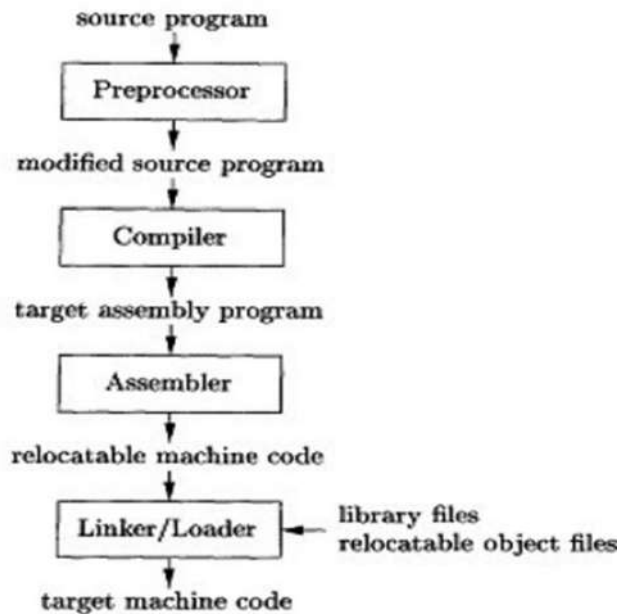


Translator:

It is a program that translates one language to another Language. Examples of translator are compiler, assembler, interpreter, linker, loader and preprocessor.

**1. INTRODUCTION TO LANGUAGE PROCESSING**

The Language Processing System can be represented as shown figure below.



PRE-PROCESSOR: Preprocessor is a program that processes its input data to produce output that is used as input to another program. The output is said to be a preprocessed form of the input data, which is often used by **compilers**.

Macro processing (Macro Expansion): A pre-processor may allow a user to define macros that are short hands for longer constructs.

There are two types of macros, object-like and function-like. Object-like macros do not take parameters; function-like macros do (although the list of parameters may be empty). The generic syntax for declaring an identifier as a macro of each type is, respectively.

```

#define <identifier> <replacement token list>           // object-like macro

#define <identifier>(<parameter list>) <replacement token list>
                                     // function-like macro, note parameters
  
```

Whenever the identifier or function call appears in the source code it is replaced with the replacement token list. This process is known as macro expansion.

File inclusion: A pre-processor may include header files into the program text.

COMPILER:

Compiler is a translator, program that translates a program written in (HLL) - the source program and translate it into an equivalent program in (MLL) - the target program. An important part of the compiler is to show errors to the programmer.



ASSEMBLER:

Programmers found it difficult to write or read programs in machine language. They begin to use a mnemonic (symbols) for each machine instruction, which they would subsequently translate into machine language. Such a mnemonic machine language is now called an assembly language.

Programs known as assembler were written to automate the translation of assembly language into machine language. The input to an assembler program is called source program and the output is a relocatable machine code (object program). Relocatable machine code is the code that can be located (and executed) from anywhere in memory, and is also called position independent code. **i.e.** The assembler produces machine code and stores it wherever it finds an empty space.



LOADER AND LINK-EDITOR (Linker):

In high level languages, some built in header files or libraries are stored. These libraries are predefined and these contain basic functions which are essential for executing the program. These functions are linked to the libraries by a program called Linker. Usually a longer program is divided into smaller subprograms called modules. And these modules must be combined to execute the program. The process of combining the modules is done by the linker.

Loader is responsible for loading programs and libraries. It is one of the essential stages in the process of executing a program, as it places programs into memory and prepares them for execution. The final output is absolute machine code that is loaded into a computer fixed storage location and may not be relocated during execution of a computer program.

Examples of Compilers:

1. Ada compilers
2. ALGOL compilers
3. BASIC compilers
4. C# compilers
5. C compilers
6. C++ compilers
7. COBOL compilers
8. Common Lisp compilers
9. Fortran compilers
10. Java compilers
11. Pascal compilers
12. PL/I compilers
13. Python compilers

Cousins of compiler

Preprocessor, interpreter, assembler, linker and loader are called as cousins of compiler.

INTERPRETER:

It is one of the translators that translate high level language to low level language. Difference between Compiler & interpreter is given below.

