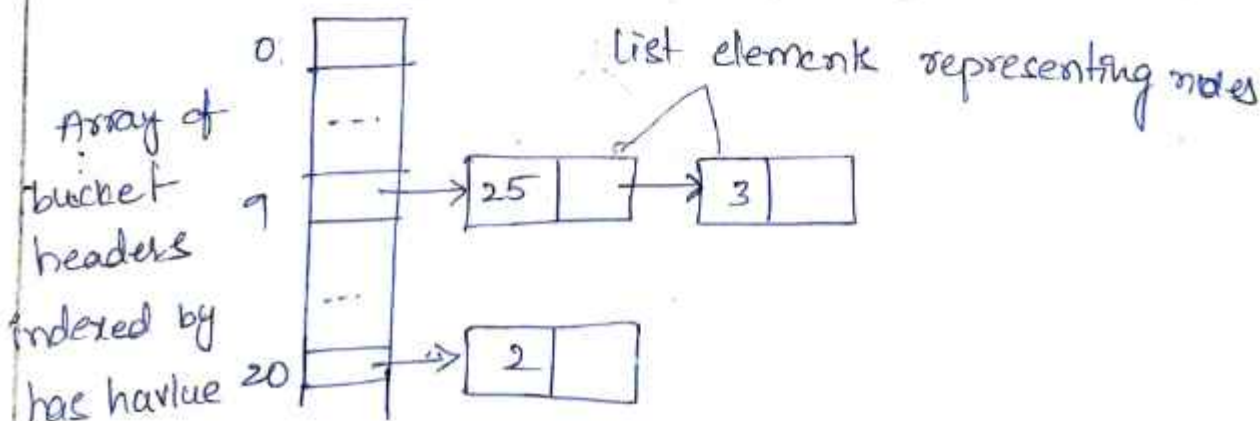


Fig: data structure for searching buckets

⑥

① syntax tree representation

- In three-address code

- eg:- source language expression $x+y*z$ is converted into sequence of 3-address instructions.

$t_1, t_2 \rightarrow$ compiler generated temporary variables.

Ex:- $a + a * (b - c) + (b - c) * d$

```

graph TD
    Root["+"] --- L1["+"]
    Root --- L1R["*"]
    L1 --- L2["a"]
    L1 --- L2R["*"]
    L2R --- L3["b"]
    L2R --- L3R["c"]
    L1R --- L4["d"]
    L1R --- L4R["*"]
    L4R --- L5["c"]
    L4R --- L5R["d"]
  
```

Fig (b) :- DAG

(a) Quadruple (2) Triple (3) Indirect Triples

Types of 3-address code: Addresses and Instructions:-

→ 3-address code can be implemented by using records with fields for the addresses

→ records are called quadruples and triples.

The address can be one of the following:

- A name → Source program names are addresses in 3-address code and names are replaced by pointer to its symbol table entry)
- A constant -
- compiler generated temporary variables.

List of the common three-address instruction forms:-

- (i) Assignment instruction → $x = y \text{ op } z$, where x, y, z - addresses
- (ii) → $x = \text{op } y$ where op - is unary operation $(-, \wedge)$
- (iii) copy instruction → $x = y$
- (iv) unconditional jump → goto L
- (v) conditional jump → if x goto L
→ if false x goto L
→ if $x \text{ relop } y$ goto L
- (vi) procedure call → $y = \text{call } P, n$
return y
P → is the address of starting of line of procedure P
n → argument address.
y ^{takes} → return value
- (vii) Indexed copy instruction → $x = y[i]$ } x, y, i are the variables.
 $z[i] = y$ }
- (viii) Address and pointer assignment $x = \&y$
 $x = *y$
 $*x = y$

Three address code

Quadruples :-

3-address code is implemented as objects (or) records with fields for the operators and operands.

→ Quadruple has 4-fields (i) op (ii) arg1 (iii) arg2 (iv) result.

eg: * Instruction like $x=y$ & $z=-y$ do not use arg2

* operator like param use neither arg2 nor result-

* conditional & unconditional jumps put the target label in result.

eg :- $a = b * -c + b * -c$

$$t_1 = -c$$

$$t_2 = b * t_1$$

$$t_3 = -c$$

$$t_4 = b * t_3$$

$$t_5 = t_2 + t_4$$

$$a = t_5$$

(a) Three address code

	op	arg1	arg2	result
(0)	-	c		t ₁
(1)	*	b	t ₁	t ₂
(2)	-	c		t ₃
(3)	*	b	t ₃	t ₄
(4)	+	t ₂	t ₄	t ₅
(5)	=	t ₅		a

(b) Quadruples

→ The disadvantage of Quadruples is too many temporary variables are needed, it requires more amount of memory.

→ In order to overcome we are using Triples.

Triples :- Triples has only three fields (i) op (ii) arg1 (iii) arg2

$$\text{Eq. } a = b * -c + b * -c$$

$$t_1 = -c$$

$$t_2 = b * t_1$$

$$t_3 = -c$$

$$t_4 = b * t_3$$

$$t_5 = t_2 + t_4$$

	op	Arg1	Arg2
(0)	-	c	
(1)	*	b	(0)
(2)	-	c	
(3)	*	b	(2)
(4)	+	(1)	(3)
5	=	a	(4)

(b) Triples representation of $a = b * -c + b * -c$;

→ using triples we refer results of an operation $bx opy$ by its position, rather than by an explicit temporary variables.

Indirect Triples:-

Indirect triples consist of listing of pointers to triples, rather than a listing of triples themselves.

→ with triples the result of an operation is referred to by its position, so moving an instruction may require us to change all references to that result. This problem does not occur with indirect triples.

instruction

35	(0)
36	(1)
37	(2)
38	(3)
39	(4)
40	(5)

	op	arg1	arg2
0	-	c	
1	*	b	(0)
2	-	c	
3	*	b	(2)
4	+	(1)	(3)
5	=	a	(4)

Fig: Indirect triple representation of three address code

Static single Assignment Form = (SSA)

SSA is a special case of 3-address code. SSA is an intermediate representation that facilitates certain code optimizations $\rightarrow$ In SSA each assignment to a variable should be specified with distinct names

Ex := $P = a + b$

$Q = P - c$

$P = Q * d$

$P = c - d$

$Q = P + Q$

$P_1 = a + b$

$Q_1 = P_1 - c$

$P_2 = Q_1 * d$

$P_3 = c - P_2$

$Q_2 = P_3 + Q_1$

(a) Three-address code (b) static single assignment form.

Ex := Intermediate program in three-address code SSA

$\rightarrow$ The least no. of temporary variables required to create 3-address code in SSA.

* A variable can only be initialized once in L-H-S

* A variable which is initialized in L-H-S could only be used R-H-S

Control flow:

Simple if, if-else, else-if, switch, for, while, do-while.
The translation of statements such as if-else-statement and while-statement is tied with translation of Boolean Expressions.

→ Boolean Expressions are used to.

- i) change the flow of control → Boolean exps are used as conditional exps in stmts that alter the flow of control.
- ii) compute the logical values for e.g. if (E) S,

Boolean Expressions:

→ A Boolean Expression can represent true or false as value.

Boolean Expressions are composed of boolean operators

&&, ||, and !

→ Boolean Expressions are generated by the following grammar

$$B \rightarrow B || B \mid B \&\& B \mid ! B \mid (B) \mid E \text{ rel } E \mid \text{true} \mid \text{false}$$

→ AND (&) OR are left-associative

→ "NOT" has higher precedence than AND & OR

Short circuit code: (Jumping code)

In short-circuit code, the boolean operators &&, || and ! are translated into jumps.

→ In short-circuit code the 2nd ^(expression) argument is evaluated only if 1st argument does not suffice to determine the value of Expressions.

Eg: if $(x < 100 || x > 200 \&\& x \neq y)$ $x = 0;$

In this translation the BE is true if control reaches label L₂.
If the expression is false, control immediately to L₄, skipping L₂ and the assignment $x = 0$.

if $x < 100$ goto L_2
 if false $x > 200$ goto L_1
 if false $x = y$ goto L_1
 $L_2: x = 0$
 $L_1:$

eg 2: $(x == y \parallel y == z)$:

Fig: Jumping code.

Flow of control statements:

→ Translation of Boolean expressions control - statements into three-address code.

Grammars
 $S \rightarrow \text{if } (B) S_1$ (d) $S \rightarrow \text{if } (B) \text{ then } S_1$
 $S \rightarrow \text{if } (B) S_1 \text{ else } S_2$... $\rightarrow$ condition (d) Boolean expression
 $S \rightarrow \text{while } (B) S_1$

∴ Grammars for simple if, if-else, while statements

(1) $S \rightarrow \text{if } (B) \text{ then } S_1$ (B) is evaluated 1st

code for simple if:

Semantic Rule:

B.true	B.code	B.true	$B.\text{true} = \text{newlabel}()$ $B.\text{true} = S_1.\text{next} = S.\text{next}$ $B.\text{false} = S_1.\text{next} = S.\text{next}$ <u>Intermediate code:</u> $S.\text{code} = B.\text{code} \parallel \text{label}(B.\text{true}) \parallel S_1.\text{code}$
	S ₁ .code	B.false	
B.false	S.next		

