

Automata Theory And Compiler Design:-

Automata:-

It is a "abstract machine (or) a model."

Automata Theory:-

It is a study of "abstract machine (or) a
"computing device."

History of Automata Theory:-

- In 1930's, "Allen Turing" studied Abstract machine.
- He studied this theory, to know "what a computing device could do and what a computing device could not do."
- In 1940's and 1950's, "Finite Automata" was studied by some scientists.
- We use this study in "software and hardware", "computing device".
- In late 1950's, "Chomsky" studied Grammer, these Grammers are used the "development of software".
- In 1960's, "S. Cook", he extended study of "Allen Turing" and he was able to separate Problems that can be solved by computer and that

can't be solved by computer.

What was the reason to study Automata Theory?

→ It was used to ^{design and} develop the some softwares.

→ We use this in lexical ^{Analyser} software, digital circuits software

to verify some systems, ^{↳ Breaks the Programs into logical units}

to web scan systems ^{↳ key words}

We use the automata Theory. ^{Identifiers Punctuations}

Introduction To Finite Automata Theory:-

Finite Automata:-

The Automata that exists in "finite no. of states".

→ The purpose of the "state" is to "remember certain portion of system history".

Advantage:-

→ We can develop the device/system in fixed no. of states/resources.

How to represent Finite Automata:-

→ To represent finite Automata, we use

(circle) $\bigcirc$ → To represent states

(start) $\longrightarrow$ ⇒ To represent start state

It is label to the start state. It will be there in arrow.

(double circle) $\bigcirc$ → To represent the final states.
↓ Accepting states.

(arc) $\curvearrowright$ → To represent transition. [movement].

It is labelled with the some input.

How To model a software:-

→ The finite automaton for the part of lexical analyser that recognises the string "then".

→ It requires 5 states. It refers to the Prefixes of the "then". [t, th, the, then, $\bigcirc$ (one empty state)].



Structural Representation:

Grammars

→ used in the development of the processing ^{the} data with a recursive structure.

Regular Expressions

→ used for the structure or organise the data.

Eq:- UNIX style

$[A-Z][a-z]^*[]$

$[A-Z]; [A-Z] -$

$[A-Z] \rightarrow$ any capital letter

$[a-z] \rightarrow$ any small letter

$*$ $\rightarrow$ any no. of occurrences

$[] \rightarrow$ a blank space

Eq:- Delhi IN

The given regular expression represent a capitalized word

followed by blank space followed by 2 capital letters.

It is a single name followed by two letters.

For multiword city names

$[A-Z][a-z]^*[] [A-Z]$

West Bengal IN $[a-z]^*$

$[A-Z][A-Z]$

Automata & Complexity:-

It discusses limitations of computations.

→ What a computer can do at all?

This study is called decidability and problems that can be solved with sometimes are called decidable.

→ What a computer can do efficiently?

This study is called intractability and the problems that can be solved with ^{no} more time called "tractable" or "undecidable".

Central Concepts of Automata Theory:-

Alphabet: It is a finite non empty set of symbols.

It is denoted by " Σ ".

Eq:- " $\Sigma = \{0,1\}$ " -- Binary Alphabet.

$\Sigma = \{a, \dots, z\}$ -- Small letter Alphabet.

ASCII: character [A-Z - ASCII - values]

String: It is a finite sequence of input symbols chosen from an alphabet Σ .

Eq:- "0111" is a string chosen from a binary alphabet. $\Sigma = \{0,1\}$

It is represented by w .

chosen from an alphabet Σ .

1) Empty string:- It is a string that contains zero occurrences of input symbols.

→ It is represented as ϵ (Epsilon).

→ It can be a string that can be any alphabet.

2) Length of the string:-

→ It can be no. of positions of symbols in the string. → It is represented as " $|w|$ ".

Ex:- $w = 1010111$

The length of the string " $|w|$ " = 7

3) Powers of the string:-

→ If Σ is an alphabet, then Σ^k is a set of strings of length k .

Ex:- $\Sigma = \{a, b\}$

$\Sigma^0 = \phi$, $\Sigma^1 = \{a, b\}$, $\Sigma^2 = \{aa, ab, ba, bb\}$

$\Sigma^3 = \{aaa, aab, aba, abb, baa, bab, bba, bbb\}$

→ The set of all strings over some alphabet Σ is denoted by " Σ^* ".

$\Sigma^* = \Sigma^0 \cup \Sigma^1 \cup \Sigma^2 \cup \Sigma^3 \dots$

→ The set of all non empty strings is denoted by " Σ^+ ".

$$\Sigma^+ = \Sigma^1 \cup \Sigma^2 \cup \Sigma^3 \dots$$

The relation b/w the both strings is " $\Sigma^* = \Sigma^+ \cup \epsilon$ ".

4) Concatenation of strings:-

If x and y are two strings then the concatenation of x, y is represented by a " xy " where the string x followed by y .

Ex:- $x = 10001$
 $y = 00110$ } $\Sigma = \{0, 1\}$

$$xy = 1000100110$$

$$\rightarrow \epsilon x = x\epsilon = x$$

$\hookrightarrow$ empty string [identity of concatenation].

Language:- It is a set of strings chosen from the Σ^* .

Where Σ is some alphabet.

Ex:- $\Sigma = \{0, 1\}$

L = set of strings with equal no. of 0's and equal no. of 1's.

$$L = \{01, 10, 0011, 1100, 1001, 0110, 0101,$$

$$1010, \dots\}$$

$\rightarrow$ It is represented by " L ".

$\rightarrow$ If Σ is an alphabet and L is subset of Σ^* ,

then L is a language.

→ If L is a language over Σ then it is also a language over
Superset of Σ .

Ex:- Set of binary integers whose value is prime

$$L = \{10, 11, 101, 111, \dots\}$$

→ Σ^* is also a language over the alphabet Σ .

→ ϕ is empty language $\phi \neq \{\epsilon\}$

$\hookrightarrow$ empty
language

$\hookrightarrow$ language with empty
string ϵ

Problem:-

It is a question of deciding whether a string belongs to some particular language or not.

Ex:- $L = \{0, 100, 10, 000, 100, 010, 110, \dots\}$

determine 000 is part of given string or not (Yes)

→ It determine the membership of the string in a language.

→ If Σ is an alphabet and L is a language the question determining whether the string 'w' is there in language L or not is a problem.

→ There are two types of finite Automata:

1) Deterministic Finite Automata:

→ Deterministic refers that for each input symbol there is only one state to which finite automata can transition from its current state.

→ DFA is represented by "5-tuple" notation. It

consists of $A = \{Q, \Sigma, \delta, q_0, F\}$

→ It consists of i) A set of states. It is represented by " Q ".

ii) A set of input symbols, represented by " Σ ".

iii) A transition function $\delta(q, a)$ which takes two arguments "state and input symbol" and give one output.

iv) starting state, represented as " q_0 ".

v) set of final states are (or) Accepting states, represented as " F ".

Notations of DFA:-

→ There are two types of Notations:

1) Transition Diagram:-

It is a Graph that contains

i) For each state, there is a node

ii) For each input symbol and for each state transition function δ is represented by an arc labelled with an input symbol.

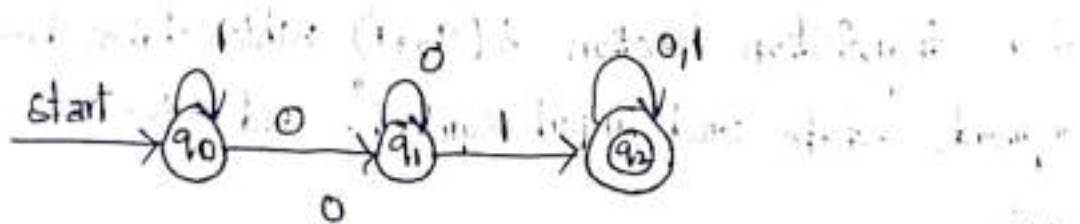
iii) There is an arrow into the starting state labelled as start.

iv) The final state or states must be in double circle $\odot$ and other states must be in single circle $\circ$.

Question:- Construct DFA which accepts all strings of "0's and 1's" that have the sequence "01".

$$\Sigma = \{0, 1\}$$

$$L = \{01, 001, 010, 011, 101, \dots\}$$



2) Transition Table:-

→ Transition Table is a tabular representation of function δ .

→ Rows corresponds to states.

→ columns corresponds to input symbols.

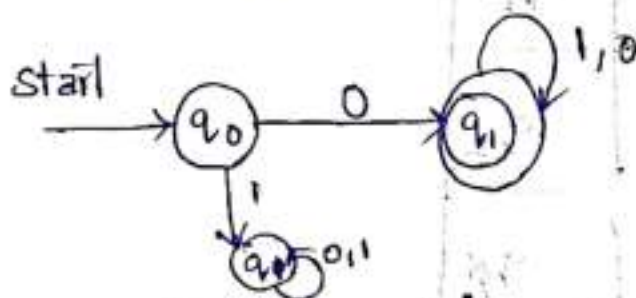
→ Entry corresponds to state, which results from $\delta(q, a)$.

	0	1
$\rightarrow q_0$	q_1	q_0
q_1	q_1	q_2
$*q_2$	q_2	q_2

→ Construct DFA which accepts all strings over alphabet $\Sigma = \{0, 1\}$ starts with "0".

$$\Sigma = \{0, 1\}$$

$$L = \{0, 01, 000, 001, 010, 011, \dots\}$$



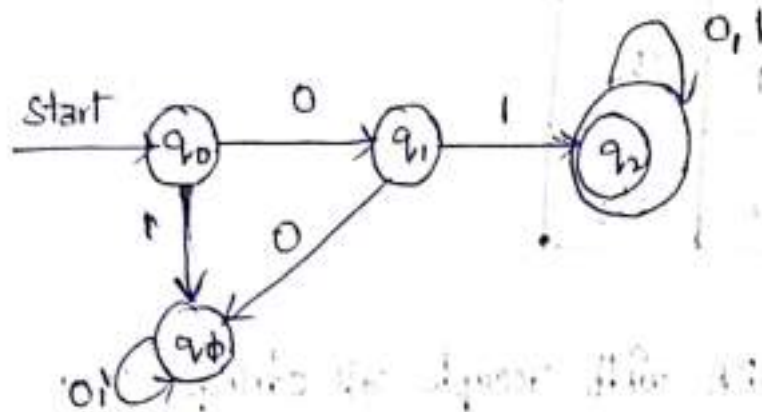
	0	1
$\rightarrow q_0$	q_1	q_2
$*q_1$	q_1	q_1
$*q_2$	q_1	q_1

	0	1
$\rightarrow q_0$	q_1	q_ϕ
q_ϕ	q_ϕ	q_ϕ
$*q_1$	q_1	q_1

→ Construct DFA which accepts all strings over alphabet $\Sigma = \{0, 1\}$ starts with "01".

$$\Sigma = \{0, 1\}$$

$$L = \{01, 10, 010, 011, 0100, 0101, \dots\}$$

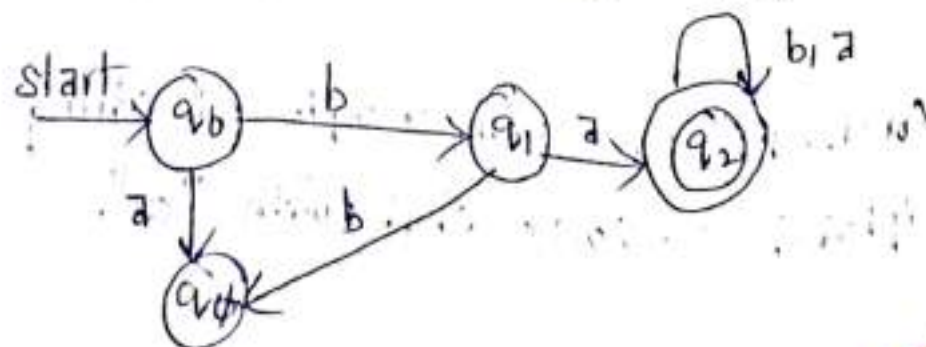


	0	1
$\rightarrow q_0$	q_1	q_ϕ
q_1	q_ϕ	q_2
q_2	q_2	q_2
q_ϕ	q_ϕ	q_ϕ

→ Construct DFA which accepts all strings over alphabet $\Sigma = \{a, b\}$ starts with "ba".

$$\Sigma = \{a, b\}$$

$$L = \{ba, bab, babb, baba, \dots\}$$

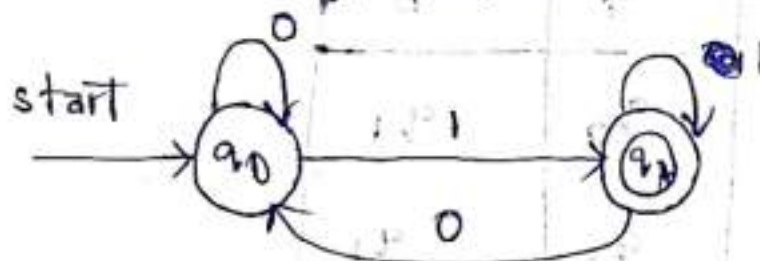


	a	b
→ q ₀	q ₀	q ₁
q ₁	q ₂	q ₀
q ₂	q ₂	q ₂
q ₃	q ₀	q ₀

→ construct DFA for all strings over alphabet $\Sigma = \{0, 1\}$ which ends with "1".

$$\Sigma = \{a, b\}$$

$$L = \{101, 11, 001, 011, 101, 111, \dots\}$$

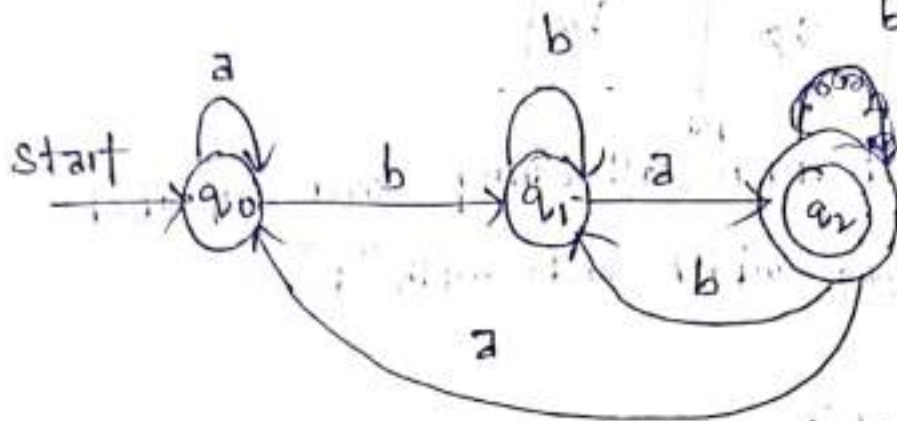


	0	1
→ q ₀	q ₀	q ₁
q ₁	q ₀	q ₁

→ Construct DFA for all strings over the alphabet $\Sigma = \{a, b\}$ which ends with "ba"

$$\Sigma = \{a, b\}$$

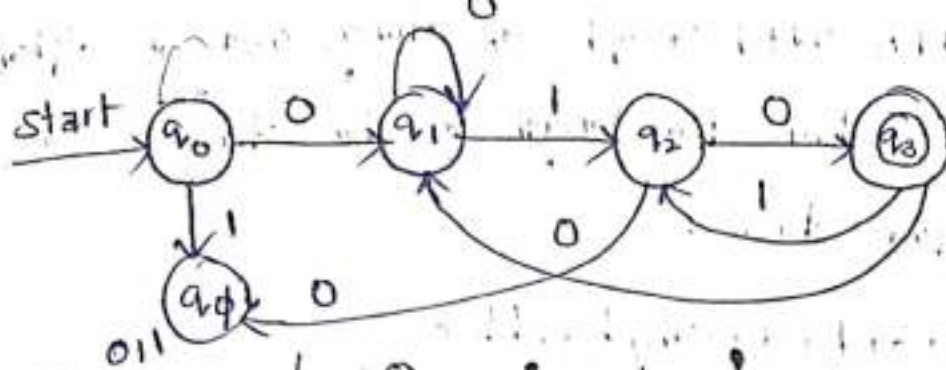
$$L = \{ba, aba, bba, aaba, abba, baba, b bba, \dots\}$$



	a	b
$\rightarrow q_0$	q_0	q_1
q_1	q_2	q_1
q_2	q_0	q_1

→ Construct DFA for all strings over the alphabet $\Sigma = \{0, 1\}$ which ends with 010

$$L = \{010, 0010, 1010, \dots\}$$

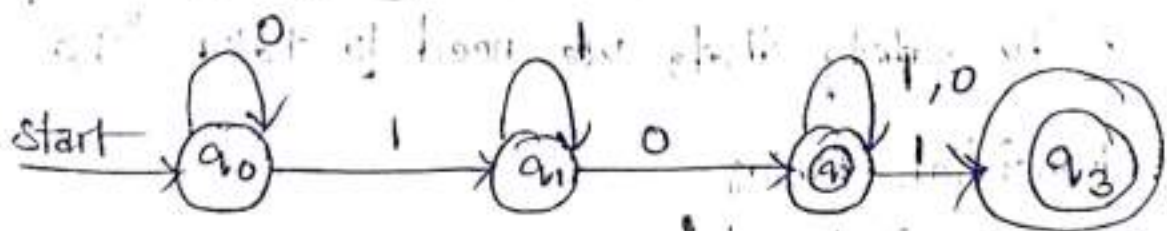


	0	1
→ q ₀	q ₁	q _φ
q ₁	q ₁	q ₂
q ₂	q ₃	q _φ
* q ₃	q ₂	q _φ
q _φ	q _φ	q _φ

→ construct DFA for all strings over the alphabet "01" which contains 10

010, 101, 1101

$L = \{ 010, 101, 1101, 1010, 1011, 0010, \dots \}$



	0	1
→ q ₀	q ₀	q ₁
q ₁	q ₂	q ₁
q ₂	q ₂	q ₃

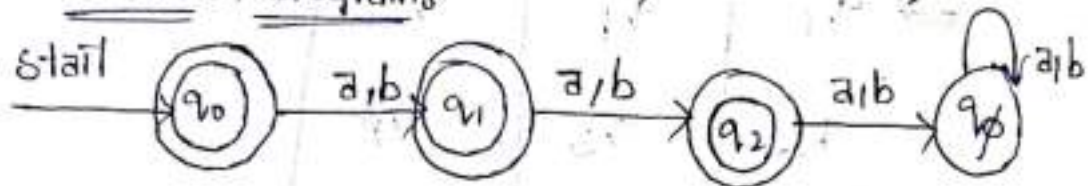
→ Construct DFA which accepts all strings over an alphabet $\Sigma = \{a, b\}$ where the length of the string is less than or equals to "2".

$L = \{ \epsilon, a, b, aa, ab, ba, bb \}$

→ It depends on the length of strings not min "LOS".

→ The length of the string = 2 $\Rightarrow$ 3 states

Transition Diagram:-



→ This DFA satisfies all the strings of the length of the string is " ≤ 2 ".

→ There is no transition of q_2 if we take self transition then we get more than 2 states.

→ But there is no numbers matching with the more than 2 in the language.

→ To satisfy that, we need to take " q_ϕ ".

Transition Table:-

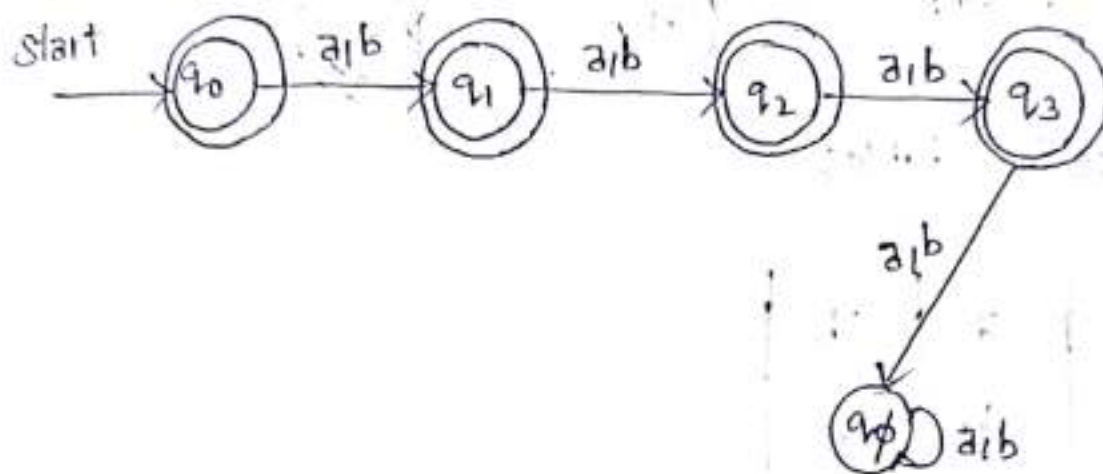
	a	b
* q_0	q_1	q_2
* q_1	q_2	q_2
* q_2	q_ϕ	q_ϕ
q_ϕ	q_ϕ	q_ϕ

→ construct DFA which accepts all strings over an alphabet $\Sigma = \{a, b\}$ where the length of the string is less than or equals to '3'.