No	Compiler	Interpreter
1	Compiler Takes Entire program as input	Interpreter Takes Single instruction as input.
2	Intermediate Object Code is Generated	No Intermediate Object Code is Generated
3	Conditional Control Statements are	Conditional Control Statements are

	Executed faster	Executed slower.
4	Memory Requirement : More(Since Object Code is Generated)	Memory Requirement is Less
5	Errors are displayed after entire program is checked. Hence, debugging is difficult as compared to interpreter.	Errors are displayed for every instruction interpreted. Hence, Debugging is easier.
6	Example : C Compiler	Example : BASIC

2. STRUCTURE OF THE COMPILER DESIGN OR PHASES OF A C O M P I L E R

A compiler operates in phases. A phase is a *logically* interrelated operation that takes source program in one representation and produces output in another representation. The phases of a compiler are shown in below

There are two phases of compilation.

Analysis Phase **or** Front-end (*Machine Independent/Language Dependent*)

Synthesis Phase **or** Back-end (*Machine Dependent/Language independent*)

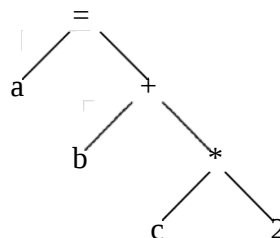
Compilation process is partitioned into no. of sub processes called '*phases*'.

Lexical Analysis:-

- It is also called as Scanner, because it reads the program.
- It is the first phase of the compiler. It gets input from the source program and produces tokens as output.
- It reads the characters one by one, starting from left to right and forms the tokens.
- **Token :**
- It represents a logically cohesive sequence of characters such as keywords, operators, identifiers, special symbols etc.
- **Example:** $a + b = 20$
 - ✓ Here, a,b,+,=,20 are all separate tokens.
 - ✓ Group of characters forming a token is called the **Lexeme**.
 - ✓ The lexical analyser not only generates a token but also enters the lexeme into the symbol table if it is not already there.

Syntax Analysis:-

- It is the second phase of the compiler. It is also known as parser.
- It gets the token stream as input from the lexical analyser of the compiler and generates syntax tree as the output.
- Syntax tree: It is a tree in which interior nodes are operators and exterior nodes are operands.
- Example: For $a=b+c*2$, syntax tree is



Semantic Analysis:-

- It is the third phase of the compiler.
- It gets input from the syntax analysis as parse tree and checks whether the given syntax is correct or not.
- It performs type conversion of all the data types into real data types.

Intermediate Code Generations:-

- It is the fourth phase of the compiler.
- It gets input from the semantic analysis and converts the input into output as intermediate code such as three-address code.
- The three-address code consists of a sequence of instructions, each of which has at most three operands.

Example: $t1 = t2 + t3$

Code Optimization:-

- It is the fifth phase of the compiler.
- It gets the intermediate code as input and produces optimized intermediate code as output.
- This phase reduces the redundant code and attempts to improve the intermediate code so that faster-running machine code will result.
- During the code optimization, the result of the program is not affected.
- To improve the code generation, the optimization involves
 - deduction and removal of dead code (unreachable code).
 - calculation of constants in expressions and terms.
 - collapsing of repeated expression into temporary string.
 - loop unrolling.
 - moving code outside the loop.

removal of unwanted temporary variables.

Code Generation:-

- It is the final phase of the compiler.
- It gets input from code optimization phase and produces the target code or object code as result.
- Intermediate instructions are translated into a sequence of machine instructions that perform the same task.
- The code generation involves
 - allocation of register and memory
 - generation of correct references
 - generation of correct data types
 - generation of missing code.