Fig: SDD for simple if-statement

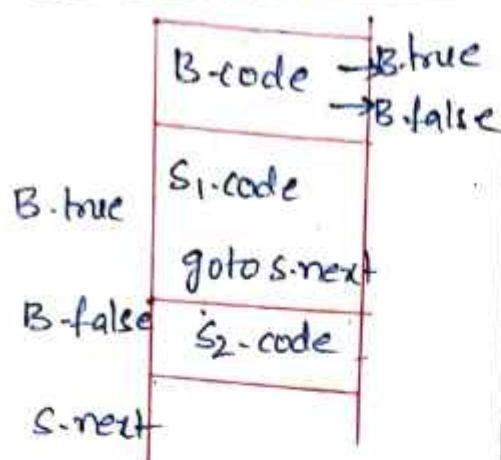
if () $\hookrightarrow$ true, newlabel() function produce three address code for B.true.
 1
 2
 3
 S

control flow translation of Boolean expression (d)
 Three address code for Boolean Expression (d) CDD (d) CDT for Boolean exp

production	Semantic Rules
$B \rightarrow B_1 B_2$	$\{ \begin{aligned} &B_1.true = B.true; \\ &B_1.false = newlabel(); \\ &B_2.true = B.true; \quad B_1.false \\ &B_2.false = B.false; \end{aligned}$ $B_1 B_2$ $B_1.code$ $B_2.code$ $B_1.true \rightarrow B_1.true$ $B_2.false \rightarrow B_2.true$ $B_2.false \rightarrow B_2.false$ Intermediate code: $B.code = B_1.code label(B_1.false) B_2.code$
$B \rightarrow B_1 \&\& B_2$	$\{ \begin{aligned} &B_1.true = newlabel(); \\ &B_1.false = B.false; \\ &B_2.true = B.true; \\ &B_2.false = B.false; \end{aligned}$ $B_1 \&\& B_2$ $B_1.code$ $B_2.code$ $B_1.true \rightarrow B_1.true$ $B_1.false \rightarrow B_1.false$ $B_2.true \rightarrow B_2.true$ $B_2.false \rightarrow B_2.false$ $B.code = B_1.code label(B_1.true) B_2.code$
$B \rightarrow !B_1$	$\{ \begin{aligned} &B_1.true = B.false; \\ &B_1.false = B.true; \\ &B.code = B_1.code \end{aligned}$
$B \rightarrow true$	$\{ B.code = gen('goto' B.true); \}$
$B \rightarrow false$	$\{ B.code = gen('goto' B.false); \}$
$B \rightarrow E_1 \text{ relop } E_2$	$\{ \begin{aligned} &B.code = E_1.code E_2.code \\ & gen('if' E_1 \text{ relop } E_2 'goto' B.true) \\ & gen('if' E_1 \text{ relop } E_2 'goto' B.false) \\ & gen('goto' E.false) \end{aligned}$
Eg if a < b goto E.true goto E.false $E \rightarrow (E_1)$	$\{ \begin{aligned} &E_1.true = E.true; \\ &E_1.false = E.false; \\ &E.code = E_1.code; \end{aligned}$

$S \rightarrow \text{if } (B) \text{ then } S_1 \text{ else } S_2$

code for if-else:

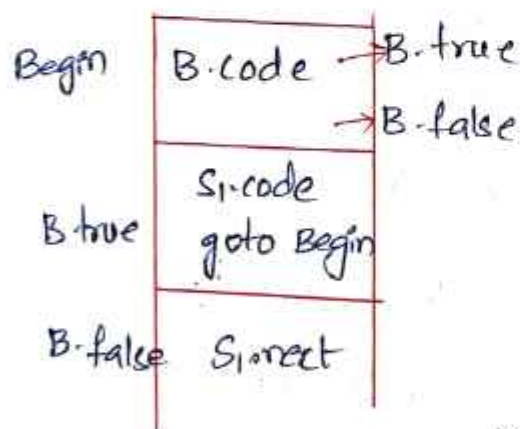


Semantic rules for if-else stmt:

$B.\text{true}_i = \text{newlabel}()$
 $(B.\text{false} = \text{newlabel}())$
 $S_1.\text{next} = S.\text{next}$
 $S_2.\text{next} = S.\text{next}$
 $S.\text{code} = \text{Intermediate Three address code} =$
 $S.\text{code} = B.\text{code} \parallel \text{label}(B.\text{true}) \parallel S_1.\text{code}$
 $\parallel \text{gen}('goto' S.\text{next})$
 $\parallel \text{label}(B.\text{false}) \parallel S_2.\text{code}$

(iii) while (B) then S₁

code for while



Semantic Rules

$\text{Begin} = \text{newlabel}()$
 $B.\text{true} = \text{newlabel}()$
 $B.\text{next} = \text{begin}$
 $B.\text{false} = S.\text{next}$

Intermediate code

$S.\text{code} = \text{label}(\text{Begin}) \parallel B.\text{code}$
 $\parallel \text{label}(B.\text{true}) \parallel S_1.\text{code}$

production

$P \rightarrow S$

Semantic Rules

$S.\text{next} = \text{newlabel}()$

$P.\text{code} = S.\text{code} \parallel \text{label}(S.\text{next})$

$S \rightarrow \text{assign}$

$S.\text{code} = \text{assign-code}$

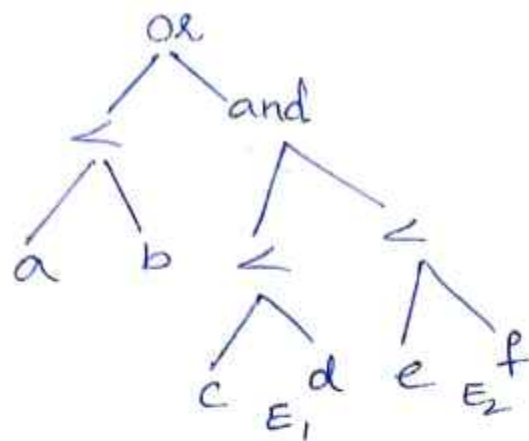
$S \rightarrow S_1 S_2$

$S_1.\text{next} = \text{newlabel}()$ $S_2.\text{next} = S.\text{next}$

$S.\text{code} = S_1.\text{code} \parallel \text{label}(S_1.\text{next}) \parallel S_2.\text{code}$

$\parallel \text{gen}('goto' \text{begin})$

eg: ① $a < b$ or $c < d$ and $e < f$



if $a < b$ goto E.true
goto E_1

E_1 : if $c < d$ goto E_2
goto E.false

E_2 : if $e < f$ E.true
goto E.false

eg: ② if $(x < 100 \parallel x > 200 \text{ \& } x \neq y)$ $x = 0$

if $x < 100$ goto L_2

goto L_3

L_3 : if $x > 200$ goto L_4

goto L_1

L_4 : if $x \neq y$ goto L_2

goto L_1

L_2 : $x = 0$

L_1 :

Types and Declarations:-

* Type checking uses logical rules to decide about the behaviour of program at runtime.

* It also ensures that types operand match type expected by the operator.

Eg:- "&&" operation Java expect its two operands to be boolean

int * float $\rightarrow$ type error

$\rightarrow$ Determine the storage needed

Translation Application:

Compilers translate a type of name into storage

Compiler also determines the amount of storage required to store the type name at run time.

Type Expression:-

Type Expression is either a basic type & formed by applying an operator called type constructs to a type expression.

$\rightarrow$ T.E are used represent the structure of type,

$\rightarrow$ T.E are primitive datatypes.

$\rightarrow$ Type name:- is a Type Expression Eg:- ~~int~~ typedef abc int.

Type Expressions are of two types.

(i) Basic type:- Basic type for language are int, real, boolean, char, float, and void. A special type, type-error is used to indicate type error.

Eg:- int x; type abc int;
 int a;
 abc b;

is a=b; $\rightarrow$ depends on language

(ii) type constructs (a) type Name:-

$\rightarrow$ type constructs applied to list of type expressions.

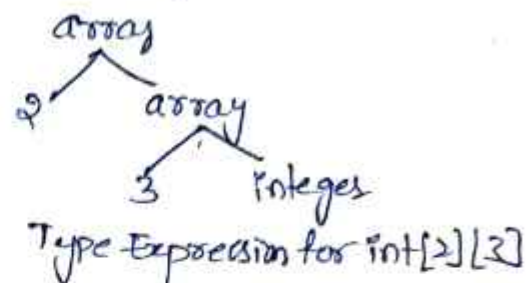
$\rightarrow$ types are formed applying an operator called type constructs to type expression.

(i) Arrays:- Arrays are specified as array (I, T) where $I \rightarrow \text{Int (or) range of integers}$
 $T \rightarrow \text{Type}$.

Eg₁:- In "C" Declaration "int a[100]" identifies type of "a" to be
 array(100, integer) $\downarrow$ int(100) a;

Eg₂:- T.E for int[2][3], "2 array with 3-integers".

obj name $\swarrow$ array (2, array(3, integer))
 $\downarrow$
 Type Expression



Eg₃:- int a[100], b[50];
a=b X type clash.

→ A type Expression can be formed by applying array constructor to a number & type expression.

(ii) Record:- A record is a data structure with name fields.

→ A type Expression can be formed by applying a record type constructor to the field name and their types.

Eg₁:- Struct st

{
 int s;
 float f;
 };

Struct st s1;

record (f, (int, float))

record (s1, integer)

record (s, float)

Eg₂:- Named records are product with named elements for record structure with 2 named fields.

length (an integer) and word (of type array (10, char))
 the record is of type.

record (length X Int) X (word X array (10, char))

struct

{ int length

; char word[10];

→ type Expression may contain variable whose values are T.E eg: int a=5;

product:- s and t are 2 T.E then their cartesian product sxt is a
 typ Expression eg:- int * int.

Function:- Function maps a collection of types to another represented by $D \rightarrow R$
 where D is domain R is range of function.

eg:- int f1 (int x, char y, float z)

{
 return m;

o/p: int

Domain = (int x char x float)

Range -

→ T.E "int * int → char" represent a function that takes 2-integers & returns a char value

Type Equivalence:-

→ Two type are said to be equivalent if and only if an operand of one type in an expression is substituted for one of the other type, without type conversion.

Type equivalence are of two types.

1) Name equivalence:-

The two type expression are said to be name equivalence if they have same name & label.