$L = \{\epsilon, a, b, aa, ab, ba, bb, aaa, aab, aba, abb, baa, bab, bba, bbb\}$

→ The length of the string is 3 $\Rightarrow$ 4 states,

Transition Diagram:-



Transition Table:-

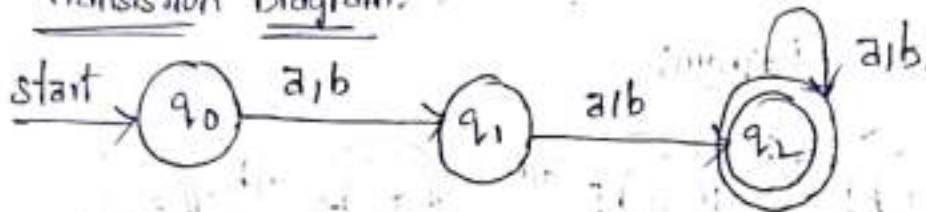
	a	b
* q ₀	q ₁	q ₁
* q ₁	q ₂	q ₂
* q ₂	q ₃	q ₃
* q ₃	q _φ	q _φ
q _φ	q _φ	q _φ

$L = \rightarrow$ Construct DFA which accepts all the strings over the alphabet $\Sigma = \{a, b\}$ where the length of the string is "greater than or equal to 3"

$L = \{aaa, aab, aba, abb, \dots\}$

$\rightarrow$ The length of the string is $2 = 3$ states

Transition Diagram:



Transition Table:

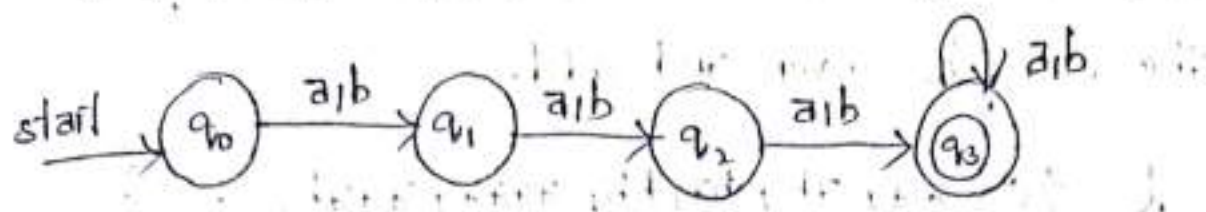
	a	b
$\rightarrow q_0$	q_1	q_1
q_1	q_2	q_2
$*q_2$	q_2	q_2

$\rightarrow$ Construct DFA which accepts all the strings over the alphabet $\Sigma = \{a, b\}$ where the length of the string is "greater than or equal to 3"

$L = \{aaa, aab, aba, abb, aaaa, aaab, \dots\}$

The length of the string is $3 = 4$ states:

Transition Diagram:-



Transition Table:-

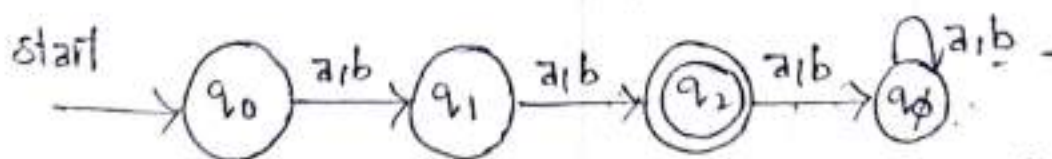
	a	b
→ q ₀	q ₁	q ₁
q ₁	q ₂	q ₂
q ₂	q ₃	q ₃
q ₃	q ₃	q ₃

→ construct DFA which accepts all the strings over the alphabet $\Sigma = \{a, b\}$ where the length of the string is "2".

$$\Sigma = \{aa, ab, ba, bb\}$$

→ The length of the string is 2 = 3 states

Transition Diagram:-



Transition Table:-

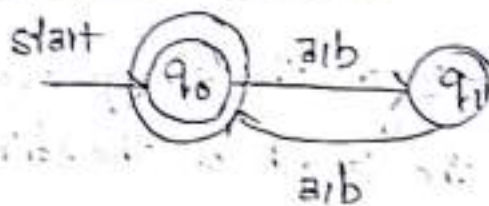
	a	b
→ q ₀	q ₁	q ₁
q ₁	q ₂	q ₂
q ₂	q ₃	q ₃
q ₃	q ₃	q ₃

→ construct the DFA which accepts all the strings over an alphabet $\Sigma = \{a, b\}$ where the length of the string is even and odd.

$L = \{ \epsilon, aa, ab, ba, bb, aaaa, aaab, \dots \}$

→ length of the string is = even and odd $\Rightarrow$ Takes only 2 states.

Transition Diagram:- For even



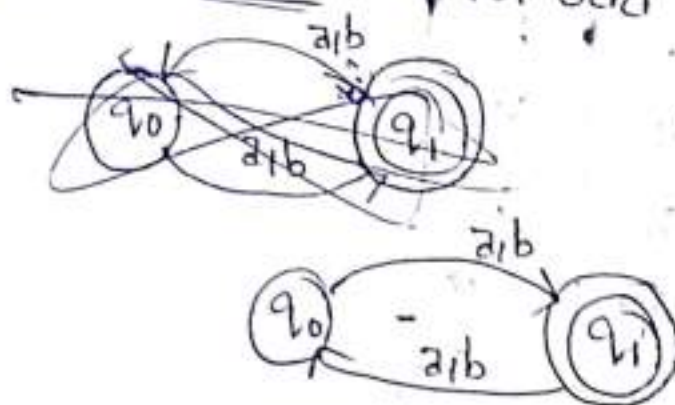
Even numbers [q_0 is final state]

Odd numbers [q_1 is final state]

Transition Table:-

	a	b
$\rightarrow q_0$	q_1	q_1
q_1	q_0	q_0

Transition Diagram:- For odd



Transition Table:-

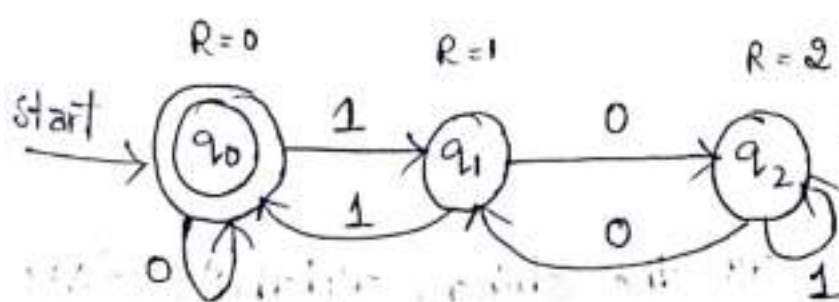
	a	b
$\rightarrow q_0$	q_1	q_1
$* q_1$	q_0	q_0

→ Construct DFA which accepts all the strings over an alphabet $\Sigma = \{0, 1\}$ where the binary integers are divisible by 3. $\hookrightarrow$ The lengths of the

$L = \{0, 11, 110, 1001, 1100, 1111, \dots\}$

→ The length of string is depends on, the "remainder of the number".

→ The length of the string is "3".



Transition Table:-

	0	1
$\rightarrow q_0$	q_0	q_1
q_1	q_2	q_1
q_2	q_1	q_2

Decimal number	Binary equivalent	Remainder
0	0	0
1	1	1
2	10	2
3	11	0
4	100	1
5	101	0



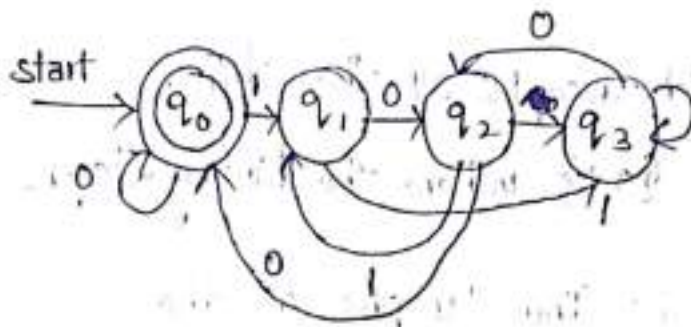
→ construct DFA which accepts all the strings over...

alphabet $\Sigma = \{0,1\}$ where binary integers are divisible by "4".

$$d = \{0, 100, 1000, 1100, 1110, \dots\}$$

remainder $R = 0, 1, 2, 3$

- Decimal binary remainder



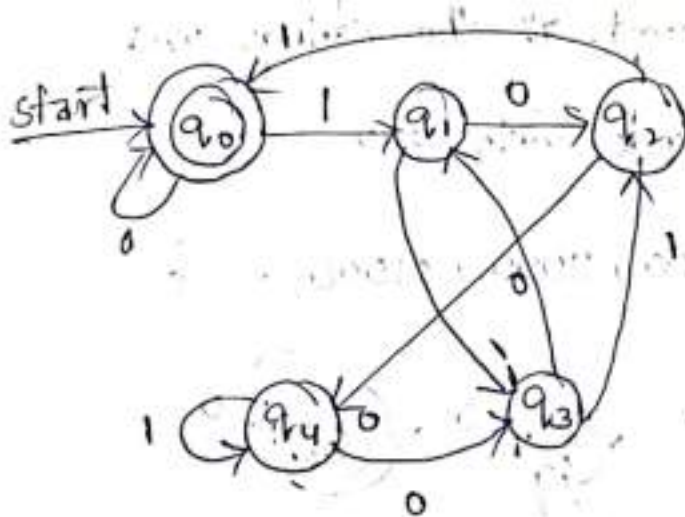
Transition Table

	0	1
$\rightarrow q_0$	q_0	q_1
q_1	q_2	q_3
q_2	q_0	q_1
q_3	q_2	q_3

→ construct DFA all the strings alphabet $\Sigma = \{0,1\}$ where are divisible by "5".

$$d = \{0, 101, 1010, 1111, \dots\}$$

$R = 0, 1, 2, 3, 4$



Decimal binary remainder

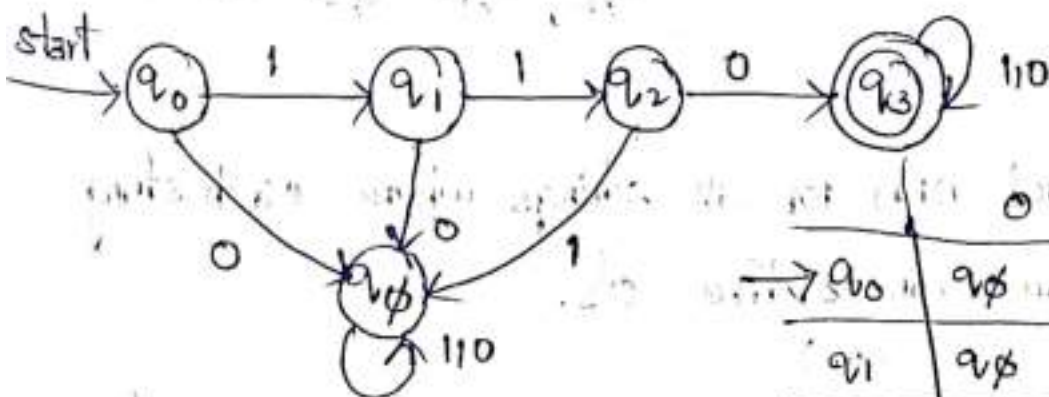
0	0	0
1	1	1
2	10	2
3	11	3
4	100	4
5	101	0
6	110	1
7	111	2
8	1000	3
9	1001	4

Transition Table:-

	0	1
→ q ₀	q ₀	q ₁
q ₁	q ₂	q ₃
q ₂	q ₄	q ₀
q ₃	q ₁	q ₂
q ₄	q ₃	q ₄

→ construct DFA which accepts all the strings over an alphabet $\Sigma = \{0, 1\}$ starts with 110.

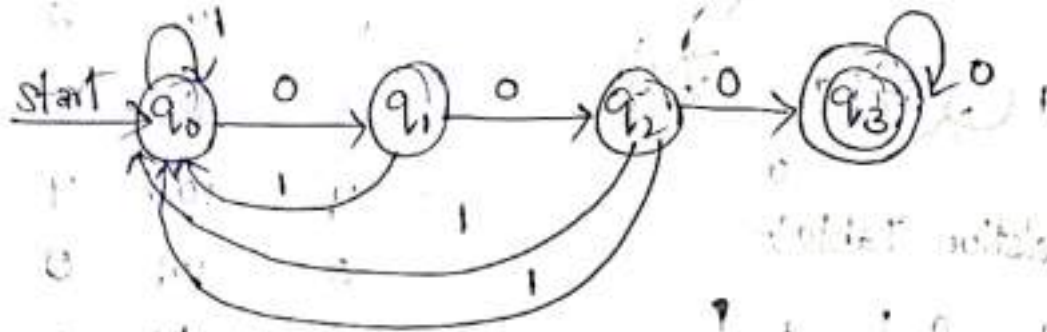
$L = \{110, 1101, 1100, 11000, 11001, 11010, \dots\}$



	0	1
→ q ₀	q _φ	q ₁
q ₁	q _φ	q ₂
q ₂	q ₃	q _φ
q ₃	q ₃	q ₃
q _φ	q _φ	q _φ

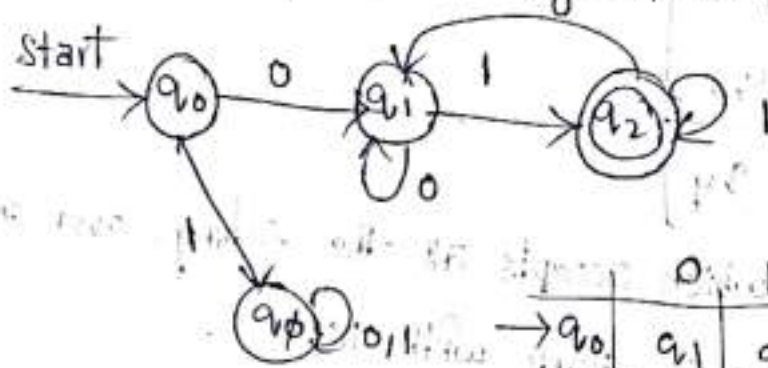
→ construct DFA, which accepts all the strings over an alphabet $\Sigma = \{0, 1\}$ ends with '000'

$L = \{000, 1000, 0000, 11000, 01000, \dots\}$



→ construct DFA $\Sigma = \{0, 1\}$ starts with 0 and ends with 1.

$L = \{01, 001, 011, 0001, 0101, 0111, 0001, \dots\}$



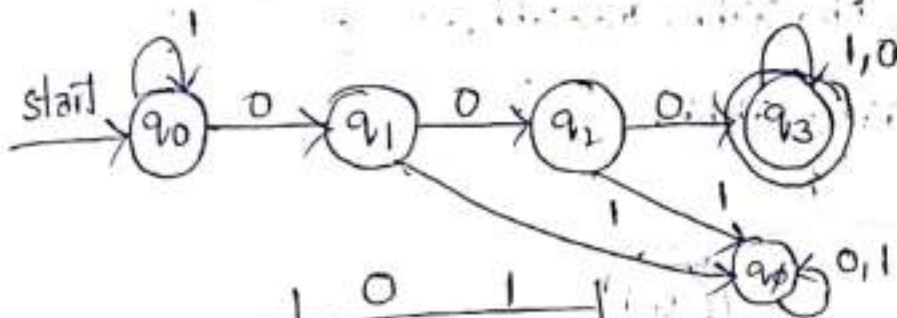
	0	1
q_0	q_1	q_ϕ
q_1	q_1	q_2
q_2	q_1	q_2

→ construct DFA for all strings where each string has 'three consecutive 0's'.

(or)
Construct DFA for all strings which have/contains '000'.

$d = \{ 000, 1000, 0001, 0010, 00010, 10001, \dots \}$

length of the string = 3 then $n+1$ states = 4

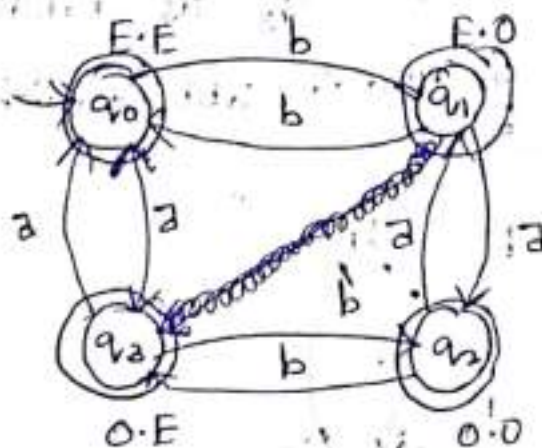


	0	1
$\rightarrow q_0$	q_1	q_0
q_1	q_2	q_4
q_2	q_3	q_4
$*q_3$	q_3	q_3

→ construct DFA's for:

- 1) Even number of a's and even number of b's.
- 2) Even number of b's and odd number of a's.
- 3) Even number of a's and even number of b's.
- 4) odd number of a's and odd number of b's.

1) Even number of a's and odd number of b's.

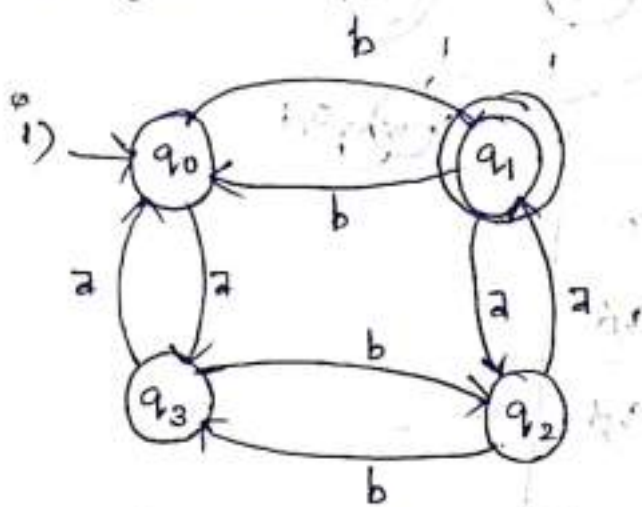


i) $L = \{b, bbb, aab, aabbb, \dots\}$

ii) $L = \{a, aaa, abb, aaabb, \dots\}$

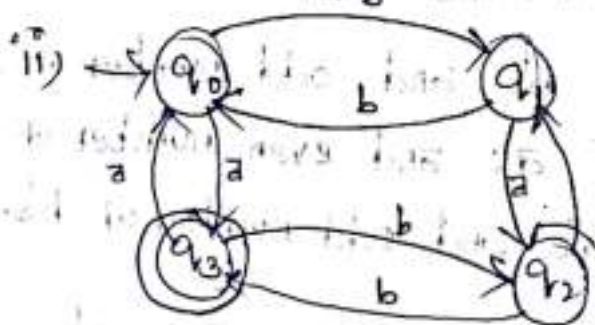
iii) $L = \{e, aa, bb, aabb, aaabb, \dots\}$

iv) $L = \{ab, aab, abb, \dots\}$



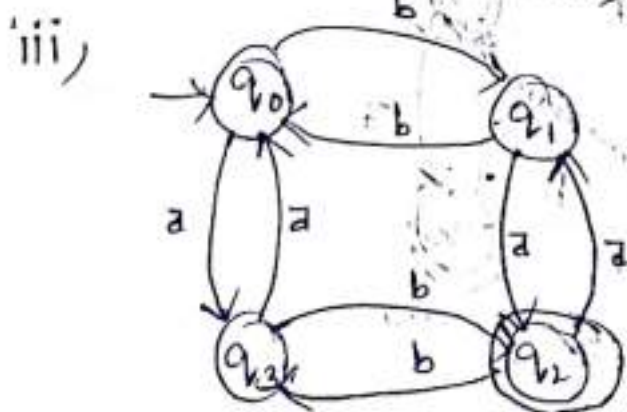
here we take " q_1 " as final state.