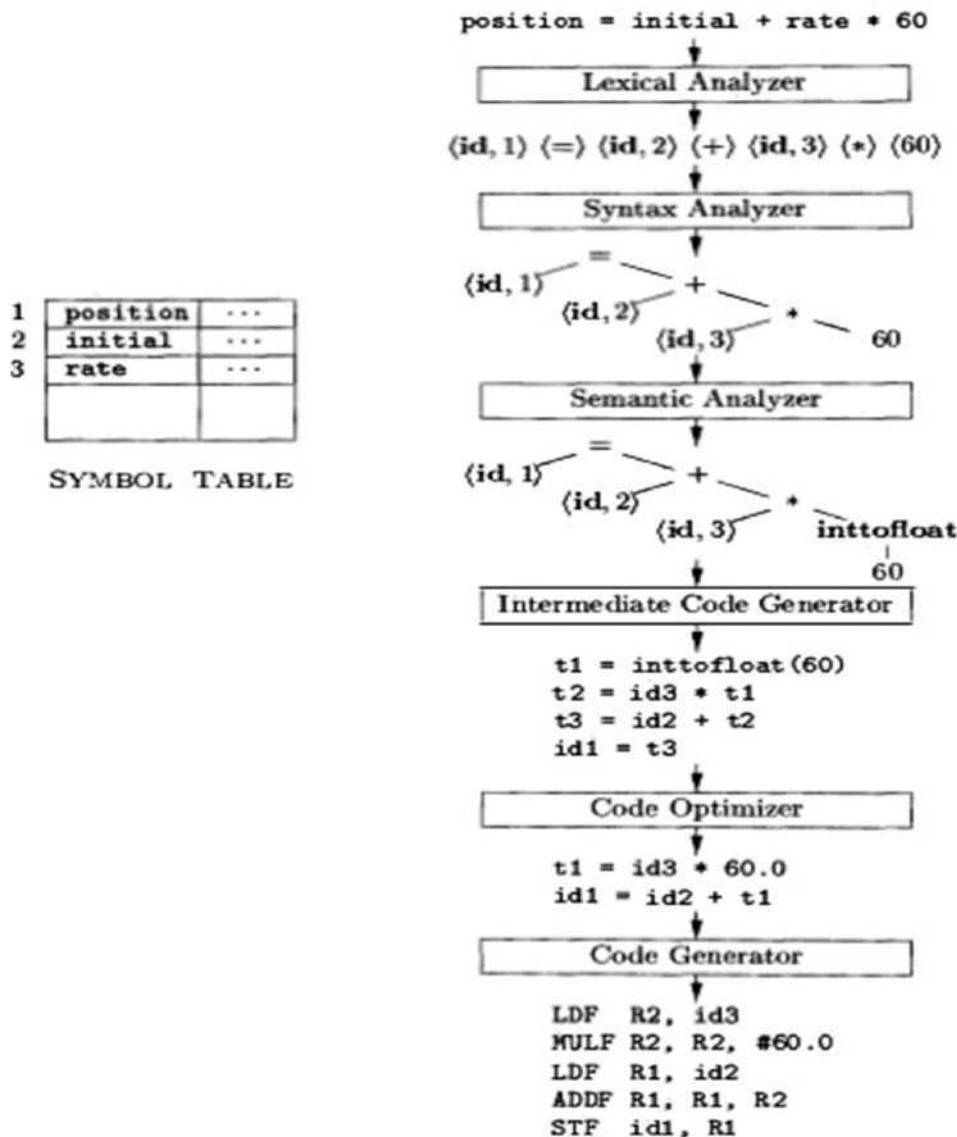
Symbol Table Management (or) Book-keeping:-

- Symbol table is used to store all the information about identifiers used in the program.
- It is a data structure containing a record for each identifier, with fields for the attributes of the identifier.
- It allows to find the record for each identifier quickly and to store or retrieve data from that record.
- Whenever an identifier is detected in any of the phases, it is stored in the symbol table.

Error Handlers:-

- Each phase can encounter errors. After detecting an error, a phase must handle the error so that compilation can proceed.
- In lexical analysis, errors occur in separation of tokens.
- In syntax analysis, errors occur during construction of syntax tree.
- In semantic analysis, errors occur when the compiler detects constructs with right syntactic structure but no meaning and during type conversion.
- In code optimization, errors occur when the result is affected by the optimization.
- In code generation, it shows error when code is missing etc.

Example: To illustrate the translation of source code through each phase, consider the statement *position := initial + rate * 60*. The figure shows the representation of this statement after each phase.



Front end and Back end of a compiler:

Compiler can be grouped into front and back ends:

Front end: analysis

These normally include lexical and syntax analysis, semantic analysis, the creation of the symbol table and the generation of intermediate code. It also includes error handling that goes along with each of these phases. The phases which are dependent on source language and independent of target machine are grouped together in front end.

Back end: synthesis

It includes code optimization phase and code generation along with the necessary error handling and symbol table operations. The phases which are independent of source language and dependent on target language/ target machine are grouped together into back end.

Advantages of frontend & back end:

- By keeping the same front end and attaching different back ends one can produce compiler for same source language on different machines.
- By keeping different front ends and same back end one can compile several different languages on the same machine.

Compiler passes

A collection of phases is done only once (single pass) or multiple times (multi pass)

- **Single** pass: usually requires everything to be defined before being used in source program.
- **Multi** pass: compiler may have to keep entire program representation in memory.