Eg:-
 typedef int value;
 typedef int total;
 :
 value var1, var2;
 total var3, var4;

Eg2: type def struct Node
 {
 int x;
 } node;
 node * first & second;
 struct node * last1, * last2;

→ In the above eg1, var1 and var2 are name equivalence because their types are same.

→ var3 & var4 also Name equivalence.

→ but var1 & var3 are not name equivalent because their types are different

Structural equivalence:-

→ If two expression are the basic type (a)

→ Formed by applying the same constructs to structurally types equivalent types then those expression are called structurally equivalent

- (i) It checks the structure of type
- (ii) Determines equivalence by whether they have same constructs applied to structurally equivalent types

Eg := type array (I_1, T_1) and array (I_2, T_2) structurally equivalent if $I_1 = I_2$ & $T_1 = T_2$

$I_i \rightarrow$ Index of array

$T_i \rightarrow$ type of

array [a(100), b(50)]

array a(100), b(100)

$\downarrow$
structurally equivalent

eg1: type def int value $\leftarrow x$
typedef int number $\leftarrow y$

x : array (50, int)

y : array (100, int)

eg2:

S_1	S_2	Equivalence	Reason	Reason
char	char	S_1 is equivalent to S_2		similar basic type
pointer (char)	pointer (char)	S_1 is equivalent to S_2		similar constructs pointer to the char type.

Declarations :=

$D \rightarrow T \text{ id} ; D/E$

$T \rightarrow B \text{ c} / \text{record 'id'}$

$B \rightarrow \text{int} / \text{float}$

$C \rightarrow \epsilon / [\text{num}] C$

D → Sequence of declarations.

T → basic & array and record types

B C → 'component' - generates zero or more integers within the brackets.

- Array type consists of basic type specified "B", followed by array component C.

Eg:- `int [10][11]`

- Record type is sequence of declaration for field of the record all surrounded by curly braces

record {int, a}

Storage layout for local Names:-

- compiler converts the typenames into the storage.
- and determines the amount of storage needed to store the typename at runtime.
- at compile time we can use these amount to assign a ^{type} name to relative address.

$$\text{relative address} = \text{offset} + \text{program counter}.$$

- Relative & types are saved in symbol table entry for type name.
- Data of varying length such as string or whose size cannot determined until runtime such as dynamic arrays.

→ The width of a type is no. of storage units needed for objects of that type.

SDT computes types and their widths for basic and array types.

$T \rightarrow B \{ t = B.type; w = B.width; \}$

$C \rightarrow \{ t.type = C.type; T.width = C.width; \}$

$B \rightarrow \text{int} \{ B.type = \text{integer}; B.width = 4; \}$

$B \rightarrow \text{float} \{ B.type = \text{float}; B.width = 8; \}$

$C \rightarrow \epsilon \{ C.type = t; C.width = w; \}$

$C \rightarrow [\text{num}] t, \{ C.type = \text{array}(\text{num-value}, C_1.type); C.width = \text{num-value} \times C_1.width; \}$

Fig 1:- SDT for computing their types & widths

→ These declaration are represented with DAG & parse tree

Ex:- parse tree for $\text{int}[2][3]$

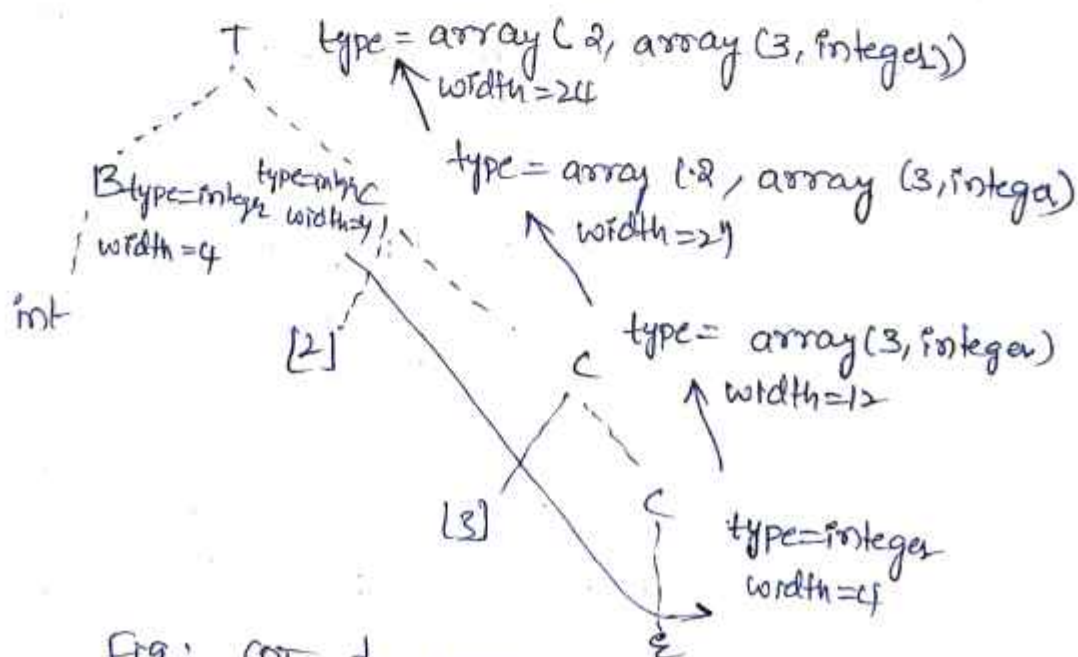


Fig:- SDT of array types

Sequence of declarations:

- In procedures all the declaration are passed at a time.
- all declarations in single procedure to be pass as a group.

$P \rightarrow_D \{ \text{offset} = 0; \}$

$D \rightarrow \begin{array}{l} T \text{ id}; \quad \{ \text{top.put}(\text{id.lexeme}, T.\text{type}, \text{offset}); \\ \text{(a) id } T; \\ \text{D; D} \quad \text{offset} = \text{offset} + T.\text{width}; \end{array}$

$D \rightarrow \epsilon$

offset — is variable to keep track of the next available relative address.

$D \rightarrow T \text{ id};$ D — creates a symbol table entry for
executing $\text{top.put}(\text{id.lexeme}, T.\text{type}, \text{offset})$

top — The current symbol table

top.put → ~~new~~ creates a symbol table entry for
id.lexeme with type & relative address.

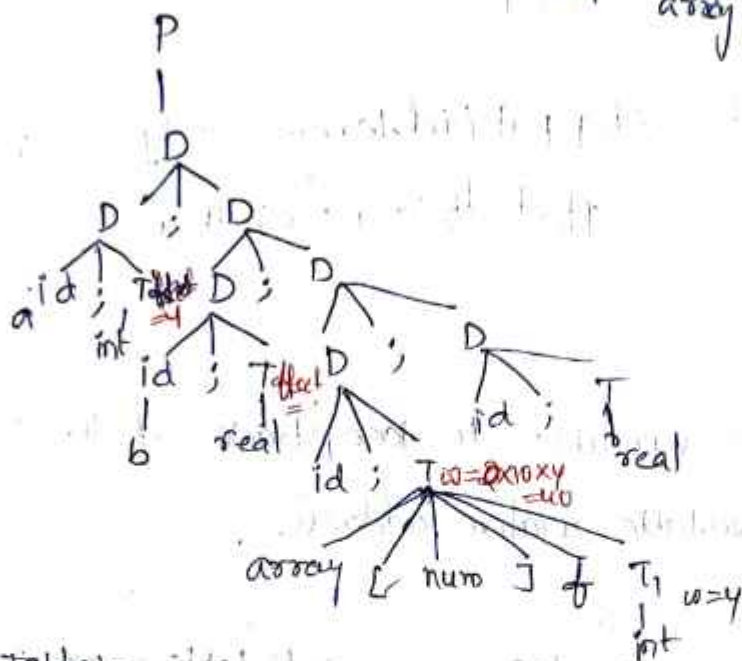
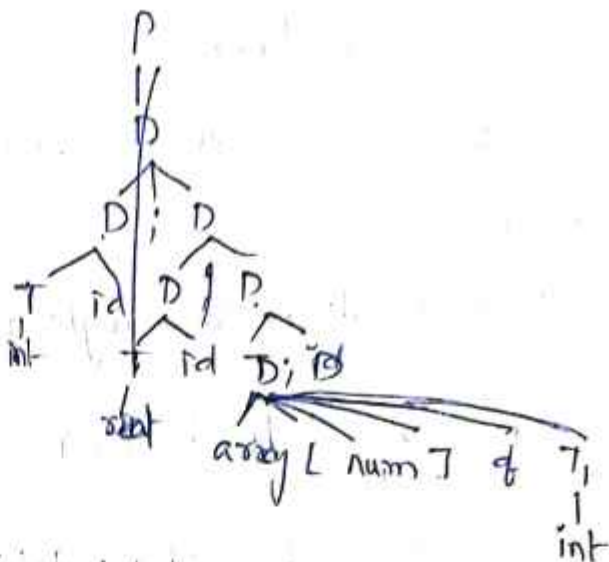
Field in Records & classes:

$T \rightarrow \text{record } \{ D \}$

The field in this record is type specified by the sequence of declaration grouped by D.

- The field names with record must be distinct
- The offset & relative address for fieldname is relative to the data area for that record


```
a: int;  
b: real;  
c: array[10] of integer;  
d: ↑real;  
}
```



Name _____

TYPE

offset

a

Integers

Q. [0-3]

$$\text{off set} = 0 + 4 = 4$$

b

real

4

$$2x + y + 8 = 12$$

2

```
array(10, int)
```

12

$$df_{\text{cel}} = 12 + 40 = 52$$

d

real

52

Type checking:-

- Compiler checks whether the program is following type rules or not.
- information about data types is maintained & Computed by compiler.
- Type checker is a module of a compiler devoted to typechecking tasks.
- To do typechecking a compiler needs to assign a TE to each component of source program.
- Compiler determines TE conform to collection of logical rules that is called type system for the source language.
- Typechecking. catch the errors in program.
- Assign types to values.
- Simple situation:- check types of objects & report a type error in case of a violation.
- more complex:- Incorrect types may be corrected (type Coercing).