$L = \{b, bbb, aab, aabbb, \dots\}$



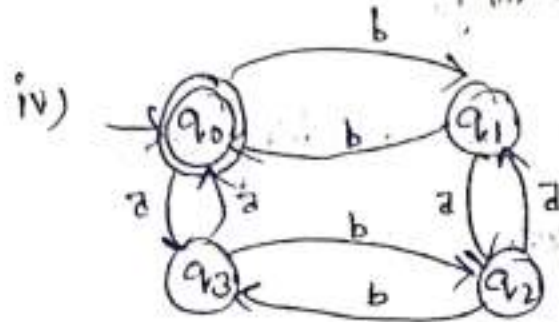
here we take " q_3 " as final state

$L = \{a, aaa, abb, aaabb, \dots\}$



here we take " q_2 " as final state

$$L = \{ \epsilon, aa, bb, aabb, aaaabb, \dots \}$$



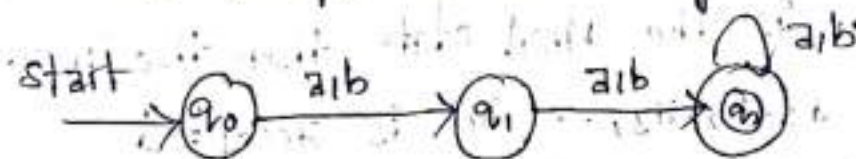
Here we take "q0" as final state

$$L = \{ ab, aab, abbb, \dots \}$$

→ Design a DFA that $L(M) = \{ x/x \text{ is a string of 0's and 1's and } |x| \geq 2 \}$

$$L = \{ aa, ab, ba, bb, aaa, aab, aba, abb, \dots \}$$

→ The length of the string is ≥ 2 = 3 states



	a	b
→ q ₀	q ₁	q ₁
q ₁	q ₂	q ₂
q ₂	q ₂	q ₂

How DFA process a string:

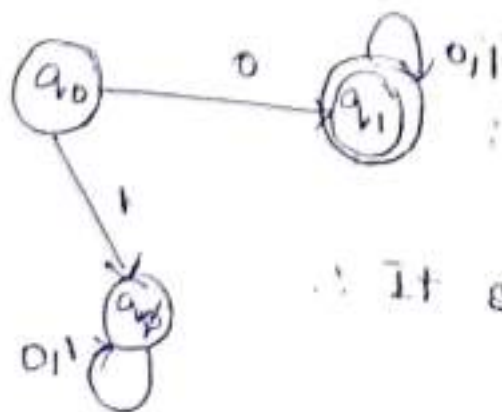
Let $a_1, a_2, a_3, \dots, a_n$ is a string (or)

Some input sequences.

- start DFA from initial state (q_0)
- consult transition from $\delta(q_0, a_1)$ to find state to which transition takes place from q_0 on input symbol a_1 .
- for the next input symbol evaluate $\delta(q_1, a_2)$ to find next state (q_2).
- continue this process to find $q_3, q_4, \dots, q_n$.
- If q_n is the final state then the input sequence $a_1, a_2, a_3, \dots, a_n$ is accepted otherwise it is rejected.

Example:-

0010010



∴ It satisfies the data.

Extended Transition Function:-

An extended transition function describes, what happens when we start from any state and follow some input sequence.

→ Extended Transition function is denoted by $\hat{\delta}$.

→ Extended Transition function we give state and input sequence.

$$\hat{\delta}(q, w) = \delta(\hat{\delta}(q, x), a) \dots = p$$

Where $w = xa$

Language of DFA:-

The language of DFA is denoted by $L(A)$.

→ It is defined as

$$L(A) = \{ w \mid \delta(q_0, w) \in F \}$$

$$L(A) = \{ w \mid \delta(q_0, w) \text{ is in } F \}$$

→ If L is a language of DFA $(L(A))$ for some DFA "A" then L is called a "regular language".

Non-Deterministic Finite Automata:- (NFA)

→ In which the finite automata can exist in any number of states on input from a state.

→ NFA is also a "5-tuple" notation. It consists of set of $(Q, \Sigma, \delta, q_0, F)$.

→ Q : set of states

Σ : set of input symbols

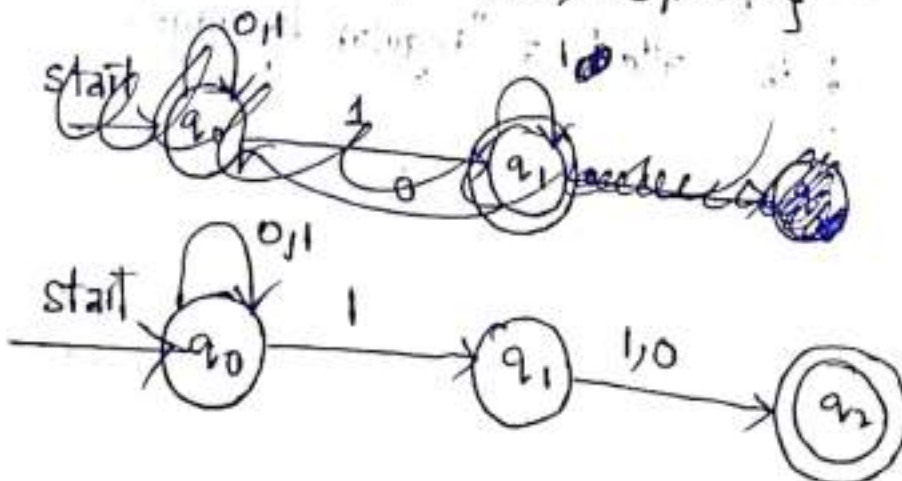
δ : A transition function, $\delta(q_0, i)$, which takes two arguments "state and input symbol" and give one output.

q_0 : starting state.

F : set of final states (or) Accepting states.

→ Construct NFA which accepts all strings whose second last symbol is "1".

$L = \{10, 11, 010, 011, 111, 110, \dots\}$



	0	1
$\rightarrow q_0$	q_0 q_1	q_1
q_1	q_2	q_2
$*q_2$	ϕ	ϕ

Extended Transition Function:-

The Language of NFA:-

It is also denoted by $L(A)$ and it is defined as $L(A) = \{w \mid \hat{\delta}(q_0, w) \cap F \neq \emptyset\}$

Extended Transition Function:-

$$\hat{\delta}(q_0, w) = \{r_1, r_2, \dots, r_n\}$$

Where, $r_1, r_2, \dots, r_n$ are final states.

Finite Automata with epsilon Transitions:-

$\rightarrow$ It is represented as "NFA- ϵ ". (or) " ϵ -NFA".

$\rightarrow$ It is a "5 tuple" notation. where

$$A = (Q, \Sigma, \delta, q_0, F) \text{ transition.}$$

Q = set of states

$$\Sigma = \Sigma \cup \{\epsilon\}$$

$\delta =$ A transition function, $\delta(q, a)$, which takes two arguments "state and input symbol" and give one output.

$q_0 =$ starting state

$F =$ set of Final states/Accepting states.

Uses of Epsilon:-

i) Each epsilon along the path is "Invisible", i.e., it contributes nothing to string along the path.

→ Construct epsilon NFA which accepts decimal numbers consisting of

ii) A string of digits

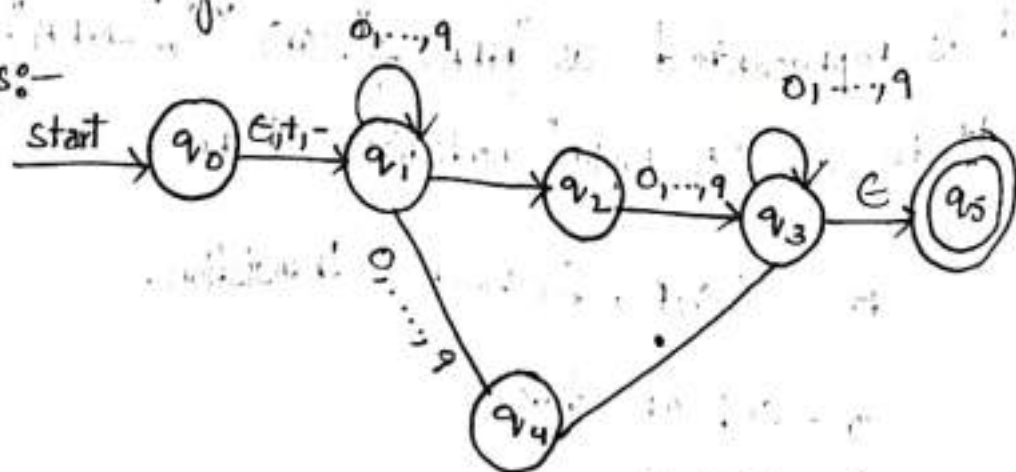
iii) A decimal point

iv) Another string of digits.

v) Either these string of digits can be empty but

vi) Atleast one of string of digits must be non empty.

Ans:-



Extended Transition Function:-

$$\hat{\delta}(q_0, w) = \bigcup_{j=1}^m (\epsilon\text{-closure}(r_j))$$

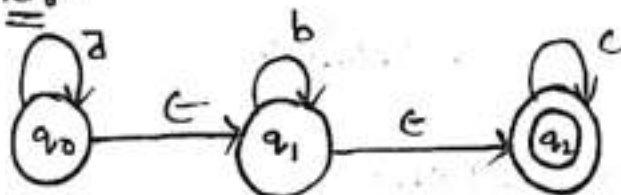
$$= \epsilon\text{-closure}(r_1) \cup \epsilon\text{-closure}(r_2) \cup \dots$$

$$\epsilon\text{-closure}(r_m)$$

Epsilon Closure:-

ϵ -closure of a state Q consists of every state that can be reached from Q along a path with " ϵ ".

Example:-



ϵ -closure(q_0):- There is an 'invisible' transition to the 'itself'. " $\{q_0, q_1, q_2\}$ "

ϵ -closure(q_1): " $\{q_1, q_2\}$ "

ϵ -closure(q_2): " $\{q_2\}$ "

Eliminating ϵ -Transitions:-

Conversion of "NFA- ϵ " (or) "ε-NFA" NFA without Transitions of ϵ :-

→ Find ϵ -closures for all states of the given

ϵ -NFA.

→ Find Transitions for each state/input symbol.

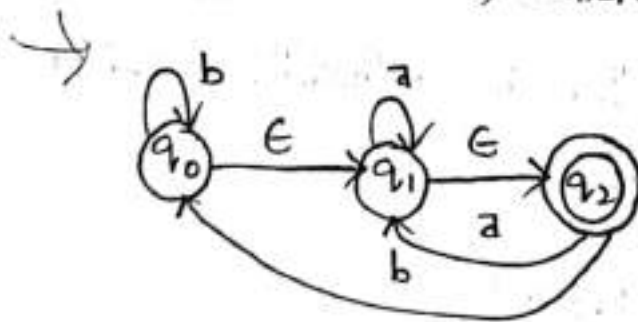
→ ~~construct~~ $\hat{\delta}(q, a) = \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q), a))$

→ construct NFA without epsilon Transitions based on the computations of Transition function.

→ Decide Final state as (F')

i) $F' = F \cup \{q_0\}$ if $\epsilon\text{-closure}(q_0)$ is having F as a member.

ii) $F' = F$, otherwise



Step 1:- Find ϵ -closures for all states

$$\epsilon\text{-closure}(q_0) = \{q_0, q_1, q_2\}$$

$$\epsilon\text{-closure}(q_1) = \{q_1, q_2\}$$

$$\epsilon\text{-closure}(q_2) = \{q_2\}$$

Step 2:- Find the Transitions for each state on input symbol.

$$\begin{aligned}\Rightarrow \hat{\delta}(q_1, a) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_0), a))) \\ &= \epsilon(\text{closure}(\delta((q_0, q_1, q_2), a))) \\ &= \epsilon(\text{closure}(\delta(q_0, a) \cup \delta(q_1, a) \cup \delta(q_2, a))) \\ &= \epsilon(\text{closure}(\phi \cup q_1 \cup q_2)) \\ &= \epsilon(\text{closure}(q_1)) \\ &= \{q_1, q_2\}\end{aligned}$$

Step 3

$$\begin{aligned}\Rightarrow \hat{\delta}(q_1, b) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_0), b))) \\ &= \epsilon(\text{closure}(\delta((q_0, q_1, q_2), b))) \\ &= \epsilon(\text{closure}(\delta(q_0, b) \cup \delta(q_1, b) \cup \delta(q_2, b))) \\ &= \epsilon(\text{closure}(q_0 \cup \phi \cup q_0)) \\ &= \epsilon(\text{closure}(q_0)) \\ &= \{q_0, q_1, q_2\}\end{aligned}$$

$$\begin{aligned}\Rightarrow \hat{\delta}(q_1, a) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_1), a))) \\ &= \epsilon(\text{closure}(\delta((q_1, q_2), a))) \\ &= \epsilon(\text{closure}(\delta(q_1, a) \cup \delta(q_2, a)))\end{aligned}$$

$$= \epsilon(\text{closure}(q_1 \cup q_1))$$

$$= \epsilon(\text{closure}(q_1))$$

$$= \{q_1, q_2\}$$

$$\Rightarrow \hat{\delta}(q_1, b) = \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_1), b)))$$

$$= \epsilon(\text{closure}(\delta(q_1, q_2), b))$$

$$= \epsilon(\text{closure}(\delta(q_1, b) \cup \delta(q_2, b)))$$

$$= \epsilon(\text{closure}(q_0 \cup q_0))$$

$$= \epsilon(\text{closure}(q_0))$$

$$= \{q_0, q_1, q_2\}$$

$$\Rightarrow \hat{\delta}(q_2, a) = \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_2), a)))$$

$$= \epsilon(\text{closure}(\delta(q_2), a))$$

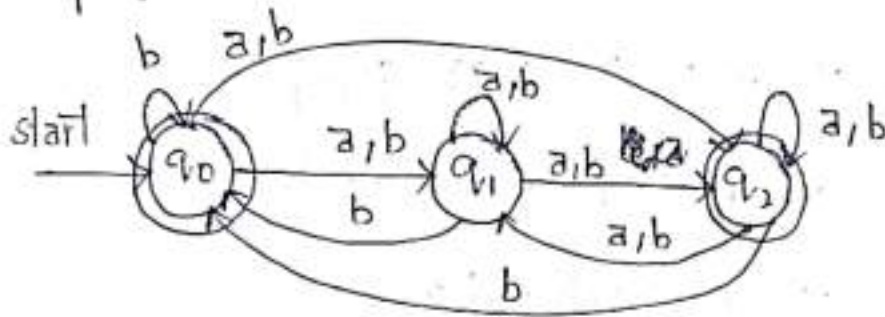
$$= \epsilon(\text{closure}(\delta(q_2, a)))$$

$$= \epsilon(\text{closure}(q_1))$$

$$= \{q_1, q_2\}$$

$$\begin{aligned}
 \Rightarrow \hat{\delta}(q_2, b) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_2), b))) \\
 &= \epsilon(\text{closure}(\delta(q_2), b)) \\
 &= \epsilon(\text{closure}(\delta(q_2, b))) \\
 &= \epsilon(\text{closure}(q_0)) \\
 &= \{q_0, q_1, q_2\}
 \end{aligned}$$

Step 3:-



Step 4:-

$F' = F \cup \{q_0\}$ if $\epsilon\text{-closure}(q_0)$ is
 having F as a member.
 Otherwise $F' = F$.

$$F = \{q_2\}$$

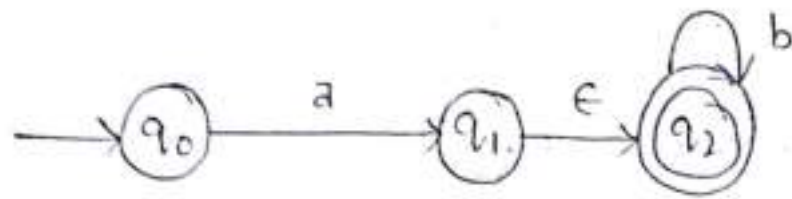
$$\epsilon\text{-closure}(q_0) = \{q_0, q_1, q_2\}$$

$$F' = F \cup \{q_0\}$$

$$= q_2 \cup \{q_0\}$$

$$= \{q_2, q_0\}$$

→ convert following ϵ -NFA to the NFA without ϵ .



Ans:-

Step 1:- Find ϵ closures for all states

$$\epsilon\text{-closure}(q_0) = \{q_0\}$$

$$\epsilon\text{-closure}(q_1) = \{q_1, q_2\}$$

$$\epsilon\text{-closure}(q_2) = \{q_2\}$$

Step 2:- Find the Transitions for each state on input symbol.

$$\begin{aligned} \Rightarrow \hat{\delta}(q_0, a) &= \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0), a)) \\ &= \epsilon\text{-closure}(\delta(\emptyset, a)) \\ &= \epsilon\text{-closure}(q_1) \\ &= \{q_1, q_2\} \end{aligned}$$

$$\begin{aligned} \Rightarrow \hat{\delta}(q_0, b) &= \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0), b)) \\ &= \epsilon\text{-closure}(\delta(q_0, b)) \\ &= \epsilon\text{-closure}(\emptyset) \\ &= \emptyset \end{aligned}$$

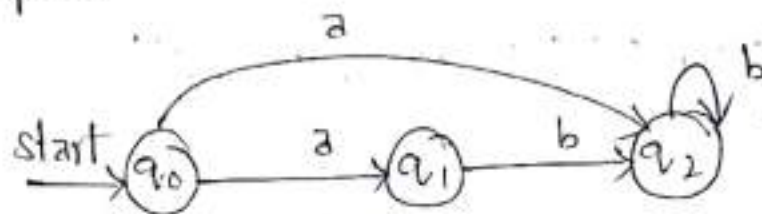
$$\begin{aligned}
 \Rightarrow \hat{\delta}(q_1, a) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_1), a))) \\
 &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_1), a))) \\
 &= \phi
 \end{aligned}$$

$$\begin{aligned}
 \Rightarrow \hat{\delta}(q_1, a) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_1), b))) \\
 &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_1), b))) \\
 &= \epsilon(\text{closure}(q_2)) \\
 &= q_2 \cup q_2
 \end{aligned}$$

$$\begin{aligned}
 \Rightarrow \hat{\delta}(q_2, a) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_2), a))) \\
 &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_2), a))) \\
 &= \phi \cup q_2 \cup \phi
 \end{aligned}$$

$$\begin{aligned}
 \Rightarrow \hat{\delta}(q_2, b) &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_2), b))) \\
 &= \epsilon(\text{closure}(\delta(\epsilon\text{-closure}(q_2), b))) \\
 &= \epsilon(\text{closure}(q_2)) \\
 &= q_2 \cup q_2
 \end{aligned}$$

Step 3:-



Step 4:-

$F' = F \cup \{q_0\}$ if ϵ -closure(q_0) is
having F as a member
Otherwise $F' = F$

$F = \{q_2\}$ ϵ -closure(q_0) = q_0

Since F is not a member of
 ϵ -closure(q_0) so, $F' = F$

$F' = \{q_2\}$

Conversion of NFA to DFA:-

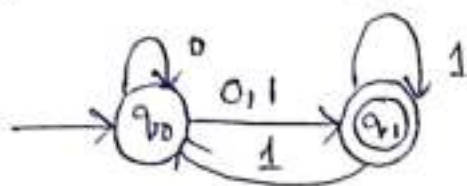
Step 1: start with initial state and find its
transitions.

Step 2: ~~Star~~ If there are any new states find
transitions for each of them.

Step 3:- Continue step 2 until there are no new transitions.

Step 4:- Mark all the states which has the final state of NFA as final state of DFA.

1) convert the following NFA into DFA / Find the equivalent DFA for the given NFA.



Step 1:- Find the transitions which starts with initial state.

	0	1
→ q ₀	[q ₀ , q ₁]	q ₁
[q ₀ , q ₁]	[q ₀ , q ₁]	q ₁
q ₁	∅	[q ₀ , q ₁]

Step 2:- Find the transitions for the new states occurred

$$\delta(q_0, 0) = \{q_0, q_1\}$$

$$\delta(q_0, 1) = \{q_1\}$$

$$\delta([q_0, q_1], 0) = \delta(q_0, 0) \cup \delta(q_1, 0)$$

$$= [q_0, q_1] \cup \phi$$

$$= \{q_0, q_1\}$$

$$\delta([q_0, q_1], 1) = \delta(q_0, 1) \cup \delta(q_1, 1)$$

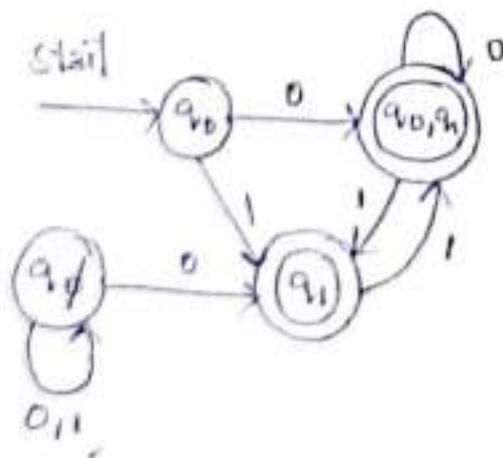
$$= q_1 \cup q_1$$

$$= q_1$$

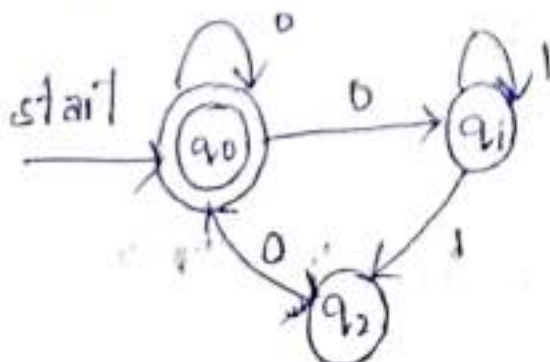
$$\delta(q_1, 0) = \phi$$

$$\delta(q_1, 1) = \{q_0, q_1\}$$

Step 3:- Continue the previous step if there are no new states, after that we need draw the transition diagram from the table.