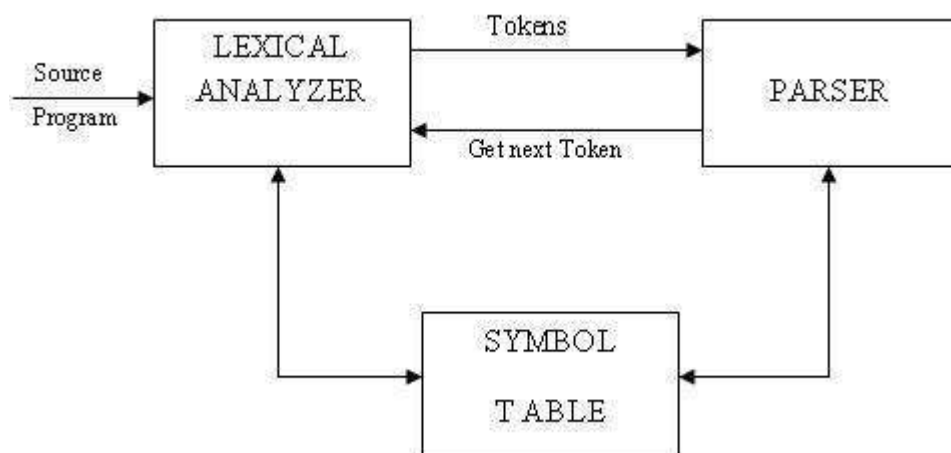
Several phases can be grouped into one single pass and the activities of these phases are interleaved during the pass. For example, lexical analysis, syntax analysis, semantic analysis and intermediate code generation might be grouped into one pass.

3) LEXICAL ANALYSIS

Lexical analysis is the process of converting a sequence of characters into a sequence of tokens. A program or function which performs lexical analysis is called a lexical analyzer or scanner. A lexer often exists as a single function which is called by a parser or another function.

The Role Of The Lexical Analyzer:

- The lexical analyzer is the first phase of a compiler.
- Its main task is to read the input characters and produce as output a sequence of tokens that the parser uses for syntax analysis.



- Upon receiving a “get next token” command from the parser, the lexical analyzer reads input characters until it can identify the next token.

Why Lexical Analyzer is separated from Syntax Analyzer?

- To make the design of compiler simple.
- To improve the efficiency of the compiler.
- To enhance the computer portability.

Token, Pattern, Lexeme:

Tokens

A token is a string of characters, categorized according to the rules as a symbol (e.g., IDENTIFIER, NUMBER, COMMA). Token is the class or category of input string. It defines the type of lexeme. The process of forming tokens from an input stream of characters is called **tokenization**. A token can look like anything that is useful for processing an input text stream or text file. Consider this expression in the C programming language: `sum=3+2;`

Lexeme	Token type
sum	Identifier
=	Assignment operator
3	Number
+	Addition operator
2	Number
;	End of statement

Lexeme:

Collection or group of characters forming tokens is called Lexeme. Lexeme follows a pattern to become a token.

Pattern:

A pattern is a description of the form that the lexemes of a token may take.

In the case of a keyword as a token, the pattern is just the sequence of characters that form the keyword. For identifiers and some other tokens, the pattern is a more complex structure that is matched by many strings.

Example is given below.

Token	lexeme	pattern
const	const	const
if	if	if
relation	<, <=, =, < >, >=, >	< or <= or = or < > or >= or letter followed by letters & digit
i	pi	any numeric constant
nun	3.14	any character b/w "and "except"
literal	"core"	pattern

Attributes for Tokens

Some tokens have attributes that can be passed back to the parser. The lexical analyzer collects information about tokens into their associated attributes. The attributes influence the translation of tokens.

- i) Constant : value of the constant
- ii) Identifiers: pointer to the corresponding symbol table entry.

Error Recovery Strategies In Lexical Analysis:

The following are the error-recovery actions in lexical analysis:

- 1) Deleting an extraneous character.

2) Inserting a missing character.

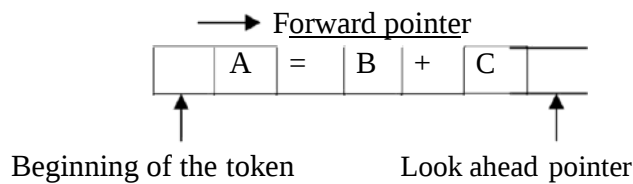
3) Replacing an incorrect character by a correct character.

4) Transforming two adjacent characters.

5) **Panic mode recovery:** Deletion of successive characters from the token until **error** is resolved.

INPUT BUFFERING

The scanner look one or more characters beyond the next lexeme before we can be sure we have the right lexeme. As characters are read from left to right, each character is stored in the buffer to form a meaningful token as shown below:



Scanner uses a two-buffer scheme that handles large look-aheads safely. We then consider an improvement involving "sentinels" that saves time checking for the ends of buffers.