Static	Dynamic
→ Type checking done at compile time.	→ perform during program Execution.
→ properties can be verified before program run.	→ permits programmer to be less concern with C, Pascal strongly.
→ Can catch many common Errors.	→ Mandatory in some situations such as array, bounds check.
→ Desirable when faster Execution importance	→ more robust and clearer code required.

Eg: Pascal, C-type

→ Type checking have been used to ~~generally~~ improve the Security of System.

Rules for Type checking:-

Type checking has two forms

i) Synthesis

ii) Interference

i) Type Synthesis :-

→ It derives the expressions from the types of its subexpressions.

→ It must be declared before they are used.

Ex:- The type of $E_1 + E_2$ is defined in terms of the types of E_1 and E_2 .

If f has type $s \rightarrow t$ and x has type s ,
then expression $f(x)$ has type t

Ex:- `add(int a, float b)`

{

}

`fn1(float c, int d) → t`

{
 `add(2, 2.5)` ↓

}

`add(float, int)`

[∵ int changes to float,
float changes to int]

→ Here f and x denote expression, $s \rightarrow t$ denote a function from s to t .

→ This rule for functions with one argument carries over to functions with several arguments.

ii) Type inference:-

→ It generally determines the type of language construct from the way it is used.

→ Ex:- $E_1 + E_2$ i.e., $\begin{matrix} 2 + 5 \\ \text{int} \quad \text{int} \end{matrix} = \text{Datatype will be int}$

$\begin{matrix} \text{abc} + \text{abc} \\ (\text{string}) \quad (\text{string}) \end{matrix} = \text{Datatype will be string}$

→ There is no need to declare variables.

→ Type inference are used in meta languages.

If $f(x)$ is an expression

then for some α and β , f has a type $\alpha \rightarrow \beta$

and x has type α

Type Conversion or type Casting:-

→ A type cast is basically a Conversion from one type to another.

→ There are two types of Conversions

- 1) Implicit type Conversion
- 2) Explicit type Conversion

1) Implicit type Conversion:- (smaller to bigger)

→ If a compiler Converts one data ^{type} into another type of data automatically.

→ There is no data loss

Ex:- short a = 20;

int b = a; // Implicit Conversion

Assign:- bool → char → short int → int → long → float

2) Explicit type Conversion:-

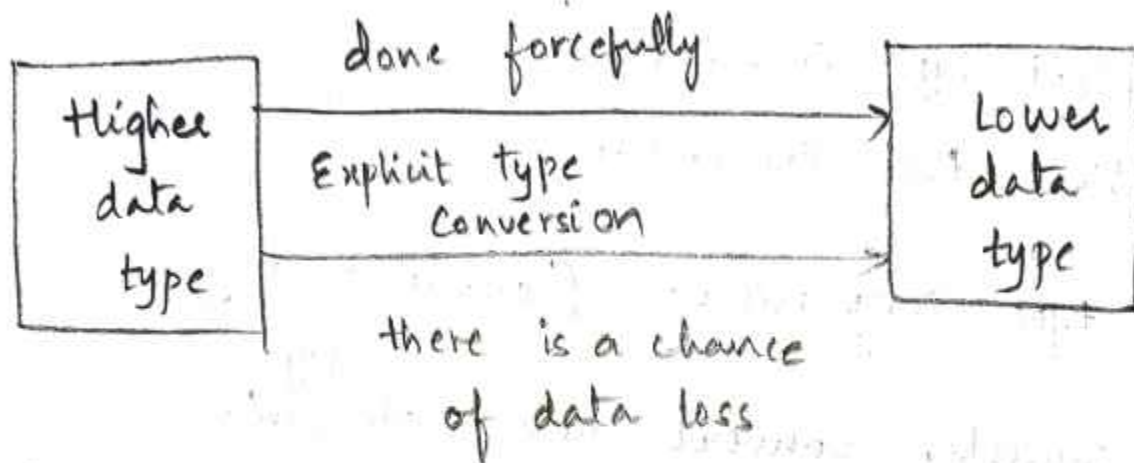
→ When data of one type is converted explicitly to another type with the help of predefined functions.

→ There is a data loss.

→ Conversion done forcefully.

→ Some Conversions cannot be made Implicitly

int to short (∵ int range is more than short
so there is a chance of data loss)



Ex:- $t_1 = (\text{float}) 2$

$t_2 = t_1 * 3.14$

Ex:- if ($E_1 \cdot \text{type} = \text{integer}$ and $E_2 \cdot \text{type} = \text{Integer}$)
 $E \cdot \text{type} = \text{integer};$

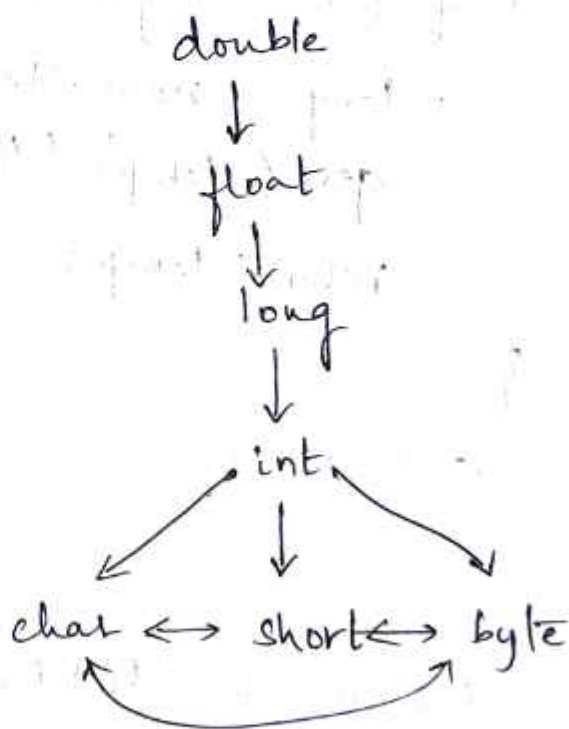
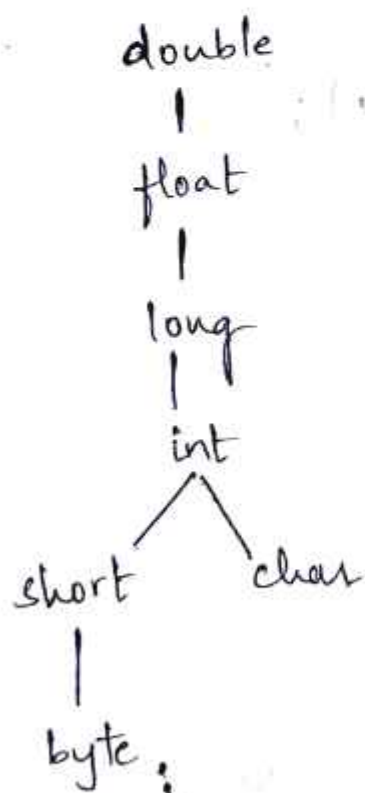
else if ($E_1 \cdot \text{type} = \text{float}$ and $E_2 \cdot \text{type} = \text{Integer}$) ---
 $E \cdot \text{type} = \text{float};$

— — — —

→ two Conversions

- i) Widening Conversions
- ii) Narrowing Conversions.

- Widening conversions generally preserve information
- narrowing conversions generally lose information
- widening rules - any lower can be widened to higher type.
- a char can be widened to int or float but char cannot be widened to short.
- narrowing rules - a type s can be narrowed to type t if there is a path from t.



1) $\text{max}(t_1, t_2)$ takes two types t_1 and t_2 and returns the maximum of two types in widening hierarchy.

2) $\text{widen}(a, t, w)$ generates type conversions if needed to widen the contents of an address a of type t into a value of type w .

```
Addr { widen( Addr a, Type t, Type w)
```

```
    if (t=w) return a;
```

```
    else if (t = integer and w = float) {
```

```
        temp = new Temp();
```

```
        gen( temp = '(float)' a );
```

```
        return temp;
```

```
    }
```

```
    else error;
```

```
}
```

→ Semantic Action $E \rightarrow E_1 + E_2$

```
 $E \rightarrow E_1 + E_2$  {  $E.type = \text{max}(E_1.type, E_2.type);$ 
```

```
     $a_1 = \text{widen}(E_1.addr, E_1.type, E.type);$ 
```

```
     $a_2 = \text{widen}(E_2.addr, E_2.type, E.type);$ 
```

```
     $E.addr = \text{new Temp}();$ 
```

```
     $\text{gen}(E.addr = 'a_1' + 'a_2');$  }
```


Overloading of Functions and operators:-

An overloaded Symbol has different meanings depending on its context. Overloading is resolved when a unique meaning is determined for each occurrence of a name.

Ex:- The + operator in Java denotes either string concatenation or addition, depending on the type of its operands.

```
void err() { --- }
```

```
void err(String s) { - - }
```

Intermediate code for switch statements (a) Three Address Code

Translation of switch statements:

Switch statement syntax:

```

switch (E)
{
  case  $v_1$  :  $s_1$ 
  case  $v_2$  :  $s_2$ 
  ...
  case  $v_{n-1}$  :  $s_{n-1}$ 
  default :  $s_n$ 
}

```

eg := switch (x+y)

```

{
  case 1: a = a + 2;
  break;
  case 4: b = b * 5;
  break;
  case 6: c = c / 2;
  break;
  default: d = d - 2;
  break;
}

```

Translation of switch statements:

Code to evaluate E into t
goto test

L_1 : code for s_1
goto next

L_2 : code for s_2
goto next

...