2) Determine the equivalent DFA to the given NFA.



Step 1:-

$\rightarrow q_0$ ~~$[q_0, q_1]$~~

	0	1
$\rightarrow q_0$	$[q_0, q_1]$	ϕ
$[q_0, q_1]$	$[q_0, q_1]$	$[q_1, q_2]$
$[q_1, q_2]$	q_0	$[q_1, q_2]$
q_ϕ	q_ϕ	q_ϕ

Step 2:-

$$\delta(q_0, 0) = \{q_0, q_1\}$$

$$\delta(q_0, 1) = \phi$$

$$\begin{aligned} \delta([q_0, q_1], 0) &= \delta(q_0, 0) \cup \delta(q_1, 0) \\ &= q_0 \cup \phi \cup q_1 \cup \phi \\ &= q_0 \cup \{q_0, q_1\} \end{aligned}$$

$$\begin{aligned} \delta([q_0, q_1], 1) &= \delta(q_0, 1) \cup \delta(q_1, 1) \\ &= \phi \cup [q_1, q_2] \\ &= [q_1, q_2] \end{aligned}$$

$$\delta(q_1, q_2) =$$

$$\delta(q_1, 0) = \phi$$

$$\delta(q_1, 1) = q_1, q_2$$

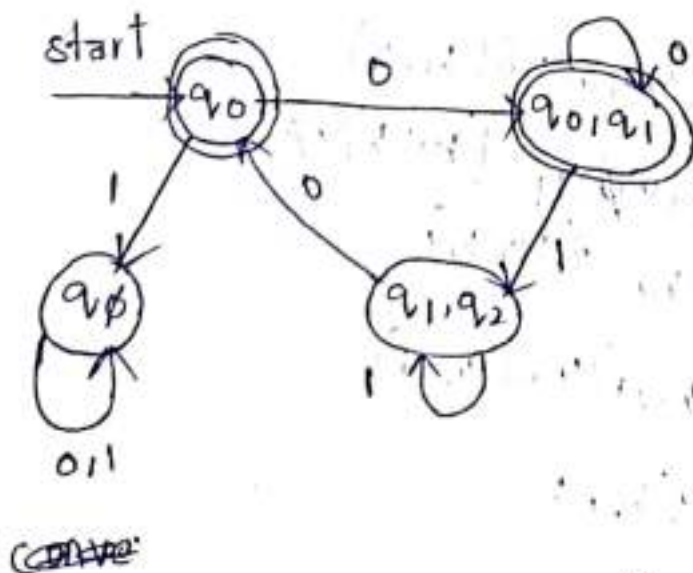
$$\delta(q_2, 0) = q_0$$

$$\delta(q_2, 1) = \phi$$

$$\begin{aligned} \delta((q_1, q_2), 0) &= \delta(q_1, 0) \cup \delta(q_2, 0) \\ &= (\phi) \cup q_0 \\ &= q_0 \end{aligned}$$

$$\begin{aligned} \delta((q_1, q_2), 1) &= \delta(q_1, 1) \cup \delta(q_2, 1) \\ &= (q_1, q_2) \cup \phi \\ &= q_1, q_2 \end{aligned}$$

Step 3:-



conversion of E-NFA to DFA:-

Step 1: Take epsilon-closure for the starting state of NFA as the starting state of DFA

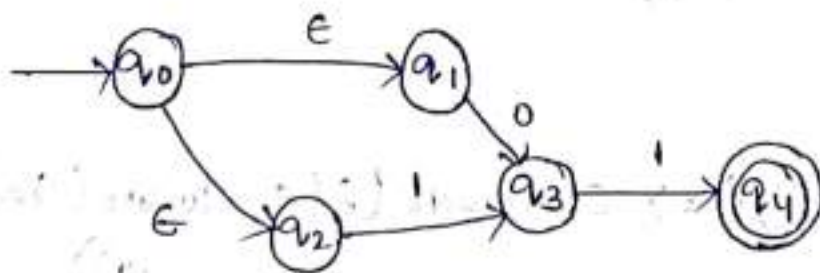
Step 2: Find the states for each input from the current state.

Step 3: If we find a new state, take it as a current state repeat step 2.

Step 4: & repeat step 2 and step 3 until there is no new state.

Step 5: Mark the states of DFA as final states which contains the final state of NFA.

1) convert the given e-NFA into DFA



Step 1: finding the ϵ -closures for the starting state of NFA as the starting state of DFA

$$\epsilon\text{-closure}(q_0) = \{q_0, q_1, q_2\}$$

	0	1
$\rightarrow [q_0, q_1, q_2]$	$\{q_3\}$	$\{q_3\}$
q_3	ϕ	q_4
q_4	ϕ	ϕ
q_ϕ	ϕ	ϕ

$$\hat{\delta}([q_0, q_1, q_2], 0) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, q_1, q_2), 0))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, q_1, q_2), 0))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_0, 0) \cup \delta(q_1, 0) \cup \delta(q_2, 0))$$

$$\Rightarrow \epsilon\text{-closure}(\phi \cup q_3 \cup \phi)$$

$$\Rightarrow \epsilon\text{-closure}(q_3)$$

$$\Rightarrow q_3$$

$$\hat{\delta}([q_0, q_1, q_2], 1) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, q_1, q_2), 1))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_0, q_1, q_2), 1)$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_0, 1) \cup \delta(q_1, 1) \cup \delta(q_2, 1))$$

$$\Rightarrow \epsilon\text{-closure}(\phi \cup \phi \cup q_3)$$

$$\Rightarrow \epsilon\text{-closure}(q_3)$$

$$\Rightarrow q_3$$

$$\hat{\delta}([q_3], 0) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_3), 0))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_3), 0)$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_3, 0))$$

$$\Rightarrow \phi$$

$$\hat{\delta}([q_3], 1) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_3), 1))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_3), 1)$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_3, 1))$$

$$\Rightarrow q_4$$

$$\hat{\delta}([q_4], 0) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_4), 0))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_4), 0)$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_4, 0))$$

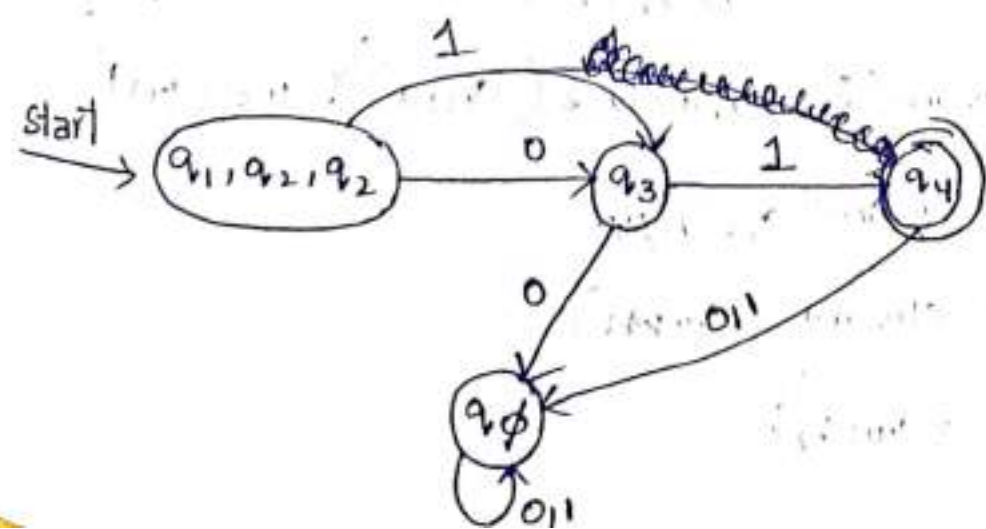
$$\Rightarrow \phi$$

$$\hat{\delta}[q_4, 1] \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_4), 1))$$

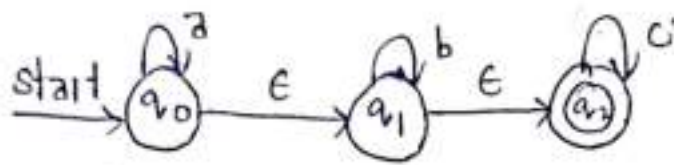
$$\Rightarrow \epsilon\text{-closure}(\delta(q_4), 1)$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_4, 1))$$

$$\Rightarrow \phi$$



→ Determine the equivalent DFA for the given ϵ -NFA.



i) ϵ -closure(q_0) = $\{q_0, q_1\}$

ii) ϵ -closure(q_1) = $\{q_1, q_2\}$

iii) ϵ -closure(q_2) = $\{q_2\}$

	a	b	c
$\{q_0, q_1\}$	$\{q_0, q_1\}$	$\{q_1, q_2\}$	$\{q_2\}$
$\{q_1, q_2\}$	ϕ	$\{q_1, q_2\}$	

$$\hat{\delta}([q_0, q_1], a) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, q_1), a))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_0, a) \cup \delta(q_1, a))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(q_0, a) \cup \delta(q_1, a) \cup \delta(q_2, a))$$

$$\Rightarrow \epsilon\text{-closure}(q_0 \cup \phi)$$

$$\Rightarrow \epsilon\text{-closure}(q_0)$$

$$\Rightarrow \{q_0, q_1\}$$

$$\hat{\delta}([q_0, q_1], b) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, q_1), b))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, b) \cup \delta(q_1, b) \cup \delta(q_2, b)))$$

$$\Rightarrow \epsilon\text{-closure}(\emptyset \cup q_1 \cup \emptyset)$$

$$\Rightarrow \epsilon\text{-closure}(q_1)$$

$$\Rightarrow \{q_1, q_2\}$$

$$\hat{\delta}([q_0, q_1, q_2], c) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, q_1, q_2), c))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_0, c) \cup \delta(q_1, c) \cup \delta(q_2, c)))$$

$$\Rightarrow \epsilon\text{-closure}(\emptyset \cup \emptyset \cup q_2)$$

$$\Rightarrow \{q_2\}$$

$$\hat{\delta}([q_1, q_2], a) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_1, q_2), a))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_1, a) \cup \delta(q_2, a)))$$

$$\Rightarrow \epsilon\text{-closure}(\emptyset \cup \emptyset)$$

$$\Rightarrow \emptyset$$

$$\hat{\delta}([q_1, q_2], b) \Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_1, q_2), b))$$

$$\Rightarrow \epsilon\text{-closure}(\delta(\epsilon\text{-closure}(q_1, b) \cup \delta(q_2, b)))$$

$$\Rightarrow \epsilon\text{-closure}(q_1 \cup \emptyset)$$

$$\Rightarrow \{q_1, q_2\}$$

NFA - Applications: Text Search.

→ It is used for web search (www) and "finding information from text".

Finding strings in Text:-

- 1) Inverted indices
 - 2) Finite Automata
- ↓
It is NFA
- maintain "some set" of words & occurrences
• It requires more memory to travel to search indices

1) Take a transition on an initial state upon each input symbol. → That's why it is used.

2) To accept the string $a_1, a_2, \dots, a_k$ Take transition from q_0 to q_1 on a_1 → This is not suitable for where documents require rapid change.

1) Take Transition from q_1 to q_2 on a_2 and so on $\dots$ Then q_k becomes the final state

Regular Expression: It is a simple expression to describe a language that the DFA accept.

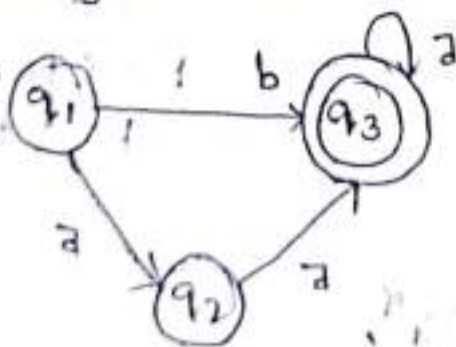
Finite Automata:-

Step 1: Take equation for each state based on incoming edges.

Step 2: Add ϵ to eqn of initial state.

Step 3: Simplify the eqns. using Arden's method and find regular expression for the final state.

i) convert given DFA into regular expression.



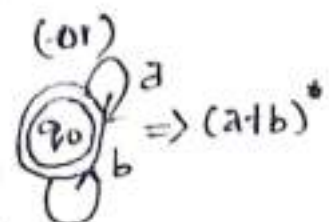
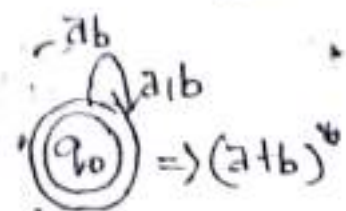
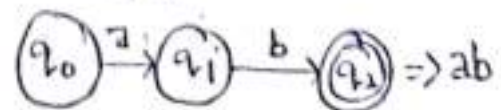
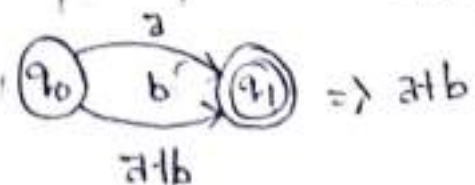
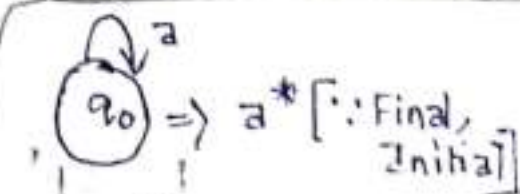
Step 1: Incoming edges

1) $q_1 = \emptyset$ — (1)

2) $q_2 = q_1 a$ — (2)

3) $q_3 = q_1 b + q_2 a + q_3 a$ — (3)

Step 2: $q_1 = \epsilon$ — (1)



Step 3:- $R = Q + RP \Rightarrow R = QP^*$

Substitute eq ① in eq ②

$$q_1 = q_1 a$$

$$= \epsilon a$$

$$q_1 = a \quad \text{--- ④}$$

Substitute eq ④ & ③ in eq ③

$$q_3 = \epsilon b + q_1 a + q_3 a$$

$$= \epsilon b + a a + q_3 a$$

$$\frac{q_3}{R} = \frac{b + aa}{Q}$$

$$+ \frac{q_3}{R} \cdot \frac{a}{P}$$

[This is in the form of

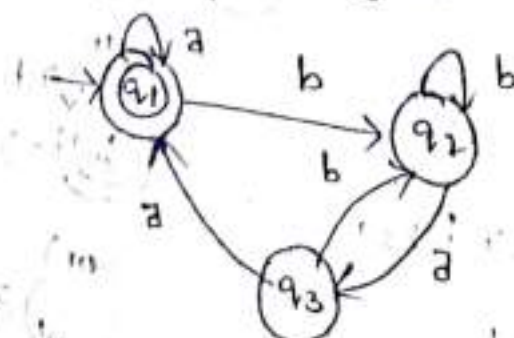
$$R = Q + RP \\ R = Q + RP]$$

$$R = QP^*$$

$$q_3 = (b + aa)a^*$$

This is regular expression.

②



$$1) q_1 = q_1 a + q_3 a \quad \text{--- ①}$$

$$2) q_2 = q_2 b + q_1 b + q_3 b \quad \text{--- ②}$$

$$3) q_3 = q_2 a \quad \text{--- ③}$$

step 2: $q_1 = q_1 a + \epsilon$

step 3: $R = a + RP$

substitute eq (3) in eq (2)

$$q_2 = q_2 b + q_1 b + q_3 b$$

$$q_2 = q_2 b + q_1 b + q_2 ab$$

$$q_2 = q_1 b + q_2 b + q_2 ab$$

$$q_2 = q_1 b + q_2 (b + ab)$$

$$[\because R = a + RP]$$

$$q_2 = (q_1 b) (b + ab)^*$$

$\downarrow$

$$R = aP^*$$

substitute (4) in (3)

$$q_3 = q_2 a$$

$$q_3 = (q_1 b) (b + ab)^* a \quad \text{--- (5)}$$

substitute eq (5) in eq (1)

$$q_1 = q_1 a + q_3 a + \epsilon$$

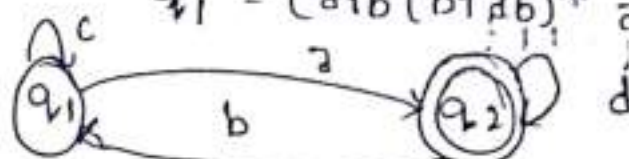
$$q_1 = q_1 a + (q_1 b) (b + ab)^* a + \epsilon$$

$$q_1 = \epsilon + q_1 (a + b(b + ab)^* a)$$

$$q_1 = \epsilon (a + b(b + ab)^* a)^*$$

$$q_1 = (a + b(b + ab)^* a)^*$$

(6)



Sol: Step 1: $q_1 = q_2 b + q_1 c \rightarrow (1)$

$$q_2 = q_1 a + q_2 d \rightarrow (2)$$

$$\text{Step 2: } q_1 = q_2 b + q_1 c + \epsilon \rightarrow (3)$$

$$\text{Step 3: } q_1 = (q_2 b + \epsilon) + q_1 c$$

$$q_1 = (q_2 b + \epsilon) c^* \Rightarrow RP = R = \emptyset P^*$$

substitute eq (3) in eq (2)

$$q_2 = q_2 d + q_1 a$$

$$q_2 = q_2 d + (q_2 b + \epsilon) c^* a$$

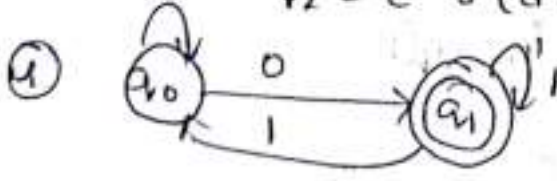
$$q_2 = q_2 d + q_2 b c^* a + c^* a$$

$$q_2 = q_2 (d + b c^* a) + c^* a$$

$$\therefore [R = RP + Q]$$

$\therefore$ Now write this in $R = \emptyset P^*$ form

$$q_2 = c^* a (d + b c^* a)^*$$



$$\text{Step 1: } q_0 = q_0 1 + q_1 1 \rightarrow (1)$$

$$q_1 = q_1 1 + q_0 0 \rightarrow (2)$$

$$\text{Step 2: } q_0 = q_0 0 + q_1 1 + \epsilon$$

$$R = RP + Q$$

$$a_0^1 = (a_1 1 + \epsilon) \cdot 0^* = (3)$$

Substitute eq (3) in eq (2)

$$a_1 = a_1 1 + (a_1 1 + \epsilon) 0^*$$

$$a_1 = a_1 1 + a_1 0^* 0 + 0^*$$

$$a_1 = a_1 (1 + 0^* 0) + 0^* 0$$

$$= a_1 = (0^* 0) (1 + 0^* 0)^*$$

DFA to regular expression by: State elimination method:-

Step 1:- Initial state should not have incoming edges, if exists create a new state make it as initial state.

Step 2:- Finite automata should have only one final state, if there are multiple final states then create a new state and make it as final.

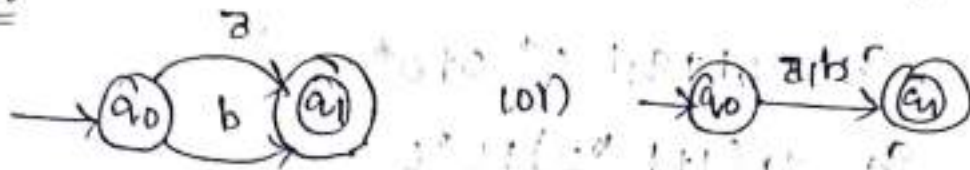
Step 3:- Final state should not have any outgoing edges if exists create a new state and make it as final state.

Step 4:- Eliminate every state except initial and final states.

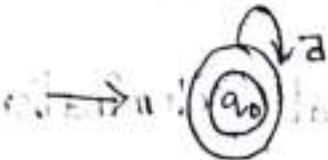
Step 5: Finally FA should have only initial state and final states.

Step 6: The label from initial state to final state becoming the regular expression.