Buffer Pairs

- A buffer is divided into two N-character halves, as shown below

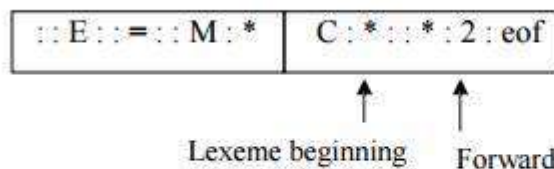


Fig. 1.9 An input buffer in two halves

- Note: Expression for $E=MC^2$ is written according to FORTRAN language.
- Each buffer is of the same size N, and N is usually the number of characters on one disk block. E.g., 1024 or 4096 bytes.
- Using one system read command we can read N characters into a buffer.
- If fewer than N characters remain in the input file, then a special character, represented by **eof**, marks the end of the source file.
- Two pointers to the input are maintained:
 1. Pointer **lexeme_beginning**, marks the beginning of the current lexeme, whose extent we are attempting to determine.
 2. Pointer **forward** scans ahead until a pattern match is found.
Once the next lexeme is determined, forward is set to the character at its right end.
- The string of characters between the two pointers is the current lexeme. After the lexeme is recorded as an attribute value of a token returned to the parser, lexeme_beginning is set to the character immediately after the lexeme just found.

Advancing forward pointer:

Advancing forward pointer requires that we first test whether we have reached the end of one of the buffers, and if so, we must reload the other buffer from the input, and move forward to the beginning of the newly loaded buffer. If the end of second buffer is reached, we must again reload the first buffer with input and the pointer wraps to the beginning of the buffer.

Code to advance forward pointer:

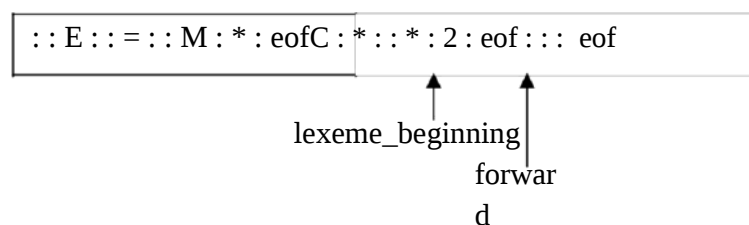
```

if forward at end of first half then begin
    reload second half;
    forward := forward + 1
end
else if forward at end of second half then
    begin reload second half;
    move forward to beginning of first half
end
else forward := forward + 1;

```

Sentinels

- For each character read, we make two tests: one for the end of the buffer, and one to determine what character is read. We can combine the buffer-end test with the test for the current character if we extend each buffer to hold a sentinel character at the end.
- The sentinel is a special character that cannot be part of the source program, and a natural choice is the character eof.
- The sentinel arrangement is as shown below:



Note that eof retains its use as a marker for the end of the entire input. Any eof that appears other than at the end of a buffer means that the input is at an end.

Code to advance forward pointer:

```

forward := forward + 1;
if forward ↑ = eof then begin
    if forward at end of first half then begin
        reload second half;
        forward := forward + 1
    end
    else if forward at end of second half then begin
        reload first half;
        move forward to beginning of first
        half end
    else /* eof within a buffer signifying end of input
        */ terminate lexical analysis
    end
end

```

SPECIFICATION OF TOKENS

There are 3 specifications of tokens:

- 1) Strings
- 2) Language
- 3) Regular expression

Strings and Languages

An **alphabet** or character class is a finite set of symbols.

A **string** over an alphabet is a finite sequence of symbols drawn from that alphabet. A

language is any countable set of strings over some fixed alphabet.

In language theory, the terms "sentence" and "word" are often used as synonyms for "string." The length of a string s , usually written $|s|$, is the number of occurrences of symbols in s . For example, banana is a string of length six. The empty string, denoted ϵ , is the string of length zero.

Operations on strings

The following string-related terms are commonly used:

1. A **prefix** of string s is any string obtained by removing zero or more symbols from the end of string s . For example, ban is a prefix of banana.
2. A **suffix** of string s is any string obtained by removing zero or more symbols from the beginning of s .
For example, nana is a suffix of banana.
3. A **substring** of s is obtained by deleting any prefix and any suffix from s .
For example, nan is a substring of banana.
4. The **proper prefixes, suffixes, and substrings** of a string s are those prefixes, suffixes, and substrings, respectively of s that are not ϵ or not equal to s itself.
5. A subsequence of s is any string formed by deleting zero or more not necessarily consecutive positions of s . For example, baan is a subsequence of banana.