L_{n-1} : code for s_{n-1}
goto next

L_n : code for s_n
goto next

test: if $t = v_1$, goto L_1
if $t = v_2$, goto L_2

if $t = v_{n-1}$, goto L_{n-1} , goto L_n
next:

Three Address Code :=

if $t = v_1$, goto L_{n-1}
goto L_n

next:

1. $t_1 = x + y$

2. If ($t_1 = 1$) goto 8

3. If ($t_1 = 4$) goto 11

4. If ($t_1 = 6$) goto 14

5. $t_2 = d - 2$

6. $d = t_2$

7. goto Next

8. $t_3 = a + 2$

9. $a = t_3$

10. goto Next

11. $t_4 = b * 5$

12. $b = t_4$

13. goto Next

14. $t_5 = c / 2$

15. $c = t_5$

16. goto Next

17. Next

if $t_i = v_{n-1}$ goto l_{n-1}
goto l_n

next:

Intermediate code for ~~procedures~~ procedures: = (d) Three address code

$D \rightarrow \text{define } T \text{ id } (F) \{S\}$ | $S \rightarrow \text{adds stmts that returns the value of an expression.}$

$F \rightarrow E \mid T \text{ id}, F$

$S \rightarrow \text{return } E;$

$E \rightarrow \text{id } (A);$

$A \rightarrow E \mid E, A$

$\rightarrow E \rightarrow \text{adds function calls, with actual parameters A.}$

Non-terminals D and T generates declarations and types.

$\rightarrow$ Function definition generated by D consists of keyword define, a return type, the function name, formal parameters in parenthesis and function body consisting of statements.

$\rightarrow$ Non-terminal F generates zero or more formal parameters.

where formal parameters consists of a type followed by identifiers

$\rightarrow$ Non-terminal S & E generate statements & expressions.

float add()

 float a;

 float a, int b;

 {

 return add();

 }

$\rightarrow$ In three-address code, a function call is unraveled into the evaluation of parameters in preparation for a call, followed by call itself and the parameters are passed by value.

Eg: If the given function is in the form of

$P(A_1, A_2, A_3, \dots, A_n)$ Eg: $n = f(a[i]);$

param A_1

param $\dots A_2$

 |

param $\dots A_n$

call P, n

Translated into three-address code as follows

1) $t_1 = i * 4$

2) $t_2 = a[t_1]$

3) param t_2

4) $t_2 = \text{call } f, 1$

5) $n = t_2$

$P \rightarrow$ is function name.

$n \rightarrow$ no. of arguments.

- The first 2 lines compute the value of expression $a[i]$ into temporary 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 the function call to t_3 .

→ Functions types:-

- The type of function must encode the return type and types of the return type and the types of the formal parameters.
- Let "void" be a special type that represent no parameters or no return type.
- Whenever the function is called the function name is name entered into the symbol table for use in the rest of the program.
- The formal parameters are stored in the Activation Record. For storing formal parameters the Activation Records are used.

Eg 2:- void main()

```
{
  int x, y
  ...
  swap(&x, &y);
}
```

```
void swap(int *a, int *b)
{
  int i;
  i = *b;
  *b = *a;
  *a = i;
}
```

Three address code

1. call main
2. param &x
3. param &y
4. call swap, 2

Eg 3:- float add() or float add(int a)

float add(int a, float b)

```
{
  return add(); or return add(x);
  return add(x, y);
}
```

Formal parameters

Actual parameters

UNIT-V

Machine-Independent Optimizations:- The principal sources of optimization, Introduction of Data-flow Analysis, foundations of Data-flow Analysis, constant propagation, partial Redundancy Elimination, Loops in flow Graphs.

Optimization is the process of transformation of code to an efficient code. Efficiency is in terms of space requirement and time for its execution without changing the meaning of the given code.

The following constraints are to be considered while applying the techniques for code improvement.

- ① The transformation must preserve the meaning of the program, that is, the target code should ensure semantic equivalence with source program.
- ② Program efficiency must be improved by a measurable amount without changing the algorithm used in the program.
- ③ When the technique is applied on a special format, then it is worth transforming.

Optimization can be classified as

1) Local optimizations:-

Optimizations performed within a single basic block are termed as local optimizations. These techniques are simple to implement and does not require any analysis since we do not require any information relating to how data and control flows.

2) Global optimizations:-

Optimizations performed across basic blocks is called global optimization. These techniques are complex as it requires additional analysis to be performed across basic blocks, this analysis is called data-flow analysis.

Optimization techniques can be applied to intermediate code or on the final target code. It is a complex and a time-consuming task that involves multiple subphases, sometimes applied more than once.

* most compilers allow the optimization to be turned off to speed up compilation process.

* To apply optimization it is important to do control flow analysis and data flow analysis followed by transformations.

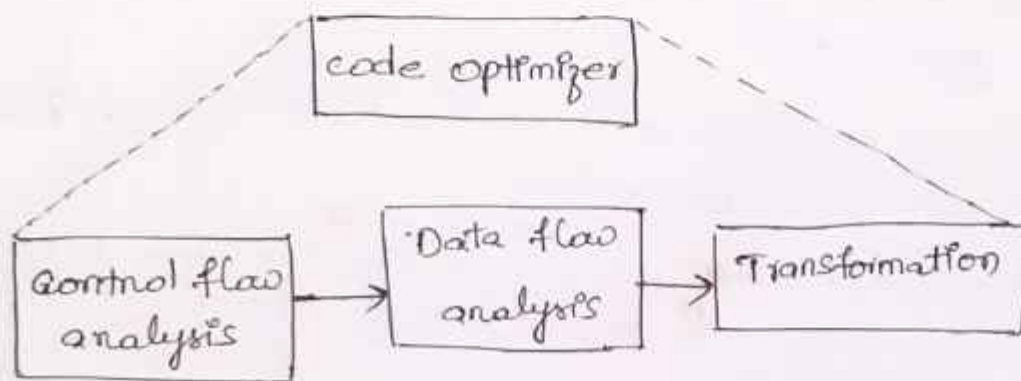


Fig: code optimization model.

Control flow Analysis:— It determines the control structure of a program and builds a control flow graph.

Data flow Analysis:— It determines the flow of scalar values and builds data flow graphs. The solution to flow analysis propagates data flow information along a flow graph.

Transformation:— Transformations help in improving the code without changing the meaning or functionality.

Flow Graph:—

A graphical representation of three address code is called flow graph. The node in the flow graph represent a single basic block and the edges represent the flow of control.

* These flow graphs are useful in performing the control flow and data flow analysis and to apply local or global optimization and.

Basic Block :- A basic block is a set of ~~conseq~~ consecutive statements that are executed sequentially, once the control enters into the block then every statement in the basic block is executed one after the other before leaving the block.

Ex: for the statement $a = b + c * d / e$ the corresponding a set of three address code is.

$$\begin{aligned} t_1 &= c * d \\ t_2 &= t_1 / e \\ t_3 &= b + t_2 \\ a &= t_3. \end{aligned}$$

All these statements correspond to a single basic block.

Procedure to identify the Basic Blocks.

Given a three address code, first identify the leader statements and group the leader statements with the statements up to the next leader.

* To identify the leader use the following rules:

1. First statement in the program is a leader.
2. Any statement that is the target of a conditional or unconditional statement is a leader statement.
3. Any statement that immediately follows a conditional/unconditional statement is a leader statement.

Example :- Identify the basic blocks for the following code fragment

```
main()
{
    int p=0; n=10;
    int a[n];
    while ( i <= (n-1) )
    {
        a[p] = p * p;
        p = p+1;
    } return; }
```

The three address code for initialize function is as follows:

- (1) $i := 0$ $\longrightarrow$ leader1 using rule 1
- (2) $n := 10$
- (3) $t_1 := n - 1$ $\longrightarrow$ leader2 using rule 2
- (4) if $i > t_1$ goto (12)
- (5) $t_2 := i * i$ $\longrightarrow$ leader3 using rule 3.
- (6) $t_3 := 4 * i$
- (7) $t_4 := a[t_3]$
- (8) $t_4 := t_2$
- (9) $t_5 := i + 1$
- (10) $i := t_5$
- (11) goto (3)
- (12) return. $\longrightarrow$ leader 4 using rule 3.

1) $i = 0$
2) $n = 10$

B1

Basic Block 1 includes statements (1) and (2)

(3) $t_1 := n - 1$
(4) if $i > t_1$ goto (12)

B2

Basic block 2 includes statements (3) and (4)

(5) $t_2 = i * i$
(6) $t_3 = 4 * i$
(7) $t_4 = a[t_3]$
(8) $t_4 = t_2$
(9) $t_5 = i + 1$
(10) $i := t_5$
(11) goto (3)

B3

Basic block 3 includes statements (5) - (11)

(12) return.

B4

Basic block 4 includes statement (12)

Fig: Basic Blocks.

Flow Graph :-

Flow graph shows the relation between the basic block and its preceding and its successor blocks. The block with the first statement is B₁. An edge is placed from block B₁ to B₂, if block B₂ could immediately follow B₁ during execution or satisfies the following conditions.