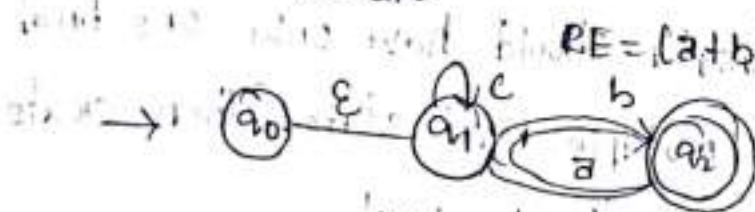
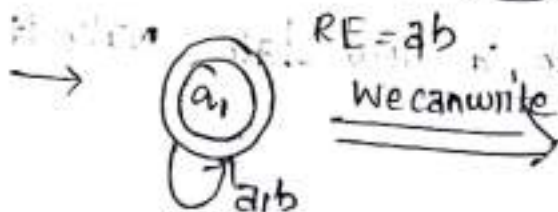
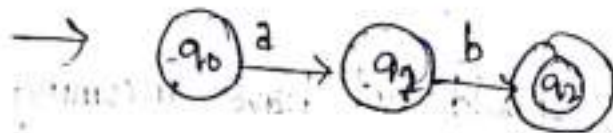
Q:-



∴ RE = a^*b^*



∴ RE = a^*



∴ $\epsilon \cdot (c^* (b^* a b^*)) = c^* b^* + c^* a b^*$

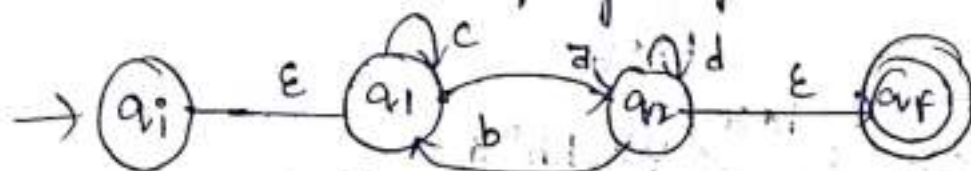


Since, there is an incoming edge for q_1 , take a new state as initial state.

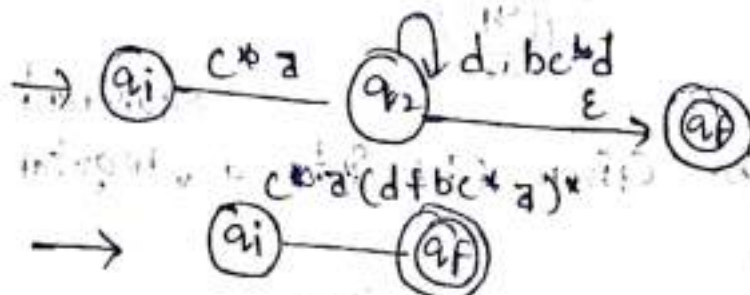
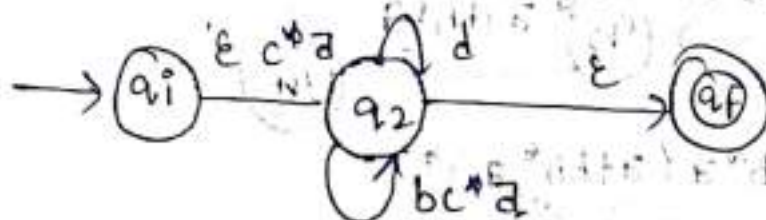


step 3:-

since there are outgoing edges from



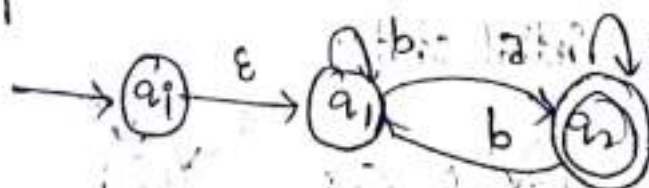
eliminate q_1 , then DFA



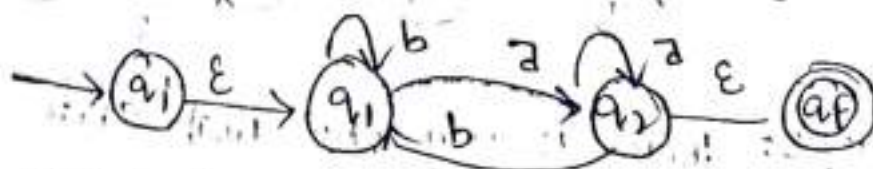
=> Write the regular expression:



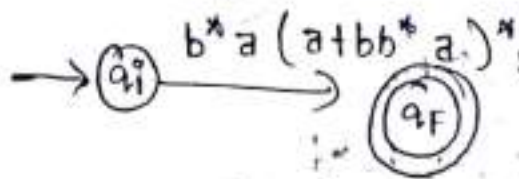
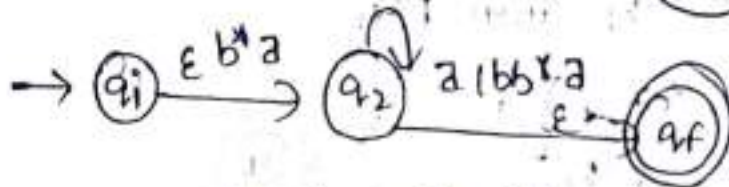
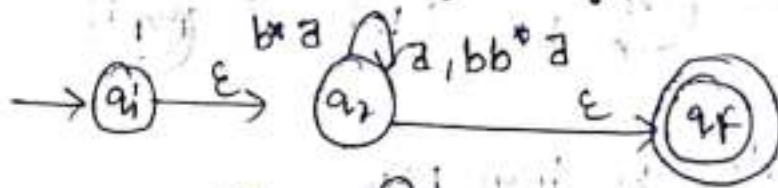
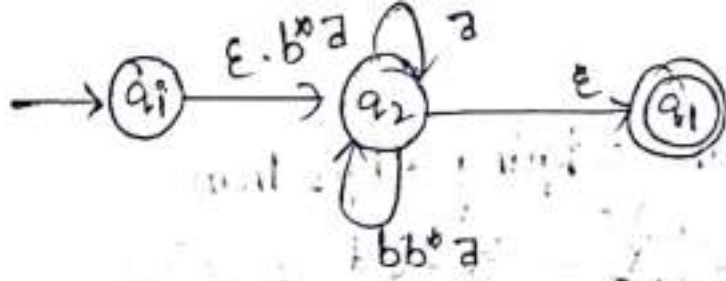
step 1:



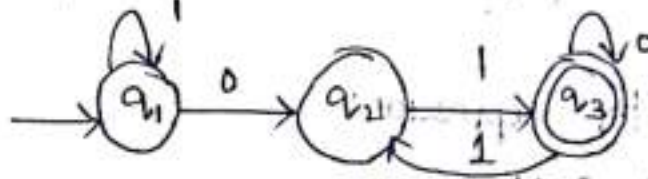
step 2:



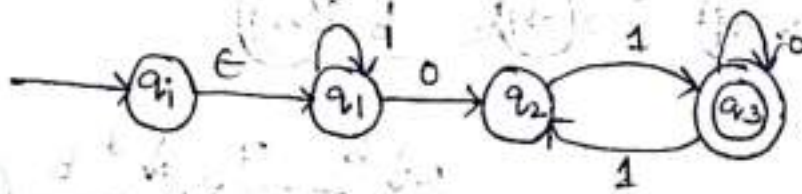
step 3: Eliminate the q_1 state



⇒ Convert the Given DFA into a regular expression.



Step 1: Since there is an incoming edge to initial, create a new state (q_i) and make it as initial state.

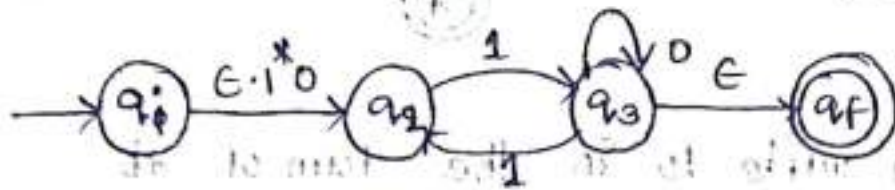


Step 2: There is only one final state.

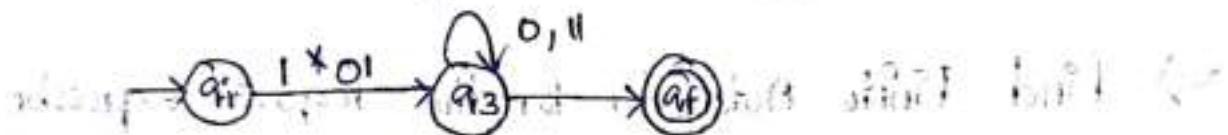
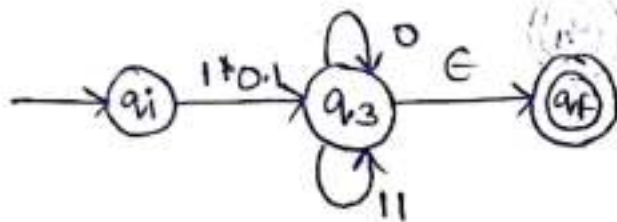
Step 3: Since there are outgoing edges to final state, create a new state and make it as final state.



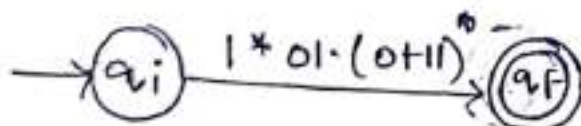
Step 4: Eliminate q_1 then the resultant DFA becomes



Eliminate q_2 then the resultant DFA becomes

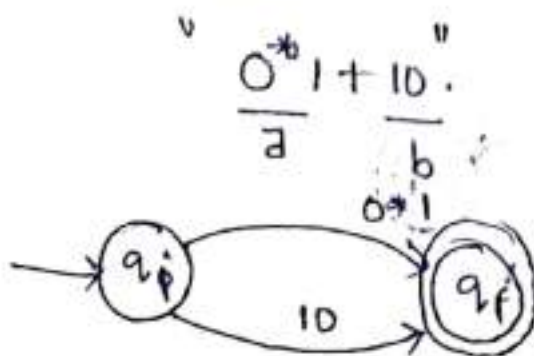


Eliminate q_3 then the resultant DFA becomes

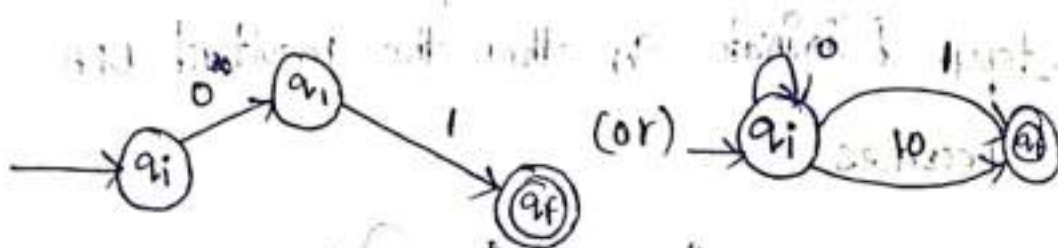


Regular Expression (RE) = $1^* 0 1 (0 + 1)^*$

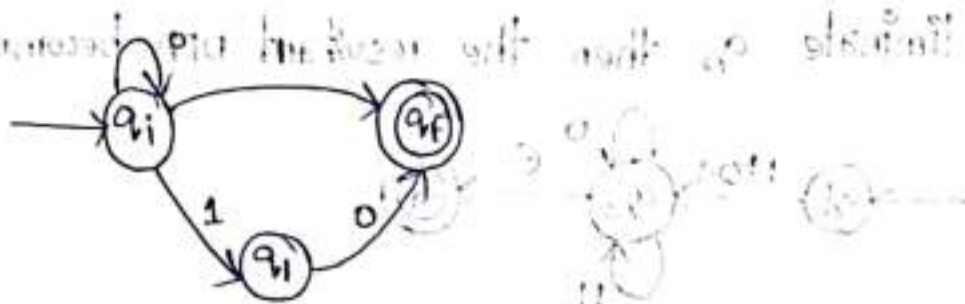
⇒ convert given regular expression into its equivalent finite Automata.



a is $\frac{0^*1}{a \ b}$ and it is in the form of ab

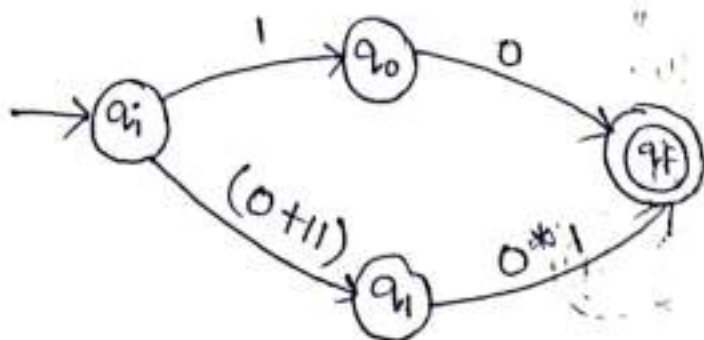
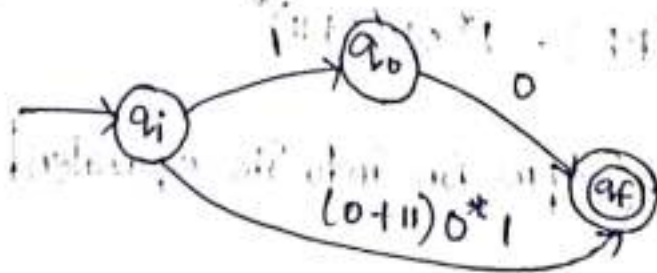
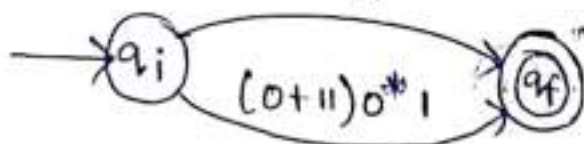


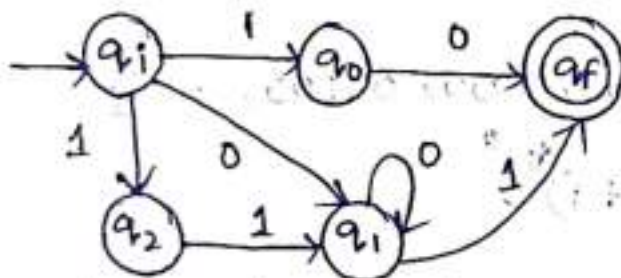
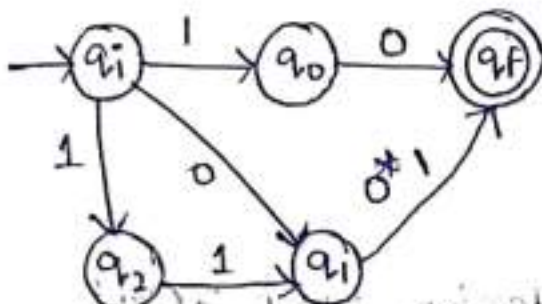
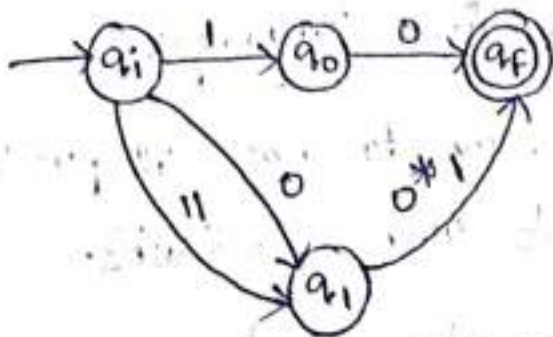
We can write 10 in the form of ab



$\Rightarrow$ Find Finite Automata for the regular expression

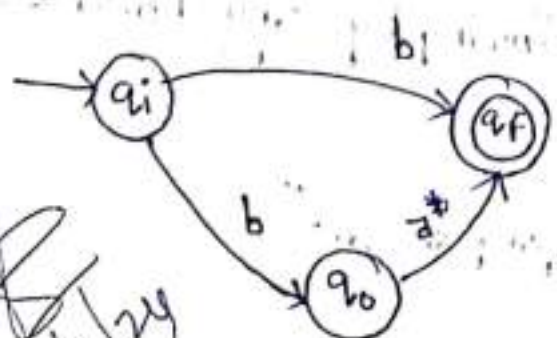
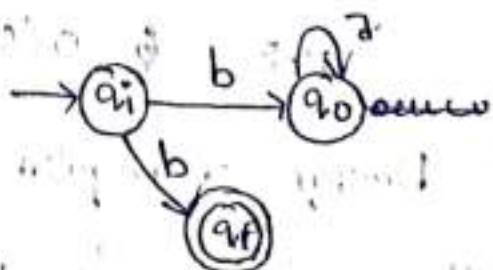
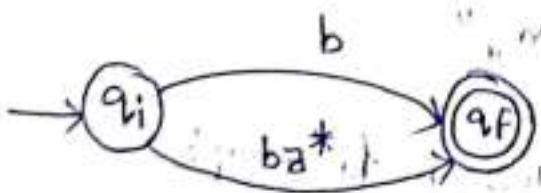
$$\frac{10 + (0+11)0^*1}{a \ b}$$





⇒ Find finite Automato for the regular expression

$$\frac{b}{a} + \frac{ba^*}{b}$$



6/3/24

Regular Expression for the Given languages:-

i) Design the regular expression, to accept all possible combinations of a's and b's over $\Sigma = \{a, b\}$.

Sol:- $L = \{ \epsilon, a, b, ab, aab, baa, aba, \dots \}$

" R.E = $(a|b)^*$ "

Transition diagram:



ii, Design a regular expression which accepts all strings which ends with "00".

Sol:- $L = \{ 00, 100, 1100, 0100, 01100, \dots \}$

" R.E = $(0+1)^*00$ "



iii, Design a regular expression for to accept all strings which starts with "0" and ends with "1".

Sol:- $L = \{ 0101, 01, 011, 01011, \dots \}$

" R.E = $0(0+1)^*1$ "

iv) Design a regular expression to accept any number of a's following by any no. of b's and any no. of c's.

Sol:- " R.E = $(a^*b^*c^*)$ "

$\Rightarrow$ v) at least one b's, one a's, one c's

$$R \cdot E = (a^* b^* c^*) [0 \vee 1] (a^* b^* c^*)$$

vi) begins with "00 (or) 11" and ends with "00 (or) 11"

$$\text{sol: } R \cdot E = (00 + 11)(0 + 1)^* (00 + 11)$$

vii) third character from the right end is "a"

$$\text{sol: } R \cdot E = (a + b)^* a (a + b)(a + b)^*$$

viii) design R.E for atleast 2 b's.

$$\text{sol: } RE = (a + b)^* b (a + b)^* b (a + b)^*$$

ix) exactly two "b's"

$$R \cdot E = a^* b a^* b a^*$$

x) Even length strings over $\Sigma = \{0\}$

$$R \cdot E = (00)^*$$

$$L = \{ \epsilon, 00, 0000, 000000, \dots \}$$

xi) odd length strings over $\Sigma = \{1\}$

$$R \cdot E = 1(11)^*$$

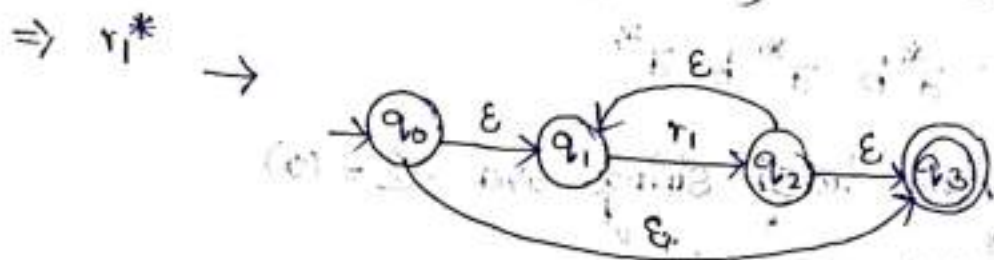
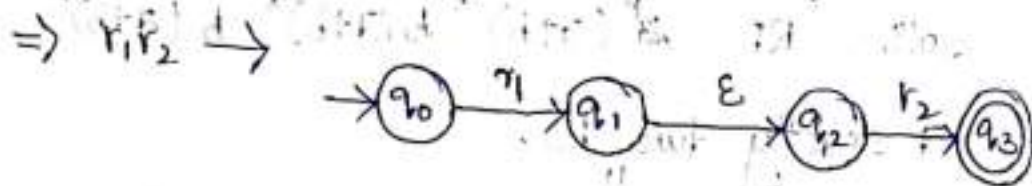
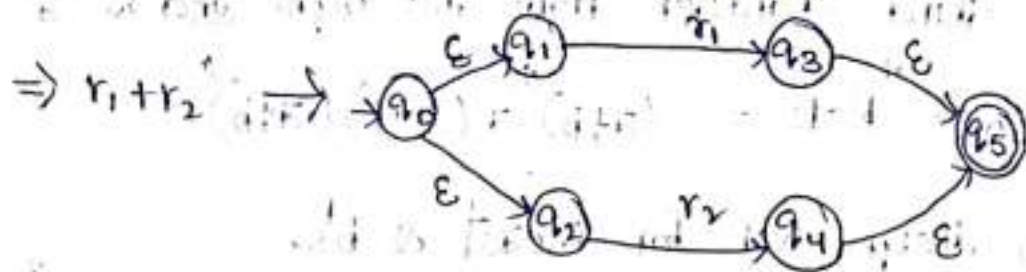
$$L = \{ \epsilon, 111, 1, 1111, 11111, \dots \}$$

Regular Expression to Finite Automata Using Subset

method:

Step 1: Convert the given expression into NFA with ϵ

Step 2: Convert NFA with ϵ into DFA.