Operations on languages:

The following are the operations that can be applied to languages: 1.Union

2.Concatenation

3.Kleene closure

4.Positive closure

The following example shows the operations on strings:

Let $L=\{0,1\}$ and $S=\{a,b,c\}$

1. Union : $L \cup S = \{0,1,a,b,c\}$
2. Concatenation : $L.S = \{0a,1a,0b,1b,0c,1c\}$
3. Kleene closure : $L^* = \{\epsilon, 0, 1, 00, \dots\}$

4. Positive closure : $L^+ = \{0,1,00,\dots\}$

Regular Expressions

Each regular expression r denotes a language $L(r)$.

Here are the rules that define the regular expressions over some alphabet Σ and the languages that those expressions denote:

1. ϵ is a regular expression, and $L(\epsilon)$ is $\{\epsilon\}$, that is, the language whose sole member is the empty string.
2. If 'a' is a symbol in Σ , then 'a' is a regular expression, and $L(a) = \{a\}$, that is, the language with one string, of length one, with 'a' in its one position.
3. Suppose r and s are regular expressions denoting the languages $L(r)$ and $L(s)$. Then,
 - a) $(r)|(s)$ is a regular expression denoting the language $L(r) \cup L(s)$.
 - b) $(r)(s)$ is a regular expression denoting the language $L(r)L(s)$.
 - c) $(r)^*$ is a regular expression denoting $(L(r))^*$.
 - d) (r) is a regular expression denoting $L(r)$.
4. The unary operator $*$ has highest precedence and is left associative.
5. Concatenation has second highest precedence and is left associative.
6. $|$ has lowest precedence and is left associative.

Regular set

A language that can be defined by a regular expression is called a regular set. If two regular expressions r and s denote the same regular set, we say they are equivalent and write $r = s$.

There are a number of algebraic laws for regular expressions that can be used to manipulate into equivalent forms.

For instance, $r | s = s | r$ is commutative; $r | (s | t) = (r | s) | t$ is associative.

Regular Definitions

Giving names to regular expressions is referred to as a Regular definition. If Σ is an alphabet of basic symbols, then a regular definition is a sequence of definitions of the form

$d_1 \rightarrow r_1$
 $d_2 \rightarrow r_2$
 $\dots\dots\dots$
 $d_n \rightarrow r_n$

1. Each d_i is a distinct name.
2. Each r_i is a regular expression over the alphabet $\Sigma \cup \{d_1, d_2, \dots, d_{i-1}\}$.

Example: Identifiers is the set of strings of letters and digits beginning with a letter. Regular definition for this set:

$\text{letter} \rightarrow A | B | \dots | Z | a | b | \dots | z$
 $\text{digit} \rightarrow 0 | 1 | \dots | 9$
 $\text{id} \rightarrow \text{letter} (\text{letter} | \text{digit}) ^*$

Shorthands

Certain constructs occur so frequently in regular expressions that it is convenient to introduce notational shorthands for them.

1. One or more instances (+):

- The unary postfix operator $+$ means “one or more instances of”.
- If r is a regular expression that denotes the language $L(r)$, then $(r)^+$ is a regular expression that denotes the language $(L(r))^+$.
- Thus the regular expression a^+ denotes the set of all strings of one or more a 's.
- The operator $+$ has the same precedence and associativity as the operator $*$.

2. Zero or one instance (?):

- The unary postfix operator $?$ Means “zero or one instance of”.
- The notation $r?$ is a shorthand for $r \mid \epsilon$.
- If r is a regular expression, then $(r)?$ is a regular expression that denotes the language $L(r) \cup \{\epsilon\}$.

3. Character Classes:

- The notation $[abc]$ where a , b and c are alphabet symbols denotes the regular expression $a \mid b \mid c$.
- Character class such as $[a-z]$ denotes the regular expression $a \mid b \mid c \mid d \mid \dots \mid z$.
- We can describe identifiers as being strings generated by the regular expression, $[A-Za-z][A-Za-z0-9]^*$

Non-regular Set

A language which cannot be described by any regular expression is a non-regular set. Example: The set of all strings of balanced parentheses and repeating strings cannot be described by a regular expression. This set can be specified by a context-free grammar.