- The last statement in B₁ is either conditional or unconditional jump statement that is followed by the first statement in B₂ or
- The first statement in B₂ follows the last statement in B₁ and is not an unconditional / conditional jump statement

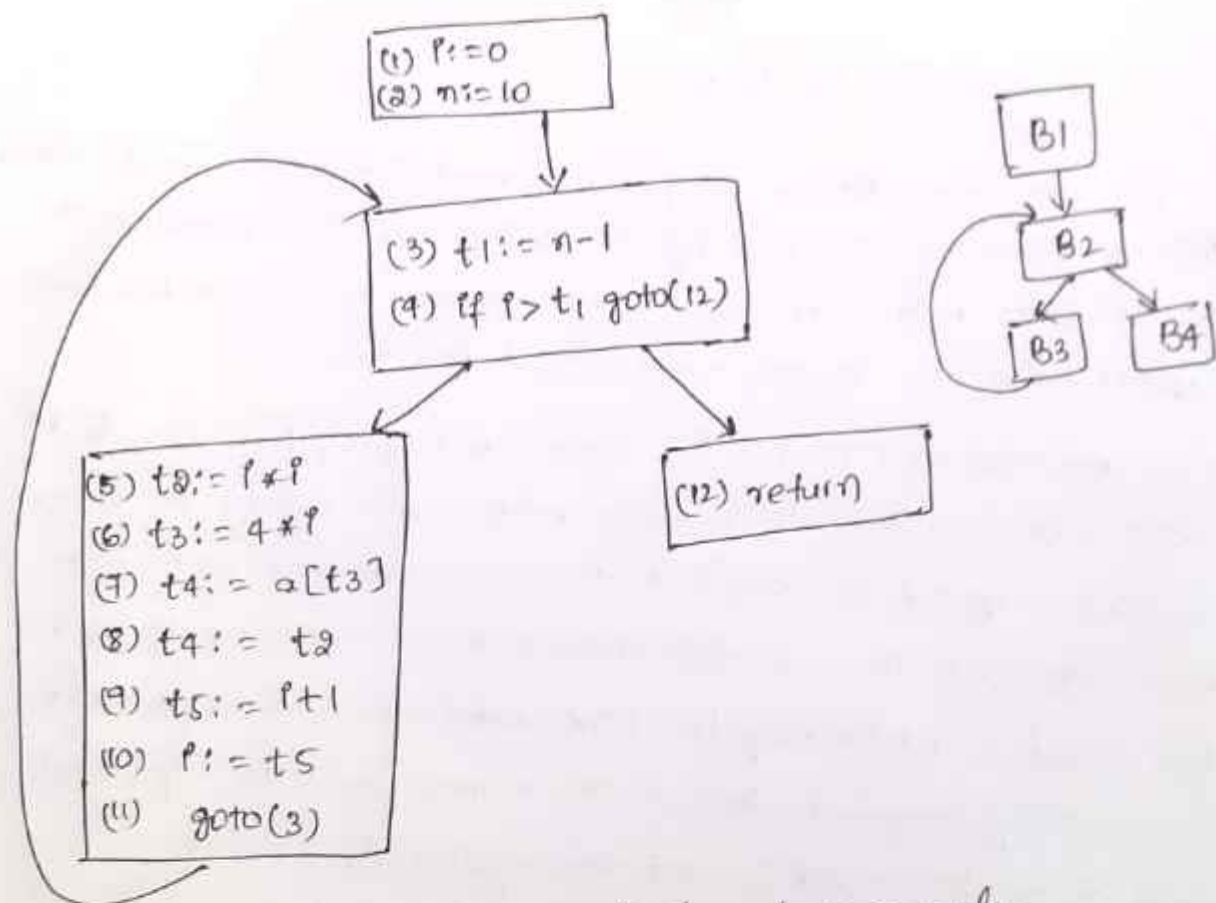


Fig: Flow graph for above example.

DAG representation of Basic Block.

A DAG is a useful data structure for implementing transformations within a basic block. It gives pictorial representation of how values computed at one statement are useful in computing the values of other variables.

* It is useful in identifying common sub-expressions within a basic block.

* A DAG has nodes, which are labeled as follows.

- ① The leaf nodes are labeled by either identifiers or constants.
If the operators are arithmetic then it always requires two r -value
- ② The labels of interior nodes correspond to the operator symbol.

Note: The DAG is not the same as a flow graph. Each node in a flow graph is a basic block, which has a set of statements that can be represented using DAG.

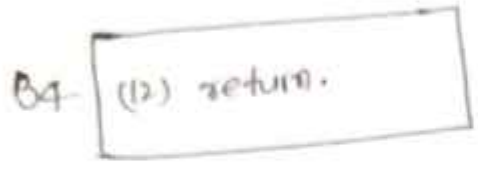
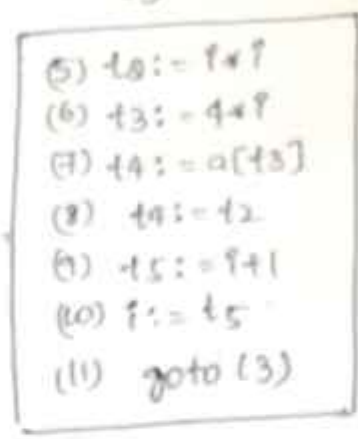
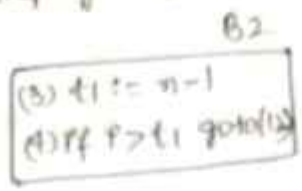
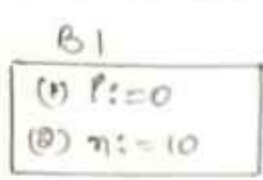
Construction of DAG.

To construct DAG, we process each statement of the block. If the statement is a copy statement, that is, a statement of the form $a = b$, then we do not create a new node, but append label a to the node with label b .

* If the statement is of the form $a = b \text{ op } c$, then we first check whether there is a node with same values as $b \text{ op } c$, if so, we append the label a to that node. If such node does not exist then we first check whether there exists nodes for b and c which may be leaf nodes or an internal nodes if recently computed, then create a new node for 'op' and add it the left child b and the right child as c .

* Label this new node with a . This would become the value of a to be used for next statements; hence, we mark the previously marked nodes with a as a_0 .

Example:- Let us consider basic block B1 to B4 and construct DAG for each block step by step



For Block B1

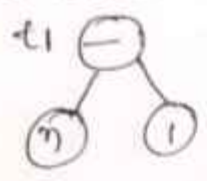
i ④

Fig: DAG for Block B1

* For the statement ~~in B1~~, first we create a leaf labeled 4 and attach identifier i to it

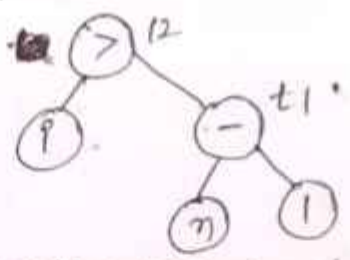
For Block B2

For first statement, which we first create nodes for n and 1, then create node for operator (-) and label it as t1.



Fig(a)

DAG for first statement in Block B2

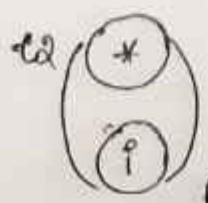


Fig(b): DAG for Block B2

* Fig(b) shows the, second statement, which is a conditional statement, we create a node for operator > and label it as 12 as when this condition is satisfied it should goto statement 12.

For Block B3

* For the first statement, we create nodes for i and use as right child, then create node for operator (*) and label it as t2.



Fig(a) DAG for Block B3.

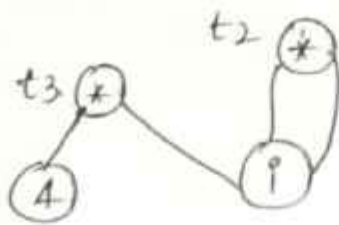
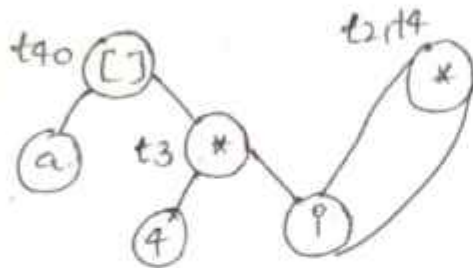
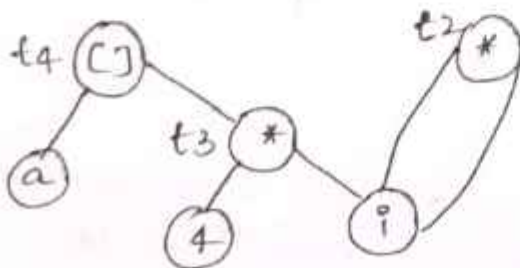


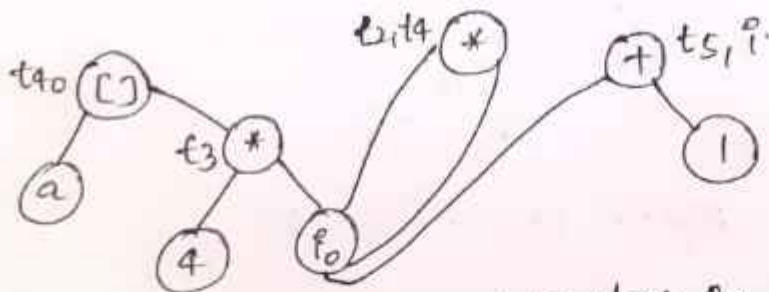
Fig: DAG for two statements in block B3.



Fig(b): DAG for three statements in block B3.



Fig(d): DAG for four statements in block B3.



Fig(e): DAG for complete block B3.

==

Principle Sources of optimization.

A transformation of a program is called local if it is applied within a basic block and global if applied across basic blocks.

* There are different types of transformations to improve the code and these transformations depend on the kind of optimization required.

① Function - preserving transformations are those transformations that are performed without changing the function it computes.
* These are primarily used when global optimizations are performed.

② Structure - preserving transformations are those that are performed without changing the set of expressions computed by the block.

* many of these transformations are applied locally.

③ Algebraic transformations are used to simplify the computation of expression set using algebraic identities.

* These can replace expensive operations by cheaper ones, for instance multiplication by 2 can be replaced by left shift.