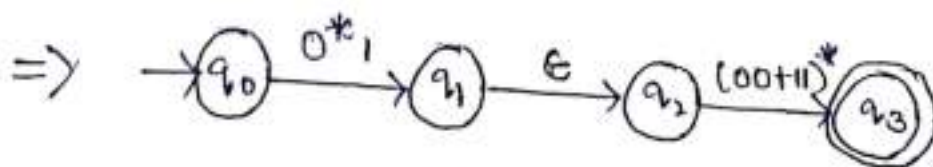


Q:- convert the given expression into the finite Automata using "subset" method.

$$\frac{0^* 1}{r_1} (00 + 11)^*_{r_2}$$

A:- This is in the form of $r_1 r_2$

where $r_1 = 0^* 1$, $r_2 = (00 + 11)^*$



$r_1 = 0^* 1$ Again this is in the form of $r_1 r_2 \Rightarrow r_1 = 0^*$, $r_2 = 1$

Applications of Regular Expressions:-

=> There are three applications of R.E

1) Regular Expressions in ^{unix} Utilities:-

→ There are various notations in unix

i) Any character [.]

ii) The sequence [1, 2, ..., 10]

=> addition = $1 + 2 + 3 + 4 + \dots + 10$

iii) [A-z], [a-z], [0-9] → A range of character

iv) [:digit:] → [0-9]

v) [:alpha:] → [A-z, a-z]

vi) [:alphanumeric:] → [A-z, a-z, 0-9]

vii) * → zero or many

viii) + → one or many

ix) ? → zero or one

x) | → concatenation

2) Regular Expressions in Lexical Analyser:-

→ One of the most important phase of lexical analyser is to scan the program and divide the program into individual units called tokens.

→ semantic analyser provides some tasks to the tokens.
→ First phase of the compiler it provides the program into tokens.

→ In tokens, there are identifier and keyword.

→ In unix, there are lex and flex they decide what the lexical analyser do in unix operating system.

g) Regular Expressions in finding patterns in the string:

→ To find the patterns in the string we use RE.

123A Main st.

$[0-9]^+[A-Z]?[A-Z][a-z]^*([A-Z](a-z)^*)^*$

(street | st | Road | Rd)

Algebraic laws of Regular Expression:

Associative laws for regular Expression RegEx

$$A + (B + C) = (A + B) + C \text{ and } A(B + C) = (A + B)C$$

Commutative laws for Regular Expression RegEx

$$A + B = B + A \text{ however } A \cdot B \neq B \cdot A$$

Identity for Regular Expression: RegEx

$$\phi + A = A + \phi = A$$

$$\epsilon \cdot A = A \cdot \epsilon = A$$

Distributivity for Regular Expression

$$A(B+C) = AB+AC \quad [\text{left}]$$

$$(B+C)A = BA+CA \quad [\text{right}]$$

$\rightarrow \phi$ is an annihilator

Idempotent $A+A=A$

closure laws for Regular Expression: RegEx

$$(A^*)^* = A^*, \phi^* = \epsilon, \epsilon^* = \epsilon, A^+ = AA^*, \text{ and } A^+ = A^*A$$

$$A^* = A^+ + \epsilon$$

DeMorgan's Law:

$$(L+B)^* = (L^*B^*)^*$$

Test for the Regular Expression whether $E=F$

1) convert the E and F into concrete R.E C & D

2) Test $L(C) = L(D)$ th. If language of C is equals to the language of D then we can say $E=F$ is TRUE otherwise FALSE.

pumping lemma:- $\rightarrow$ To prove the language is regular expression or not.

$\rightarrow$ repetition [Pumping] lemma[word] $\rightarrow$ It is entry point, using that we can ~~create~~ derive other words.

$\rightarrow$ Repetition of that word, by that word can derive the other words.

$\rightarrow$ There are two types of pumping lemma:

- 1) Regular Expression for pumping lemma
- 2) context-free languages

Theorem:-

let " L " be the Regular language, Then there exists a constant " c " such that for each and every string " w " in " L ".

$$|w| \geq c$$

We can break the w in three strings; " x, y, z "

- $|y| > 0$
- $|xy| \leq c$
- For all $k \geq 0$, the string xy^kz is also in L .

procedure to prove a language is not regular using pumping lemmas

1) consider L is a regular language

2) select constant 'c' such that $|w| \geq c$, for every string (w) in L .

3) Take a string (w) from L and divide w into x, y and z .

such that

$$|y| > 0$$

$$|xy| \leq c$$

for $k \geq 0$, xy^kz is in L .

4) If not, say L is a irregular.

problem:-

prove $L = \{a^m b^n \mid \gcd(m, n) = 1\}$ is not a regular

language.

Ans- $L = \{aabb, aaabbbb, aaaabbbbb, \dots\}$

1. Let L is a regular language

2. Constant $c = 5$, where $|w| \geq c$.

3. Let $w = \underline{aa} \underline{ab} \underline{bbb}$
 $\quad \quad \quad x \quad y \quad z$

divide w into x, y, z such that $|y| > 0$ & $|xy| \leq c$

then $x = aa$, $y = ab$, $z = bbb$

for every $k \geq 0$, compute xy^kz

let $k=0$,

$xy^0z \Rightarrow x \cdot y^0 \cdot z = xz = aabbb$, this is in d

let $k=1$,

$\Rightarrow xy^1z \Rightarrow x \cdot y^1 \cdot z = xyz = aaabbbb$ is in d

let $k=2$,

$\Rightarrow xy^2z \Rightarrow x \cdot y^2 \cdot z = xyyz = aaababbbb$ is not in d

$\therefore$ The given language d is irregular not reg.

C: prove $d = \{0^n \mid n \text{ is a perfect square}\}$ is not a regular language.

Ans: $d = \{0, 0000, 00000000, 0000000000000000\}$

1) let d is a regular language.

2) constant $c \in \mathbb{N}$, where $|w| \geq c$

3) select $w = 0000$,

divide w in xyz such that

$|y| > 0$ & $|xy| \leq c$ then

$x = \epsilon, y = 0, z = 000$

for every $k \geq 0$, compute xy^kz

let $k=0$,

$\Rightarrow xy^0z \Rightarrow xz = \epsilon \cdot 000 \Rightarrow 000$ is not in d

let $k=1$

$\Rightarrow xy^kz \Rightarrow xyz \Rightarrow \epsilon \cdot 0 \cdot 000 \Rightarrow 0000$ is in L .

let $k=2$

$\Rightarrow xy^kz = xyyz \Rightarrow \epsilon \cdot 0 \cdot 0 \cdot 0000 \Rightarrow 000000$ is not in L .

$\therefore$ The given language L is not regular.

Context-free Grammars:- (CFG) set of rules.

$\rightarrow$ Grammar is a ^{tuple} Quadruple Notations

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

$V \rightarrow$ set of non-terminals (or) variables, (uppercase letters)

$T \rightarrow$ set of terminals (lowercase letters)

$P \rightarrow$ set of productions ($\alpha \rightarrow \beta$) these are constructed to terminals & Non-Terminals ($\alpha = aA, \beta = bBa$)

$$\Rightarrow \alpha, \beta \in (V \cup T)^*$$

$S \rightarrow$ start symbol

Types of Grammars:- 4 types of Grammars

1) Type 0:- It is an unrestricted Grammar.

$\rightarrow$ The language generated by this kind of Grammar is "Recursively Enumerable Language (REL)".

$\rightarrow$ To represent this we use "Turing Machine (TM)".

Produ
 $\Rightarrow \alpha \rightarrow \beta$

where α has atleast one non-terminal

Context free Grammar (CFG):-

A context free grammar (G) is a 4 tuple notation, where $G = (V, T, P, S)$

where $V \rightarrow$ set of non-terminals or variables

$T \rightarrow$ set of terminals

$\rightarrow$ represented in uppercase letters.

$\rightarrow$ represented in lowercase letters.

$S \rightarrow$ start symbol

$P \rightarrow$ set of productions of the form $(\alpha \rightarrow \beta)$

where $\alpha \in V$, $|\alpha| = 1$ & (β can be a terminal or terminal followed by non-terminal or non-terminal followed by a terminal)

Derivations:-

$\rightarrow$ productions of the grammar are used to infer the strings of the language.

$\rightarrow$ There are two approaches

1) Recursive Inference

2) Derivations

1) Recursive Inference Approach:-

→ This is a bottom up approach (body to head).

→ We start with known symbols (terminals).

Grammar is:

$$1. E \rightarrow I$$

$$6. I \rightarrow a$$

$$2. E \rightarrow E + E$$

$$7. I \rightarrow I_0$$

$$3. E \rightarrow E * E$$

$$8. I \rightarrow I_1$$

$$4. E \rightarrow (E)$$

$$9. I \rightarrow I_a$$

$$5. I \rightarrow b$$

$$10. I \rightarrow I_b$$

→ The string to be derived is $a^*(a+boo)$

$$a^*(a+boo)$$

$$E * (E + I_0)$$

$$E * E$$

$$I * (a+boo)$$

$$E * (E + I_0)$$

$$E$$

$$E * (a+boo)$$

$$E * (E + I)$$

$$E * (I + boo)$$

$$E * (E + E)$$

$$E * (E + boo)$$

$$E * (E)$$

S.No	String Inferred	for language of	Productions used	Steps used
1)	a	I	6	-
2)	a	E	1	1)
3)	b	I	5	-
4)	bo	I	7	3)
5)	boo	I	7	4)

6) boo E 1, 5)

7) a+boo E 2, 6)

8) (a+boo) E 4, 7)

9) a*(a+boo) E 3, 2), 8)

2) Derivation Approach:-

→ This is a topdown approach.

→ let "G" be a "CFG",

$$\alpha A \beta \in (V \cup T)^*$$

Where, A is a non terminal and $A \rightarrow \gamma$ then

$$\alpha A \beta \xRightarrow{*}_G \alpha \gamma \beta$$

$$\Rightarrow a * (a + boo)$$

Grammar is:

1. $E \rightarrow I$

6. $I \rightarrow a$

2. $E \rightarrow E + E$

7. $I \rightarrow I 0$

3. $E \rightarrow E * E$

8. $I \rightarrow I 1$

4. $E \rightarrow (E)$

9. $I \rightarrow I a$

5. $I \rightarrow b$

10. $I \rightarrow I b$

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

$$\rightarrow a * (I + E) \rightarrow a * (a + E) \rightarrow a * (a + I) \rightarrow a * (a + I 0)$$

$$\rightarrow a * (a + I 0) \rightarrow a * (a + I 0 0) \rightarrow a * (a + boo)$$

Left Most Derivation And Right Most Derivation:

1) Left Most Derivation:

In which At each step of derivation

if we replace left most variable with its production body.

→ We represent leftmost derivation by $\xrightarrow{*}_{lm}$

2) Rightmost Derivation:

In which, at each step of derivation,

we replace rightmost variable with its production body.

→ we represent leftmost derivation by $\xrightarrow{*}_m$

Examples: $\rightarrow E$

$\rightarrow E * (I + boo)$

$\rightarrow E * E$

$\rightarrow E * (a + boo)$

$\rightarrow E * (E)$

$\rightarrow I * (a + boo)$

$\rightarrow E * (E + E)$

$\rightarrow a * (a + boo)$

$\rightarrow E * (E + I)$

$\rightarrow E * (E + I0)$

$\rightarrow E * (E + I00)$

$\rightarrow E * (E + boo)$

→ Derive abababa from the Given Grammar and leftmost derivation approach.

$$S \rightarrow sbs | a$$

case i)

$$\rightarrow S$$

$$\rightarrow sbs \quad \{S \rightarrow sbs\}$$

$$\rightarrow sbsbs \quad \{S \rightarrow sbs\}$$

$$\rightarrow sbsbsbs \quad \{S \rightarrow sbs\}$$

$$\rightarrow absbsbs \quad \{S \rightarrow a\}$$

$$\rightarrow ababsbs \quad \{S \rightarrow a\}$$

$$\rightarrow abababs \quad \{S \rightarrow a\}$$

$$\rightarrow abababa \quad \{S \rightarrow a\}$$

case ii)

$$\rightarrow S$$

$$\rightarrow sbs \quad \{S \rightarrow sbs\}$$

$$\rightarrow abs \quad \{S \rightarrow a\}$$

$$\rightarrow absbs \quad \{S \rightarrow sbs\}$$

$$\rightarrow absbsbs \quad \{S \rightarrow sbs\}$$

$$\rightarrow ababsbs \quad \{S \rightarrow a\}$$

$$\rightarrow abababs \quad \{S \rightarrow a\}$$

$$\rightarrow abababa \quad \{S \rightarrow a\}$$

Language of Context free Grammars

Let "G" be a context free Grammar, then its language is denoted by $L(G)$ and it is defined as " $L(G) = \{w \text{ in } T^* / S \xRightarrow[G]{*} w\}$ ".

→ The language of CFG is known as "context free language".

Sentential Form:-

Let "G" be a "CFG", $\alpha \in (V \cup T)^*$, S is a start symbol. If $S \xRightarrow[G]{*} \alpha$ then α is in "sentential form".

→ There are two types:

1) left^{most} sentential form:

If $S \xRightarrow[lm]{*} \alpha$ then α is leftmost sentential form.

2) Rightmost sentential form:

If $S \xRightarrow[rm]{*} \alpha$ then α is rightmost sentential form.

Parse tree / derivation tree / syntax tree:-

A parse tree is a tree which is constructed with the following conditions:

1) root node must be a start symbol.

2) Intermediate / Internal node, must be a variable (including starting symbol).

3) leaf node must be a terminal (or) ϵ

→ construct parse tree for the string

$a*(a+b)$ using the CFG $E \rightarrow E+E/E*E/(E)$

$(E)/I, I \rightarrow I_1I_2/I_1I_2I_3/b/a$

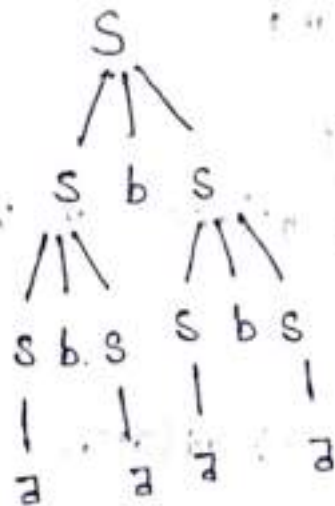
Ans-



→ construct parse tree for the string

abababa using the Grammar

$S \rightarrow sb^2/a$



Ambiguous Grammar:

A CFG "G" is said to be ambiguous for atleast one string in its language if there exist two distinct parse trees.

Example: Grammar is $E \rightarrow E + E, E \rightarrow E * E, E \rightarrow a$
construct "a + a * a" parse tree.



$\therefore$ The Given Grammar is ambiguous Grammar.

Removing Ambiguity:

Removing ambiguity is done

i) by adding precedence and

iv) by adding associativity.

→ To implement these, take an unambiguous grammar.

$$\Rightarrow a + a + a \Rightarrow (a + a) + a \Rightarrow a + (a + a)$$

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

$$\rightarrow E \rightarrow E * E \rightarrow F \rightarrow a$$

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

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

$$\rightarrow E \rightarrow T$$

$$\rightarrow F \rightarrow T \rightarrow F$$

$a + a + a$



Leftmost derivations in ambiguity:

→ There are more derivations in leftmost and Rightmost ambiguity.

→ There are only one derivation in the unambiguous grammar.

Inherent Ambiguous:-

A language is said to be ambiguity if all its
Grammer is ambiguous.

push down Automata (PDA):-

- If we add stack to the finite Automata then it will become Pushdown Automata.
- If we want to store the infinite information, we cannot store in finite automata then we use the push down Automata.

→ PDA = FA + stack

- This is used to represent the context free Grammar.
- In PDA, we use three components

1) input tape / input string

2) finite Control (Push, Pop)

3) stack

Formal Definition of PDA:-

- It is represented by "P", it is a 7 tuple notation

$$P = (Q, \Sigma, \Gamma, \delta, q_0, z_0, F)$$

where, $P = \text{PDA}$

$Q = \text{set of states}$

$\Sigma \rightarrow$ set of input symbols

$T \rightarrow$ set of stack elements

$\delta \rightarrow$ It is a transition function

$$\delta(q_0, a, x) = (q_n, \alpha)$$

$q \rightarrow$ current state

$a \rightarrow$ input to be processed

$x \rightarrow$ top of the stack

$q_n \rightarrow$ new state

$\alpha \rightarrow$ Symbol to be replaced in place

of x

$q_0 \rightarrow$ Initial state

$z_0 \rightarrow$ top of stack

$F \rightarrow$ Set of accepting states/Final states

Q:- Construct PDA for the given language

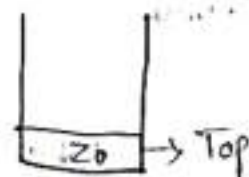
$$L = \{a^n b^n / n \geq 1\}$$

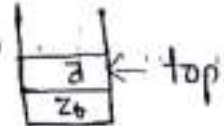
Sol:- $L = \{a^n b^n / n \geq 1\}$

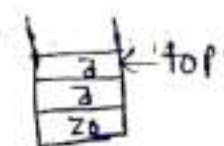
$$L = \{ab, aabb, aaabbb, \dots\}$$

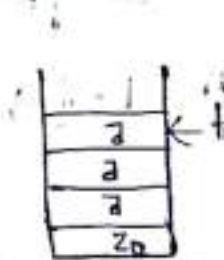
Let us assume "aaabbb" is our input tape/string,
 q_0 is the initial state.

Let z_0 be at the top of the stack.

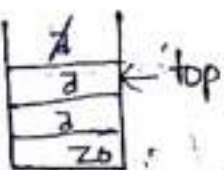


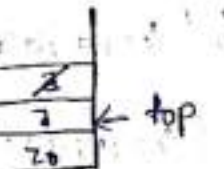
$$\delta(q_0, a, z_0) = (q_0, az_0)$$


$$\delta(q_0, a, a) = (q_0, aaa)$$


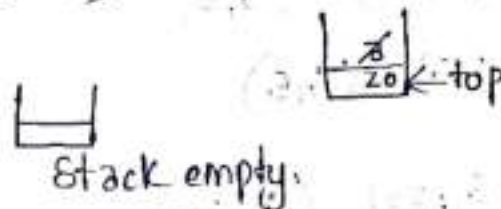
$$\delta(q_0, a, a) = (q_0, aaaa)$$


$$\delta(q_0, b, a) = (q_1, \epsilon)$$

$$\delta(q_1, b, a) = (q_1, \epsilon)$$


$$\delta(q_1, b, a) = (q_1, \epsilon)$$


$$\delta(q_1, \epsilon, z_0) = (q_2, \epsilon)$$

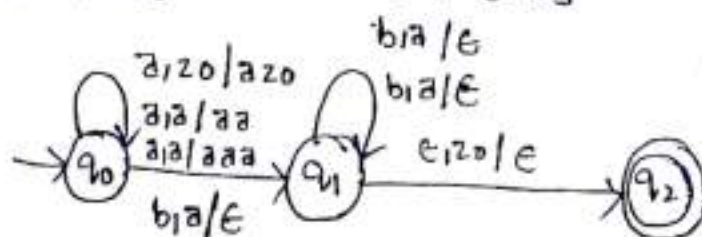


$$PDA = (Q, \Sigma, \Gamma, \delta, q_0, z_0, F)$$

$$Q \rightarrow \{q_0, q_1, q_2\} \quad \Gamma \rightarrow \{a, z_0\} \quad F \rightarrow \{q_2\}$$

$$\Sigma \rightarrow \{a, b\} \quad \delta$$

$$q_0 \rightarrow q_0 \quad z_0 \rightarrow \{z_0\}$$



Instantaneous Description of PDA:

→ To check the acceptance of string

find the acceptance of a string $aabb$ by the PDA $P = \{q_0, q_1, q_2\}$

$$P = [\{q_0, q_1, q_2\}, \{a, b\}, \{a, z_0\}, \delta, q_0, z_0, \{q_2\}]$$

$$\delta(q_0, a, z_0) = (q_0, az_0)$$

$$\delta(q_0, a, az_0) = (q_0, aaz_0)$$

$$\delta(q_0, a, aaz_0) = (q_0, aaaz_0)$$

$$\delta(q_0, b, aaaz_0) = (q_1, \epsilon)$$

$$\delta(q_0, b, aaz_0) = (q_1, \epsilon)$$

$$\delta(q_1, b, az_0) = (q_1, \epsilon)$$

$$\delta(q_1, \epsilon, z_0) = (q_2, z_0)$$

$$\delta(q_0, aabb, z_0) \vdash (q_0, aabb, az_0)$$

$$\vdash (q_0, bb, aaz_0)$$

$$\vdash (q_1, b, aa)$$

$$\vdash (q_1, \epsilon, az_0)$$

$$\vdash (q_2, z_0)$$

∴ the given string $aabb$ is accepted

by $\vdash (q_2, z_0)$ by the PDA

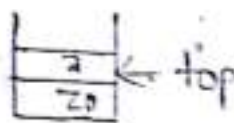
Construct PDA to accept the language $L = \{a^n b^n \mid n \geq 1\}$ and find acceptance of string "aaabbbb".