RECOGNITION OF TOKENS

Consider the following grammar fragment:

```
stmt → if expr then stmt
      | if expr then stmt else stmt
      | ε
```

```
expr → term relop term
      | term
```

```
term → id |
       num
```

where the terminals `if`, `then`, `else`, `relop`, `id` and `num` generate sets of strings given by the

following regular definitions:

digit $\rightarrow$ [0-9]

digits $\rightarrow$ digit(digit)⁺

if $\rightarrow$ if

then $\rightarrow$ then

else $\rightarrow$ else

relop $\rightarrow$ < | <= | = | <> | > | >=

id $\rightarrow$ letter(letter | digits)*

num $\rightarrow$ digit⁺ (.digit⁺)?(E(+|-)?digit⁺)?

For this language fragment the lexical analyzer will recognize the keywords if, then, else, as well as the lexemes denoted by relop, id, and num. To simplify matters, we assume keywords are reserved; that is, they cannot be used as identifiers.

Transition diagrams

It is a diagrammatic representation to depict the action that will take place when a lexical analyzer is called by the parser to get the next token. It is used to keep track of information about the characters that are seen as the forward pointer scans the input.

Transition diagram for relational operators

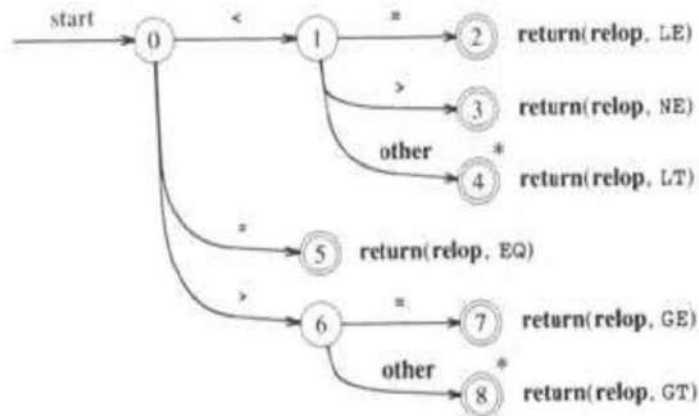
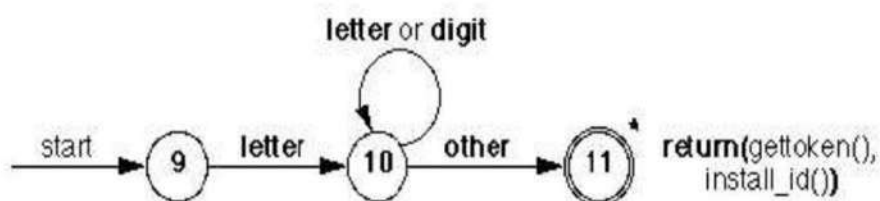


Fig. 3.12. Transition diagram for relational operators.

Transition diagram for identifiers and keywords

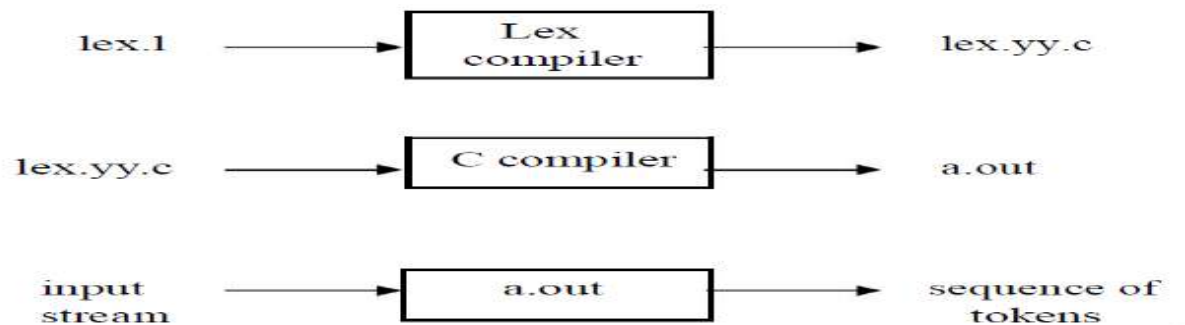


LEX

Lex is a computer program that generates lexical analyzers. Lex is commonly used with the YACC parser generator.

Creating a lexical analyzer

- First, a specification of a lexical analyzer is prepared by creating a program `lex.l` in the Lex language. Then, `lex.l` is run through the Lex compiler to produce a C program `lex.yy.c`.
- Finally, `lex.yy.c` is run through the C compiler to produce an object program `a.out`, which is the lexical analyzer that transforms an input stream into a sequence of tokens.