Function preserving Transformations.

The following techniques are function-preserving transformations

1) common sub-expression Elimination

2) copy propagation.

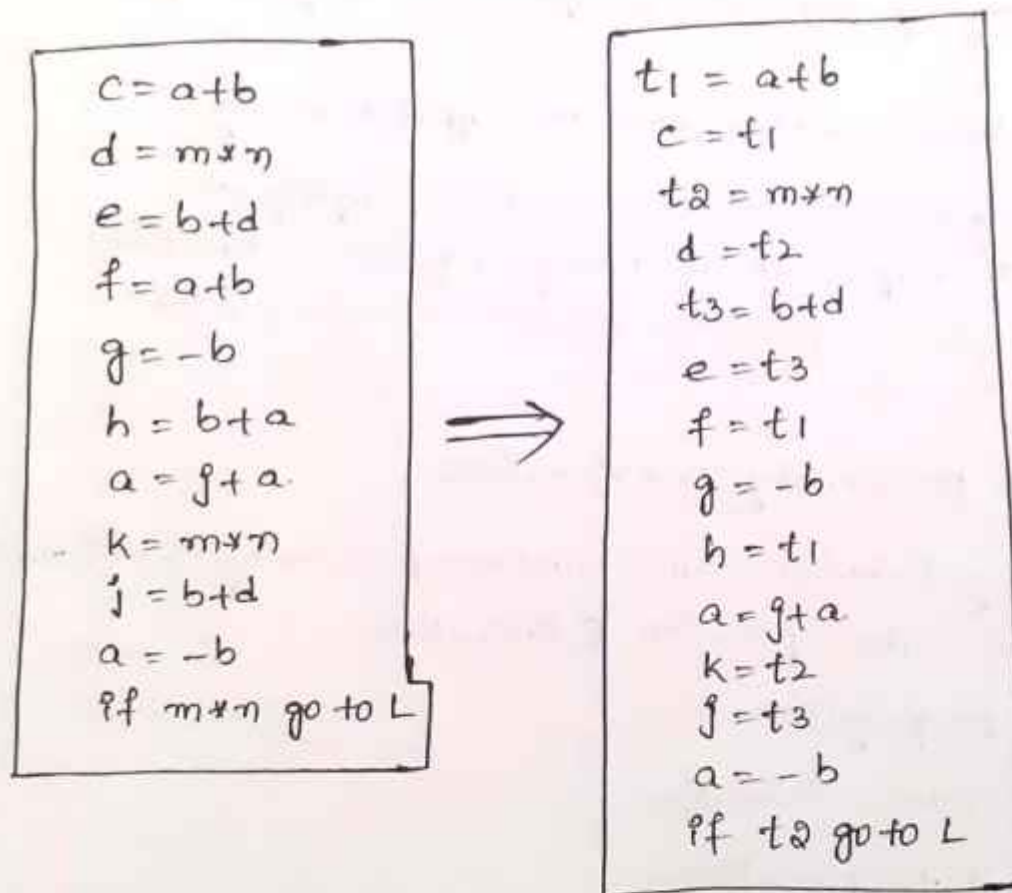
3) Dead code elimination

4) constant propagation.

(1) common sub-expression Elimination :-

An expression E is said to be common sub expression if E is computed before and the variables in the expression are not modified since its ~~comp~~ computation. If such expression is present then the re-computation can be avoided by using the result of the previous computation.

- * This technique can be applied both locally and globally
- * We need to maintain a table to store the details of expressions evaluated so far and use this information to identify and eliminate the re-computation of the same expression.
- * The common sub-expression elimination can be done within basic block by analyzing and storing the information of expression in the table until the operands in the expression are redefined.
- * If any operand in the expression is redefined, then remove the expression from the table.



Program code before and after common sub-expression elimination

- * The above figure shows the optimized code on applying common sub-expression elimination technique locally.

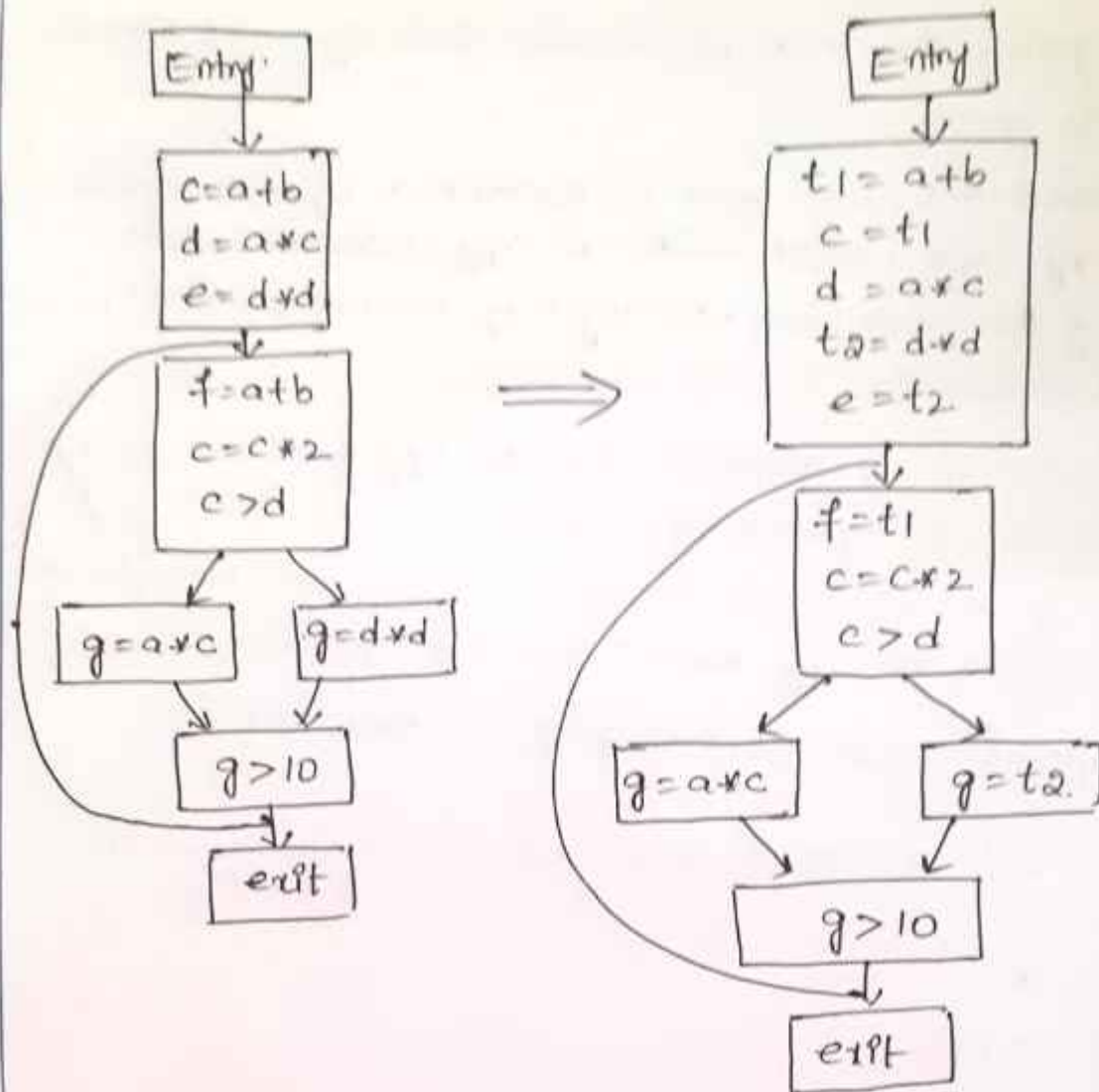


Fig: Example for global common Sub Expression Elimination.

* The above example shows the common-Sub expression elimination globally by replacing the evaluated expression with a new temporary variable t_1 and t_2 and these variables are used whenever the same expression is used.

(ii) Copy propagation :-

A copy statement is a statement that is in the form $a = b$.

* copy propagation technique is applied to replace the later use of variable 'a' with the use of 'b' if original values of 'a' do not change after this assignment.

* This technique can be applied locally and globally.

- * This optimization reduces runtime stack size and improves execution speed.
- * To implement this, we need to maintain a separate table called copy table, which holds all copy statements. While traversing the basic block, if any copy statement is found, add this information into the copy table.
- * While processing any statement, check the copy table for variable a , which can be replaced variable b .
- * If there is an assignment statement that computes the value of a then remove the copy statement from the copy table.
- * Copy propagation helps in identifying the dead code.

Example:- Let the basic block contain the following set of statements.

$Y = X$
 $Z = Y + 1$
 $W = Y$
 $Y = W + Z$
 $Y = W$

Statement NO	Instruction	Updated Instruction	copy table content
1.	$Y = X$	$Y = X$	$\{(Y, X)\}$
2.	$Z = Y + 1$	$Z = X + 1$	$\{(Y, X)\}$
3.	$W = Y$	$W = X$	$\{(Y, X), (W, X)\}$
4.	$Y = W + Z$	$Y = X + Z$	$\{(W, X)\}$
5.	$Y = W$	$Y = W$	$\{(W, X), (Y, X)\}$

- * On the first statement 1, since it is a copy statement, the information is added in to the copy table.
- * Statement 2 uses the value X indirectly from Y , hence $GET(Y, table)$ could return X and the statement is modified

Replacing Y with X.

- * The third statement is a copy statement and since Y is to replace X, we insert into the copy table W to be replaced by X.
- * When statement 4 is processed, Y value is defined, and hence, details of Y are removed from the copy table.
- * The fifth statement is a copy statement and is inserted into the copy table.

(iii) Dead code Elimination:-

Code that is unreachable or does not affect the program is said to be dead code. Such code requires unnecessary CPU time, which can be identified and eliminated using this technique.