Δ :- push for "a" and pop for every two "b's".

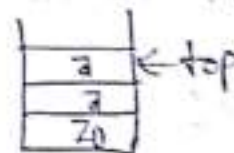
$L = \{abb, aabbbb, aaabbbbb, \dots\}$

Let $aaabbbb \in$ is the input tape, q_0 is the initial state and z_0 is at the top of stack

$\delta(q_0, a, z_0) = (q_0, az_0)$



$\delta(q_0, a, a) = (q_0, aaa)$



$\delta(q_0, b, a) = (q_1, aa)$

$\delta(q_1, b, a) = (q_2, \epsilon)$

$\delta(q_1, b, a) = (q_1, az_0)$

$\delta(q_1, b, a) = (q_2, \epsilon)$

$\delta(q_2, \epsilon, z_0) = (q_3, z_0)$

$P = [\{q_0, q_1, q_2, q_3\}, \{a, b\}, \{a, z_0\}, \delta, q_0, z_0, \{q_3\}]$



$$\delta(q_0, aaabbbbb, z_0) \vdash (q_0, aabbbbb, az_0)$$

$$\vdash (q_0, abbbbb, aa) \vdash (q_0, bbbbb, aaa)$$

$$\vdash (q_1, bbbbb, aaa) \vdash (q_2, bbbb, aa) \vdash (q_1, bbb, a)$$

$$\vdash (q_2, bb, a) \vdash (q_1, b, a) \vdash (q_2, \epsilon, z_0) \vdash (q_3, z_0)$$

$\therefore$ The given string $aaabbbbb$ is accepted

$\vdash (q_3, z_0)$ by the given PDA.

→ Construct PDA for the given language $L = \{a^n b^m c^n \mid m, n > 0\}$

Ans:- push for 'a's', Pop for 'c's' and 'b's' for no change.
 $m=2, n=3, m=1, n=2$

$$L = \{abc, abbc, aaabbbccc, aabccc, \dots\}$$

languages of PDA:-

1) Acceptance by Final state

$$L = \{w \mid \delta(q_0, w, z_0) \xrightarrow[p]{*} (q_f, \epsilon, \alpha)\}$$

continuation of $d = \{a^m b^n c^n \mid m, n \geq 0\}$

let $aabcc \in d$ as input tape, q_0 is the ^{initial} ~~decided~~ state and z_0 at the top of the stack.

$$\rightarrow \delta(q_0, a, z_0) = (q_0, az_0)$$

$$\rightarrow \delta(q_0, a, a) = (q_0, aa)$$

$$\rightarrow \delta(q_0, b, a) = (q_1, aa)$$

$$\rightarrow \delta(q_1, c, a) = (q_2, a)$$

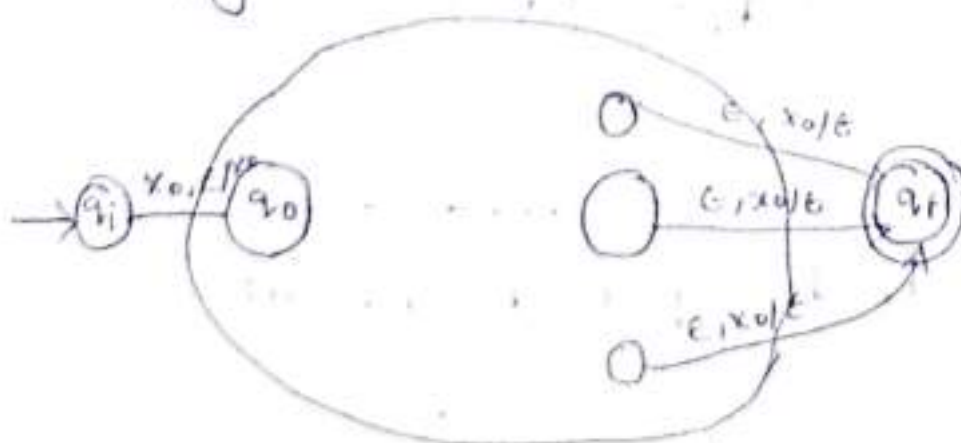
$$\rightarrow \delta(q_2, \epsilon, a) = (q_3, z_0)$$

$$\rightarrow \delta(q_3, \epsilon, z_0) = (q_3, z_0)$$

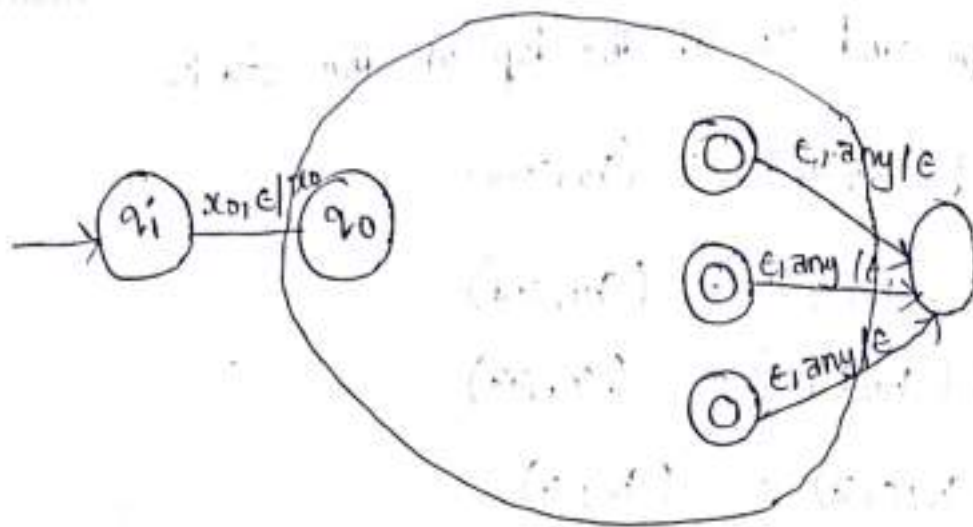
s) Acceptance by Empty stack:

$$N(P) = \{ w \mid \delta(q_0, w, z_0) \xrightarrow{*}_P (q, \epsilon, \epsilon) \}$$

2) From empty stack to final state:



4) From final state to empty stack



→ Equivalence of PDA and CFG

Conversion of CFG to PDA:-

A For non terminal: $A \rightarrow \alpha$

$$\delta(q_1, \epsilon, A) = (q_1, \alpha)$$

For terminal:

$$\delta(q_1, a, a) = (q_1, \epsilon)$$

d:

Example:-

convert the given CFG into PDA

$$S \rightarrow AB$$

$$A \rightarrow 0S/0$$

$$B \rightarrow 1S/1$$

$$\begin{array}{lll} S \rightarrow AB & B \rightarrow 1S & A \rightarrow 0 \\ A \rightarrow 0S & B \rightarrow 1 & \end{array}$$

consider $S \rightarrow AB$

$$\delta(q, \epsilon, S) = (q, AB)$$

$A \rightarrow 0S$

$$\delta(q, \epsilon, A) = (q, 0S)$$

$A \rightarrow 0$

$$\delta(q, \epsilon, A) = (q, 0)$$

$B \rightarrow 1S$

$$\delta(q, \epsilon, B) = (q, 1S)$$

$B \rightarrow 1$

$$\delta(q, \epsilon, B) = (q, 1)$$

$$\Rightarrow \delta(q, 0, 0) = (q, \epsilon)$$

$$\delta(q, 1, 1) = (q, \epsilon)$$

$$\delta(q, \epsilon, S) = \{(q, AB)\}$$

$$\delta(q, \epsilon, 0) = \{(q, 0S), (q, 0)\}$$

$$\delta(q, \epsilon, 1) = \{(q, 1S), (q, 1)\}$$

$$\delta(0, 0, 0) = \{(q, \epsilon)\}$$

$$\delta(1, 1, 1) = \{(q, \epsilon)\}$$

construct PDA

convert the given CFG to PDA:-

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

$$I \rightarrow I_0 / I_1 / I_a / I_b / a / b$$

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

For Non-Terminal $A \rightarrow \alpha$

$$\delta(q, \epsilon, A) = (q, \alpha)$$

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

$$\delta(q, \epsilon, E) = (q, E + E)$$

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

$$\delta(q, \epsilon, E) = (q, E * E)$$

$$\rightarrow E \rightarrow I$$

$$\delta(q, \epsilon, E) = (q, I)$$

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

$$\delta(q, \epsilon, E) = (q, (E))$$

$$\rightarrow I \rightarrow I_0$$

$$\delta(q, \epsilon, I) = (q, I_0)$$

$$\rightarrow I \rightarrow I_1$$

$$\delta(q, \epsilon, I) = (q, I_1)$$

$$\rightarrow I \rightarrow I_a$$

$$\delta(q, \epsilon, I) = (q, I_a)$$

$$\rightarrow I \rightarrow I_b$$

$$\delta(q, \epsilon, I) = (q, Ib)$$

$$\rightarrow I \rightarrow a$$

$$\delta(q, \epsilon, I) = (q, a)$$

$$\rightarrow I \rightarrow b$$

$$\delta(q, \epsilon, I) = (q, b)$$

For Terminal

$$T = \{a, b, +, *, 0, 1\}$$

$$\delta(q, a, a) = (q, \epsilon)$$

$$\rightarrow \delta(q, a, a) = (q, \epsilon)$$

$$\rightarrow \delta(q, b, b) = (q, \epsilon)$$

$$\rightarrow \delta(q, 0, 0) = (q, \epsilon)$$

$$\rightarrow \delta(q, 1, 1) = (q, \epsilon)$$

$$\rightarrow \delta(q, +, +) = (q, \epsilon)$$

$$\rightarrow \delta(q, *, *) = (q, \epsilon)$$

$$\Rightarrow \delta(q, \epsilon, \epsilon) = \{(q, \epsilon \epsilon), (q, \epsilon * \epsilon), (q, \epsilon I), (q, \epsilon)\}$$

$$\Rightarrow \delta(q, \epsilon, I) = \{(q, Ia), (q, Ib), (q, Ia), (q, Ib), (q, a), (q, b)\}$$

$$\therefore PDA = [\{q\}, \{a, b, 0, 1, +, *\}, \Gamma, \delta, q, z_0, \{q\}]$$

Conversion of PDA to CFG:- based on Transition functions

convert the given PDA to CFG

if there are no stack symbols - 1

$$\delta(q_0, a, z_0) = (q_0, az_0)$$

$$1 \rightarrow 2$$

$$\delta(q_0, a, a) = (q_0, aa)$$

$$2 \rightarrow 4$$

$$\delta(q_0, b, a) = (q_1, \epsilon)$$

$$PDA = \left[\{q_0, q_1\}, \{a, b\}, \{a, z_0\}, \delta, q_0, z_0, \{q_1\} \right]$$

$$\delta(q_0, b, a) = (q_1, \epsilon)$$

$$Q = \{q_0, q_1\}$$

$$\delta(q_0, \epsilon, a) = (q_1, \epsilon)$$

$$\therefore Q \Gamma Q$$

Ans:- Let us assume, s be the start symbol

find variables

$$V = \{s, [q_0 a q_0], [q_0 a q_1], [q_0 z_0 q_0], [q_0 z_0 q_1], [q_1 a q_0], [q_1 a q_1], [q_1 z_0 q_0], [q_1 z_0 q_1]\}$$

$$T = \{a, b\}$$

find productions:

$$\rightarrow \delta(q_0, a, z_0) = (q_0, az_0)$$

$$[q_0 a q_0] \rightarrow a [q_0 a q_0] [q_0 z_0 q_0]$$

$$[q_0 a q_1] \rightarrow a [q_0 a q_1] [q_1 z_0 q_1]$$

$$[q_0 z_0 q_0] \rightarrow a [q_0 a q_0] [q_0 z_0 q_0]$$

$$[q_0 z_0 q_1] \rightarrow a [q_0 a q_1] [q_1 z_0 q_1]$$

$$\rightarrow \delta(q_0, a, a) = (q_0, a, a)$$

$$[q_0 \ a \ q_0] \rightarrow a [q_0 \ a \ q_0] \ [q_0 \ a \ q_0]$$

$$[q_0 \ a \ q_0] \rightarrow a [q_0 \ a \ q_1] \ [q_1 \ a \ q_1]$$

$$[q_0 \ a \ q_1] \rightarrow a [q_0 \ a \ q_0] \ [q_0 \ a \ q_1]$$

$$[q_0 \ a \ q_1] \rightarrow a [q_0 \ a \ q_1] \ [q_0 \ a \ q_0]$$

$$\rightarrow \delta(q_0, b, a) = (q_1, \epsilon)$$

$$[q_0 \ a \ q_0] \rightarrow b [q_1 \ a \ q_0]$$

$$[q_0 \ a \ q_1] \rightarrow b [q_1 \ a \ q_1]$$

$$\rightarrow \delta(q_1, b, a) = (q_1, \epsilon)$$

$$[q_1 \ a \ q_1] \rightarrow b$$

$$S \rightarrow [q_0 \ z_0 \ q_0]$$

$$S \rightarrow [q_0 \ z_0 \ q_1]$$

$$\rightarrow \delta(q_1, \epsilon, z_0) = (q_1, \epsilon)$$

$$[q_1 \ z_0 \ q_1] \rightarrow \epsilon$$

Construct ^{CFG} ~~PDA~~ for the given PDA

$$\delta(q_1, 1, z_0) = (q_1, 1, z_0)$$

$$\delta(P, 0, z_0) = (q_1, z_0)$$

$$\delta(q_1, 1, x) = (q_1, 1, x)$$

$$\delta(q_1, 0, x) = (P, x)$$

$$P = \{q, P, q\}, \{0, 1\}, \{x, z_0\}, \\ \delta, q, z_0, \{P\}$$

$$\delta(q_1, \epsilon, x) = (q_1, \epsilon)$$

$$\delta(P, 1, x) = (P, \epsilon)$$

Sol:- let S be the start symbol

let $CFG = G = (V, \Sigma, P, S)$

Find Variables

$$V = \{S, [P \dot{x} P], [P x q], [q x P], [q x q], \\ [P z o P], [P z o q], [q z o P], [q z o q]\}$$

$$\Sigma = \{0, 1\}$$

$P \rightarrow$ Find productions: $S \rightarrow q z o q$ $S \rightarrow q z o P$

$$\rightarrow \delta(q, 1, z_0) = (q_1, x z_0)$$

$$[q z o q] \rightarrow 1 [q x q] [q z o q]$$

$$[q z o q] \rightarrow 1 [q x P] [P z o q]$$

$$[q z o P] \rightarrow 1 [q x q] [q z o P]$$

$$[q z o P] \rightarrow 1 [q x P] [P z o P]$$

$$\rightarrow \delta(q, 1, x) = (q, x x)$$

$$[q x q] \rightarrow 1 [q x q] [q x q]$$

$$[q x q] \rightarrow 1 [q x P] [P x q]$$

$$[q x P] \rightarrow 1 [q x q] [q x P]$$

$$[q x P] \rightarrow 1 [q x P] [P x P]$$

$$\rightarrow \delta(q, a, x) = (p, x)$$

$$[a \times a] \rightarrow^0 [p \times q]$$

$$[q \times p] \rightarrow^0 [p \times q]$$

$$\rightarrow \delta(q, \epsilon, x) = (q, \epsilon)$$

$$[q \times q] \rightarrow \epsilon$$

$$\rightarrow \delta(p, 1, x) = (p, \epsilon)$$

$$[p \times p] = \delta 1$$

$$\rightarrow \delta(p, 0, z_0) = (q, z_0)$$

$$[p \ z_0 \ q] = z_0 \circ [q \ z_0 \ q]$$

$$[p \ z_0 \ p] \rightarrow^0 [q \ z_0 \ p]$$

Turing Machine:

Turing machine consists of a finite control, which can be any finite set of states.

→ There is a tape divided into cells, each cell can hold finite number of symbols.

Finite
control

... B B x_1 x_2 ... x_n B B ...

$x_1, x_2, \dots, x_n \rightarrow$ Input Symbols

B $\rightarrow$ Blank symbol

- choose input string from the input alphabet and place in tape.
- Extend the length of the tape to the left and right and place blank symbol over the extended cells.
- The Blank is a tape but not input symbol.
- There is a read write head (or) tape head; Initially the tape head is at leftmost cell that holds the input.
- A move of the turing machine is a function of the state of the finite control and tape symbol scanned.
- In one move a turing machine will
 - 1) change state: The next state optionally may be same as the current state.
 - 2) Write a tape symbol to the cell scanned, this tape symbol replaces the symbol in that cell.
 - 3) move the tape head to left or right but don't allow the head to remain stationary.

Defination:-

Turing machine is a " γ tuple" notation, And the symbol is $M = (Q, \Sigma, \Gamma, \delta, q_0, B, F)$.

$\Gamma \rightarrow$ input tape / Tape symbols

$B \rightarrow$ Blank symbol

$Q \rightarrow$ set of states

$\Sigma \rightarrow$ set of input symbols

$q_0 \rightarrow$ Initial/Starting state

$F \rightarrow$ Final state set

$\delta \rightarrow$ Transition Function

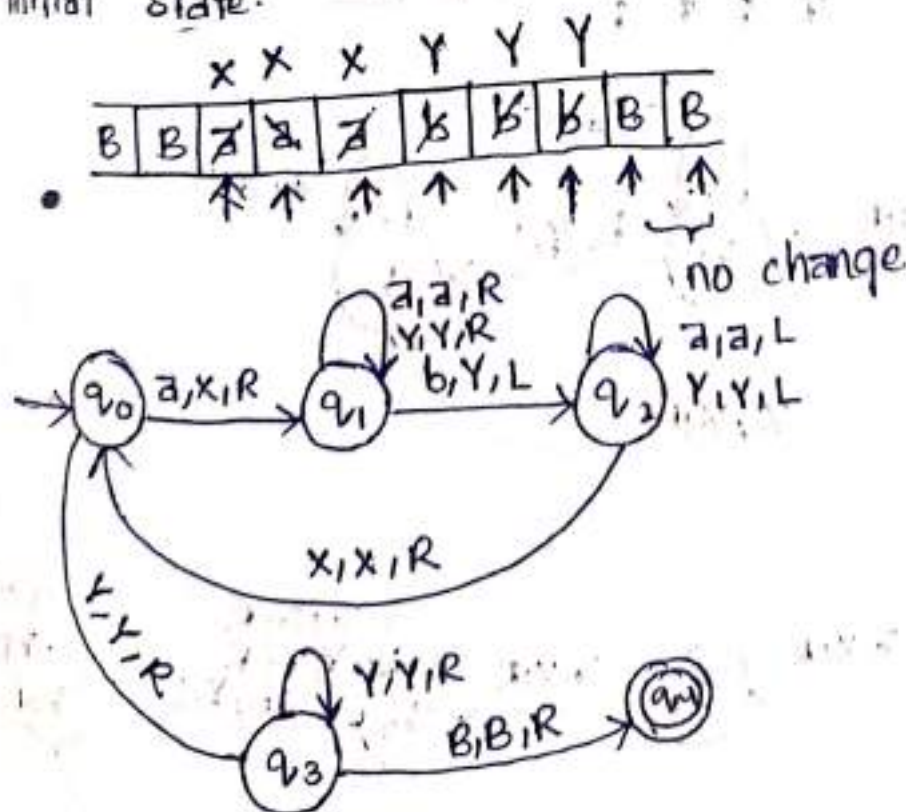
where, the transition Function maps to
 $Q \times \Sigma \Rightarrow Q \times \Gamma \times \{L, R\}$

Transition diagram in Turing machine:

Q Design Turing machine for the language $L = \{a^n b^n \mid n > 0\}$.