Lex Specification

A Lex program consists of three parts:

```

{ definitions }
%%
{ rules }
%%
{ user subroutines }
  
```

- **Definitions** include declarations of variables, constants, and regular definitions
- **Rules** are statements of the form


```

p1 {action1}
p2 {action2}
...
pn {actionn}
      
```

 where p_i is regular expression and $action_i$ describes what action the lexical analyzer should take when pattern p_i matches a lexeme. Actions are written in C code.
- **User subroutines** are auxiliary procedures needed by the actions. These can be compiled separately and loaded with the lexical analyzer.

YACC- Yet Another Compiler-Compiler

Yacc provides a general tool for describing the input to a computer program. The Yacc user specifies the structures of his input, together with code to be invoked as each such structure is recognized. Yacc turns such a specification into a subroutine that handles the input process; frequently, it is convenient and appropriate to have most of the flow of control in the user's application handled by this subroutine.

FINITE AUTOMATA

Finite Automata is one of the mathematical models that consist of a number of states and edges. It is a transition diagram that recognizes a regular expression or grammar.

Types of Finite Automata

There are two types of Finite Automata :

- Non-deterministic Finite Automata (NFA)
- Deterministic Finite Automata (DFA)

Non-deterministic Finite Automata

NFA is a mathematical model that consists of five tuples denoted by

$M = \{Q_n, \Sigma, \delta, q_0, f_n\}$

Q_n – finite set of states

Σ – finite set of input symbols

δ – transition function that maps state-symbol pairs to set of states

q_0 – starting state

f_n – final state

Deterministic Finite Automata

DFA is a special case of a NFA in which

- i) no state has an ϵ -transition.
- ii) there is at most one transition from each state on any input.

DFA has five tuples denoted by

$M = \{Q_d, \Sigma, \delta, q_0, f_d\}$

Q_d – finite set of states

Σ – finite set of input symbols

δ – transition function that maps state-symbol pairs to set of states

q_0 – starting state

f_d – final state

Construction of DFA from regular expression

The following steps are involved in the construction of DFA from regular expression:

- i) Convert RE to NFA using Thomson's rules
- ii) Convert NFA to DFA
- iii) Construct minimized DFA

UNIT –I (questions from previous year question papers)

- 1) Define pre-processor. What are the functions of pre-processor?
- 2) Briefly describe about the Lexical errors.
- 3) Explain the role of lexical analyzer in compiler construction process.
- 4) Differentiate between lexical analysis and parsing.
- 5) What are pass and phase? Write the differences between a pass and a phase in the context of compiler construction.
- 6) Write a LEX program to identify comments in the program.
- 7) Write the regular expression for the language accepting the strings which are starting with 1 and ending with 0, over the set $\Sigma = \{0,1\}$.
- 8) Write a regular expression for identifiers and reserved words. Design the transition diagrams for them.
- 9) Write a regular expression for relation operators. Design a transition diagram for them.
- 10) Write regular expressions for the set of words having a, e, i, o, u appearing in that order, although not necessarily consecutively.
- 11) Give general format for LEX program.
- 12) Compare compiler and interpreter with suitable diagrams.
- 13) Why lexical and syntax analyzer are separated out?
- 14) Write about a brief note on various tools available for compiler construction.
- 15) Explain the three general approaches for the implementation of a Lexical analyzer.[8]
- 16) Describe the lexical errors and various error recovery strategies with suitable examples.[8]
- 17) Draw a block diagram of phases of a compiler and indicate the main functions of each phase.[8]
- 18) Define lexeme, token and pattern. Identify the lexemes that make up the tokens in the following program segment. Indicate corresponding token and pattern.
void swap(int i, int j)
{
int t;
t=i;
i=j;
j=t;
}
- 19) Describe the functionality of compilers in a typical language processing system.
- 20) Give the LEX specification to identify reserved words and identified in C.
- 21) What is regular expression? Write the regular expression following. [5] i) Accepts all strings of 0's & 1's which and with 01. ii) Accepts all strings of 0's & 1's, whose 9th position from the right end is 1. iii) Equal no of 1's & 0's
- 22) What is Transition diagram? Explain with examples.
- 23) Describe the need and functionality of linkers, assemblers and loaders.
- 24) Explain the different phases of the compiler, showing the output of each phase using the example for the statement $z = (a*20) + b - c$
- 25) Construct NFA equivalent to regular expression $r = (a + b)^*ab$. [8]