Example program:-

```
1) int Var1;  
2) void Sample()  
3) {  
4)     int i;  
5)     i = 1; /* dead store */  
6)     Var1 = 1; /* dead store */  
7)     Var2 = 2;  
8)     return;  
9)     Var1 = 3; /* unreachable */  
10) }
```

The code fragment after dead code elimination

```
int Var1;  
void Sample()  
{  
    Var1 = 2;  
    return;  
}
```

→ Here the value of i is never used, and the value assigned to Var1 variable in statement 6 is dead store and the statement 9 is unreachable. These statements can be eliminated as they do not affect the program execution.

(iv) Constant propagation:-

Constant propagation is an approach that propagates the constant value assigned to a variable at the place of its use.

For example, given an assignment statement $x = c$, where c is a constant, replace later uses of x with uses of c , provided there are no intervening assignments to x .

This approach is similar to copy propagation and is applied at first stage of optimization. This method can analyze by propagating constant value in conditional statement, to determine whether a branch should be executed or not, that is, identifies the dead code.

Example: Let us consider the following example:

```
pi = 22/7
```

```
void area_per(int r)
```

```
{
```

```
    float area, perimeter;
```

```
    area = pi * r * r;
```

```
    perimeter = 2 * pi * r;
```

```
    print area, perimeter;
```

```
}
```

In this example, we can notice some simple constant propagation results, which are as follows.

→ In line 1 the variable π is constant and this value is the result of $22/7$, which can be computed at compile time and has the value 3.143.

→ In line 5 the variable π can be replaced with the constant value 3.143.

→ In line 6 since the value of π is constant, the partial result of the statement can be computed and the statement can be modified as $\text{perimeter} = 6.285 * r$;

Loop optimization Techniques.

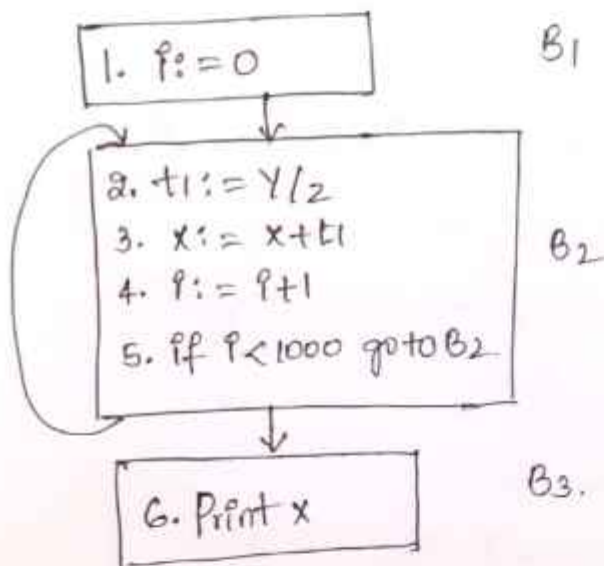
8

In most of the programs, 90% of execution time is inside the loops; hence loop optimization is the most valuable machine-independent optimization and are good candidates for code improvement.

Let us consider the example of a program that has a loop statement:

```
for (int i=0; i<1000; i++)  
    x = x + y/i;  
Print x;
```

The optimized code of this program is represented in flow graph as



Flow graph After code optimization.

The total number of statements that are executed are as follows.

<u>Statement Number</u>	<u>Frequency of Execution</u>	<u>Total No. of Statements</u>
1	1	1
2-5	1000	4000
6	1	1

on the whole, the number of statements that are executed are 4002. Instead if the code is written as follows then the total number of statements that are executed are only 1002.

1. $t_1 = y/2$
2. $x = x + t_1$
3. $x = x + t_1$
4. $x = x + t_1$
5. ...
6. ...
1001. $x = x + t_1$
1002. Print x .

The execution is three times better for later than the first form, but the space requirements is more. Inefficient code is generated in the loops for various reasons like

- Induction Variable usage to keep track of iteration
- unnecessary computations made inside the iterative loops that are not effective.
- use of high strength operators inside the loop.

The important loop optimizations are

- ① Elimination of loop invariant computations.
- ② Strength reduction
- ③ code motion
- ④ Elimination of induction variables.

(1) A Loop Invariant computation :-

A loop invariant computation is one that evaluates the same value every time the statements in loop are executed. * moving such loops ~~out~~ ^{outside} computational statements outside the loop leads to a reduction in the execution time.

Algorithm for elimination of loop invariant code.

Step 1: Identify loop-invariant code.

Step 2: An instruction is loop-invariant if, for each operand:

- (i) The operand is constant, OR
- (ii) All definitions for all the operands that reach the instruction inside the loop are made outside the loop, OR
- (iii) There is exactly one definition made using loop invariant variables inside the loop for an operand that reaches the instruction.

Step 3: move it outside the loop

- (1) move each instruction i to newly created preheader if and only if it satisfies these requirements of the loop.

Example:

```
for (int i = 0; i < 1000; i++)  
    x = x + y/z;  
Print x;
```

In the above example, the value y and z remains unchanged as there is no definition for these variables in the loop. Hence, the computation of y/z remains unchanged, which can be moved above the loop and the same code can be rewritten as follows.

```
t1 = y/z;  
for (int i = 0; i < 1000; i++)  
    x = x + t1;  
Print x;
```


(2) Induction Variables :-

Induction Variables are those variables used in a loop, their values are in lock-step, and hence, it may be possible to eliminate all except one.

* There are two types of induction variables - basic and derived.

* Basic induction variables are those that are of the form

$$I = I \pm c$$

where I is loop variable and c is some constant.

* Derived induction variables are those that are defined only once in the loop and their value is linear function of the basic induction variable.

for example :-

$$J = A * I + B$$

Here J is a variable that is dependent on the basic induction variable I and the constants A and B . This is represented as a triplet (I, A, B)

* Induction variable elimination involves three steps.

1) Detecting induction variable.

2) Reducing the strength of induction variable.

3) Eliminating induction variable.

Induction Variable Elimination :-

Some loops contain two or more induction variables that can be combined into one induction variable.

Example:- The code fragment below has three induction variables (i1, i2 and i3) that can be replaced with one induction variable, thus eliminating two induction variables. 10

```
int a[SIZE];
int b[SIZE];
void f(void)
{
    int i1, i2, i3;
    for (i1=0, i2=0, i3=0; i1 < SIZE; i1++)
        a[i2++] = b[i3++];
    return;
```

After induction variable elimination the code is rewritten as

```
int a[SIZE]
int b[SIZE]
void f(void)
{
    int i1;
    for (i1=0; i1 < SIZE; i1++)
        a[i1] = b[i1];
    return;
```

* Induction variable elimination can reduce the number of additions (or subtractions) in a loop, and improve both runtime performance and code space.

Machine - Dependent Optimization.

This optimization can be applied on target machine instructions. This includes register allocation, use of addressing modes, and peep hole optimization.

* Instructions involving register operands are faster and shorter; hence, if we make use of more registers during target code generation, efficient code will be generated.

Peephole optimization:-

Generally code generation algorithms produce code, statement by statement. This may contain redundant instructions and suboptimal constructs. The efficiency of such code can be improved by applying peephole optimization, which is simple but effective optimization on target code.

- * The peephole is considered a small moving window on the target code. The code in peephole need not be contiguous. It improves the performance of the target program by examining and transforming a short sequence of target instructions.
- * The advantage of peephole optimization is that each improvement applied increases opportunities and shows additional improvements. It may need repeated passes to be applied over the target code to get the maximum benefit.

(i) Redundant loads and stores

The code generation algorithm produces the target code, which is either represented with single operand or two operands or three operands.

- * Let us assume the instructions are with two operands. The following is an example that gives the assembly code for the statement

$x = y + z.$

1. MOV y, R_0

2. ADD z, R_0

3. MOV R_0, x

Instruction 1 moves the value of y to Register R_0 , second instruction performs the addition of value in z with the register content and the result of the operation is stored in the register.

- * The third instruction copies the register content to the location x . At this point the value of x is available in both location of x and the register R_0 .

If the above algorithm is applied on the code $a = b + c, d = a + c$ then it generates the code given below.

1. MOV b, R₀
2. ADD C, R₀
3. mov R₀, a
4. mov a, R₀
5. ADD c, R₀
6. mov R₀, d

Here we can say that 3 and 4 are redundant load and store instructions. These instructions will not affect the values before or after their execution. Such redundant statements can be eliminated and the resultant code is as follows.

1. mov b, R₀
2. ADD C, R₀
3. ADD c, R₀
4. mov R₀, d

(ii) Algebraic Simplification:-

There are few algebraic identities that occur frequently enough and are worth considering. Look at the following statements.

$$x := x + 0$$
$$x := x + 1$$

They do not alter the value of x . If we keep them as it is, later when code generation algorithm is applied on it, it may produce silly statements that are of no use. Hence, such statements whether they are in three address code or target code can be removed.

(iii) Deadcode Elimination:-

Removal of unreachable code is an opportunity for peephole optimization. A statement immediately after an

unconditional jump or a statement that never get a chance to be executed can be identified and eliminated. Such code is called the dead code.

Ex: // define $x=0$

```
...  
if(x)  
{  
    .. print value.  
}
```

If this is translated to target code as

```
If x=1 goto L1  
      goto L2.
```

L1: print value.

L2: ----

here, value will be never printed. So whatever code inside the body of "if(x)" is dead code; hence it can be removed.

(iv) Reduction in Strength:-

This optimization mainly deals with replacing expensive operations by cheaper ones for example.

→ $x^2 \Rightarrow x * x$

→ fixed-point multiplication and division by a power of 2.

⇒ shift

→ floating point division by a constant ⇒ floating-point multiplication by a constant.

⇒