A:- $L = \{ab, aabb, aaabbb, \dots\}$

let "aaabbb" be the input tape, q_0 be the initial state.



$$M = (Q, \Sigma, \Gamma, \delta, q_0, B, F)$$

$$Q = \{q_0, q_1, q_2, q_3, q_4\}$$

$$\Sigma = \{a, b\}$$

$$\Gamma = \{a, b, x, y, B\}$$

$$F = \{q_4\}$$

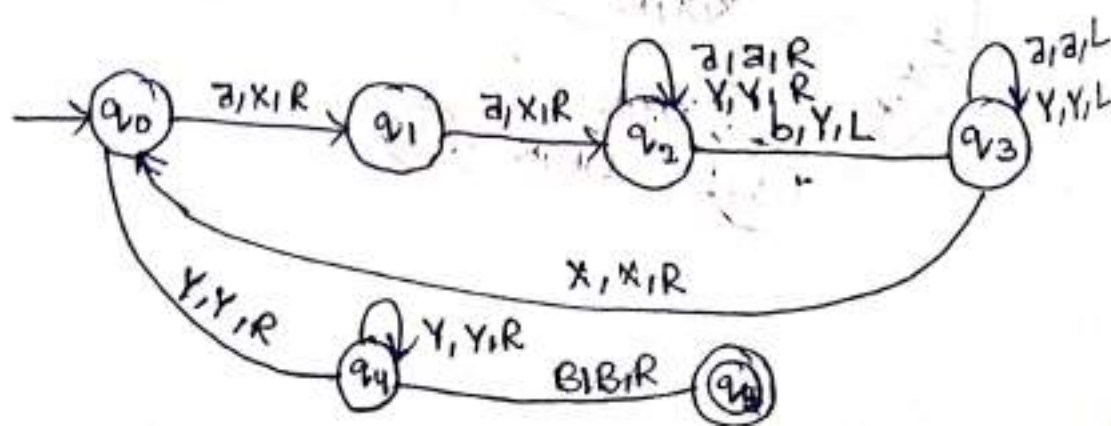
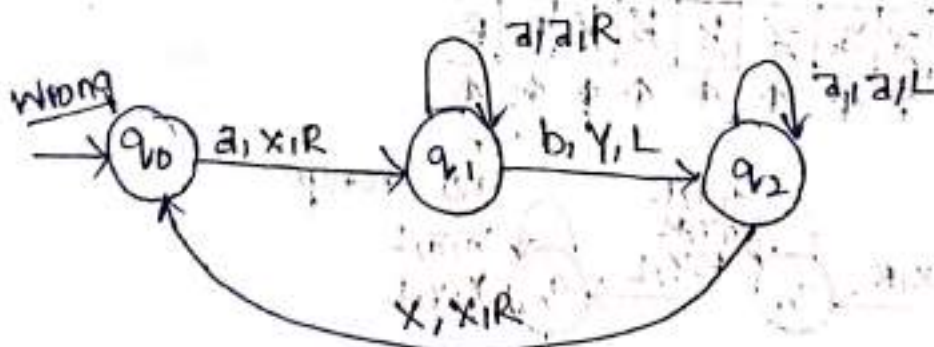
② Design Turing machine for the language

$$L = \{a^n b^n \mid n \geq 1\}$$

A:- $L = \{aab, aaabbb, aaaaaabbbb, \dots\}$

Let "aaabbb" be the input tape and q_0 be the initial state. [∵ For 2 a's we need to change to x.]

		x	x	x	x	y	y		
B	B	a	a	a	a	b	b	B	B
		↑	↑	↑	↑				



$$M = (Q, \Sigma, \Gamma, \delta, q_0, B, F)$$

$$Q = \{q_0, q_1, q_2, q_3, q_4, q_5\}$$

$$\Sigma = \{a, b\}$$

$$\Gamma = \{a, b, X, Y, B\}$$

$$F = \{q_5\}$$

Transition Table:-

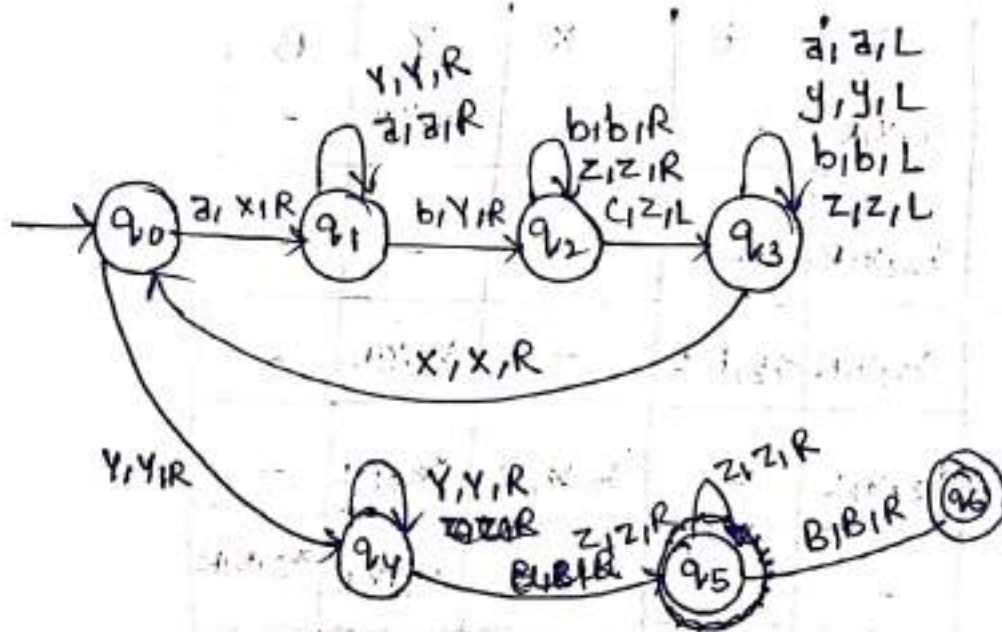
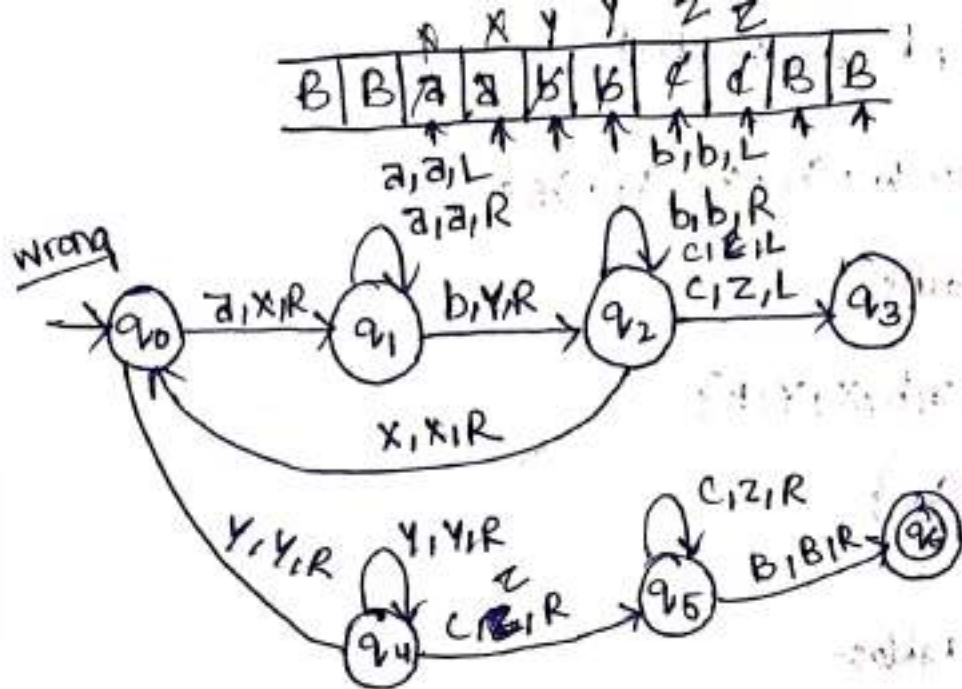
	a	b	X	Y	B
q_0	q_1, X, R	-	-	q_4, Y, R	-
q_1	q_2, X, R	-	-	-	-
q_2	q_2, a, R	q_3, b, L	-	q_2, Y, R	-
q_3	q_3, a, L	-	q_0, X, R	q_3, Y, L	-
q_4	-	-	-	q_4, Y, R	q_5, B, R
q_5	-	-	-	-	-

③ Design Turing machine for the given language

$$L = \{a^n b^n c^n \mid n > 0\}$$

A:- $L = \{abc, aabbcc, aaabbbccc, \dots\}$

let "aabbcc" be the input tape and q_0 be the initial state.



$$M = (Q, \Sigma, \Gamma, \delta, q_0, B, F)$$

$$Q = \{q_0, q_1, q_2, q_3, q_4, q_5\}$$

$$\Sigma = \{a, b, c\}$$

$$\Gamma = \{a, b, c, X, Y, B\}$$

$$F = \{q_6\}$$

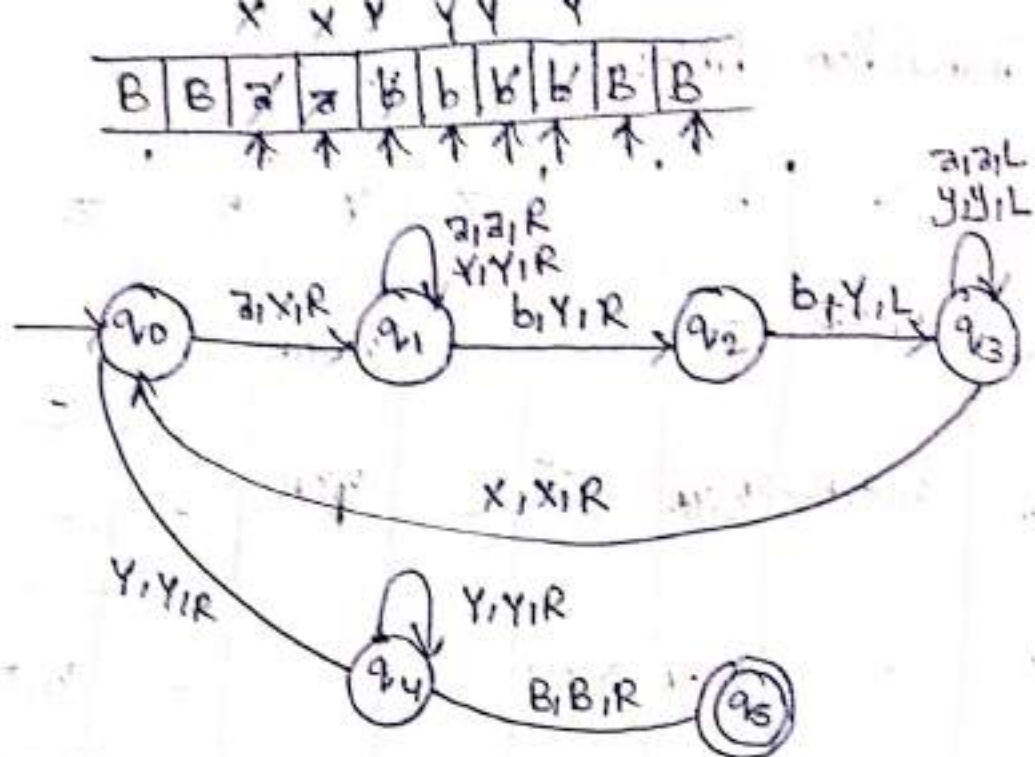
Transition Table:

	a	b	c	x	y	B	z
q_0	q_1, x, R	-	-	-	q_4, y, R	-	-
q_1	q_{11}, a, R	q_2, y, R	-	-	q_2, y, R	-	-
q_2	-	q_{21}, b, R	q_{21}, z, L	-	-	-	q_{12}, z, R
q_3	q_3, a, L	q_3, b, L	-	q_0, x, R	q_{21}, y, L	-	q_3, z, L
q_4	-	-	-	-	q_{41}, y, R	-	q_5, z, R
q_5	-	-	-	-	-	q_6, B, R	q_5, z, R
q_6	-	-	-	-	-	-	-

④ Design Turing machine for the language $L = \{a^n b^{2n} / n \geq 1\}$

A:- $L = \{abb, aabbbb, aaabbbbbbb, \dots\}$

let "aabbbb" be the input tape and q_0 be the initial state [$\because$ for 2 b's we need to change to Y]



$$M = (Q, \Sigma, \Gamma, \delta, q_0, B, F)$$

$$Q = \{q_0, q_1, q_2, q_3, q_4, q_5\}$$

$$\Gamma = \{a, b, x, y, B\}$$

$$\Sigma = \{a, b\}$$

$$F = \{q_5\}$$

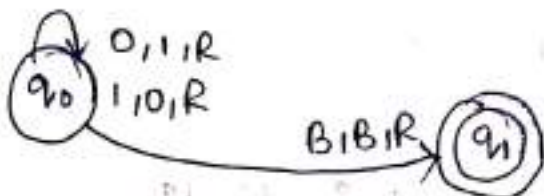
Transition Table:

	a	b	x	y	B
q_0	q_1, x, R	-	-	q_4, y, R	-
q_1	q_1, a, R	q_2, y, R	-	q_1, y, R	-
q_2	-	q_2, y, L	-	-	-
q_3	q_3, a, L	-	q_0, x, R	q_3, y, L	-
q_4	-	-	-	q_4, y, R	q_5, B, R
q_5	-	-	-	-	-

- 5) construct Turing machine for
 i) One's complement ii) Two's complement.

Ans: One's complement

B | 1 | 0 | 1 | 0 | B



⇒ Two's complement

1 0 1 0 1

0 1 0 1 0

0 1 0 1 1

- ii) Two's complement

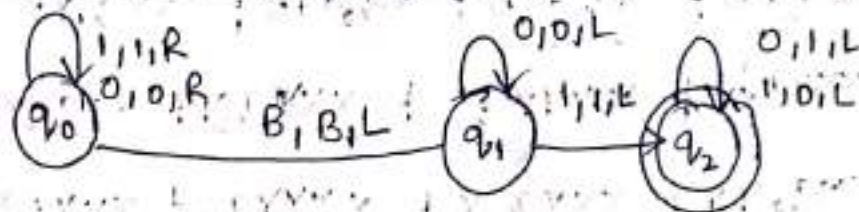
1 0 1 0

0 1 0 1

0 1 1 0

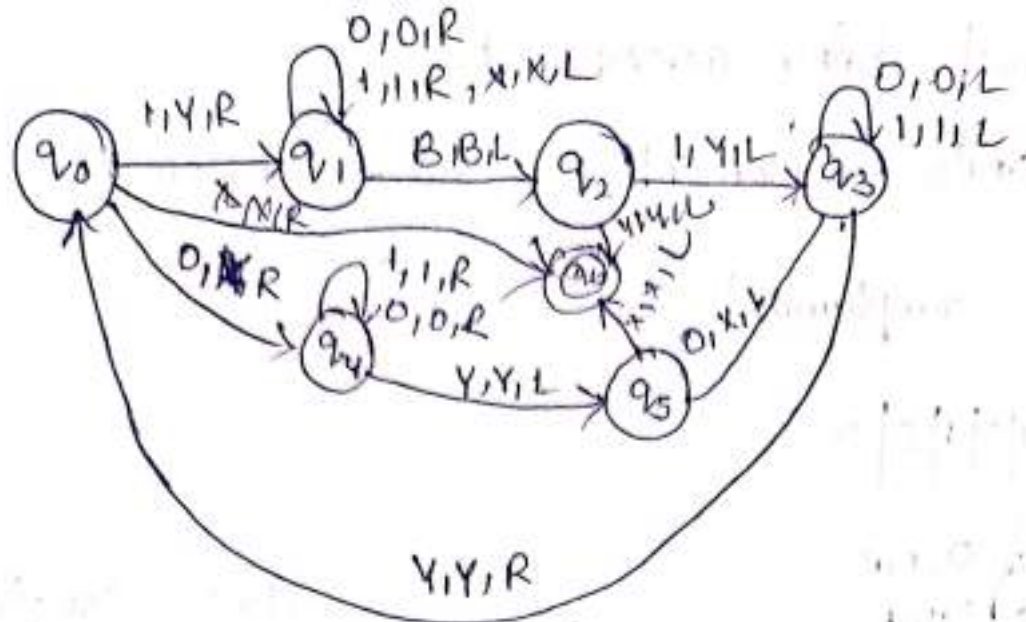
B | 0 | 1 | 1 | 0 | B

B | 0 | 1 | 0 | 1 | 1 | B



- 6) construct Turing machine for the palindrome of the 0's and 1's.

0 1 0, 0 0 0, 0 1 0 1 0 1, 1 1



Instantaneous Description For Turing Machine:-

Let "aaabbb" be the input string/tape.

Acceptance For the Turing machine:

$w = aabb$

$q_0 a a b b \vdash x q_1 a b b \vdash x a q_1 b b \vdash x a y q_2 b \vdash$

$x a q_2 y b \vdash x q_2 a y b \vdash x q_2 a y b \vdash x x q_1 y b \vdash$

$x x y q_1 b \vdash x x y y q_2 \vdash x x y q_2 y \vdash x x q_3 y y \vdash$

$x x q_0 y y \vdash x x y q_3 y \vdash x x y y q_3 \vdash x x y y q_4$

$\therefore$ The Acceptance of Turing machine is verified.

Language of Turing machine:

Language of Turing machine M is denoted by $L(M)$ and it is defined as

$$L(M) = \{w \text{ in } \Sigma^* / q_0 w \vdash^* \alpha p \beta\}$$

$\Rightarrow \alpha, \beta$ are any strings of tape symbols

$\Rightarrow p$ can be any state.

Undecidability :-

Recursive language: A language, if there exists a Turing machine, that accepts all strings in L and rejects all strings not in L .

$$L = \{a^n b^n\}$$
$$= \{ab, aabb, aaabbb, \dots\}$$

Let "abab" be the input tape symbol, if we find the "Acceptance" for this, this is not present in the language then it is a "recursive language".

Recursively Enumerable languages: If L is a language, if there exists a Turing machine that accepts all strings in L that may or may not reject all strings in the L .

$\rightarrow$ These are said to be Partially decidable.

⇒ The

Decidability:-

The recursive languages are said to be decidable.

Undecidability:-

A language which is not decidable is said to be undecidable.

→ There is no Turing machine for the undecidable languages.

→ The languages which are not recursively enumerable.

language that is not recursively enumerable:-

1) In order to prove that not a recursively enumerable we need to

i) Enumerating binary strings:- In order to prove this we need to write "binary strings".

$\{ \epsilon, 0, 1, 00, 01, 10, 11, 000, 001, \dots \}$

Here " ϵ " is the "0th" string

"0" is the "1st" string and so on.

Constructing codes for Turing machine:-

→ Assign positive integers to the symbols of transition in Turing machine.

→ For input symbols, Assign - the / Take the $x_1(0)$, $x_2(1)$ and $x_3(8)$.

→ For directions, Take the $D_1(L)$, $D_2(R)$.

Diagonalized language: It is represented by " d_d ".

It is a set of strings " w_i " and " w_i " is not in Turing machine $(or) L(M_i)$ for some M_i .

Table of Acceptance: We have write "0 and 1" in the acceptance table. Rows we take Turing machine and for columns we take strings.

The strings must be accepted by the Turing machine.

→ If Turing machine accepts take "1", if not accepts we take "0".

→ $L(M_i)$ will have the diagonal of w_i and M_i .

Proof of the language that is not recursively enumerable (Diagonalized language):

Let d_d be $L(M_i)$ for some M_i .

1) If w_i is in d_d , then w_i is in $L(M_i)$.

But, by the definition d_d , w_i is not in the ~~language~~ $L(M_i)$.

2) If w_i is not there in L_d , then it is not there in $L(M_i)$.

But by the definition of L_d , w_i is in the L_d .

$\therefore$ There is no Turing machine for the L_d .

So that we can have Turing machine for L_d , which means L_d is not recursively enumerable language. This proof is called as "proof by contradiction". This means: Firstly we assume that there will be the L_d later that we assume there will not have L_d .

Undecidable problem that is not recursively
enumerable:

Complements of recursive and recursively enumerable languages:

Theorem 1: If L is recursive then, $\bar{L}$ is also recursive.

proof:- let " L " be a recursive language.

so, " L " halts always for Turing machine " M ".

→ Let M be the Turing machine of A .

~~To some changes~~ To get M' , do some changes:

→ Make all accepting states as non-accepting states.

→ Take a new state as a accepting state and write transitions from each non-accepting state of M on each tape symbol to the accepting state.

The Turing machine M' accepts strings which are not accepted by M and rejects all strings which are accepted by M .

That means, the Turing machine M' always halts.

A is recursive.

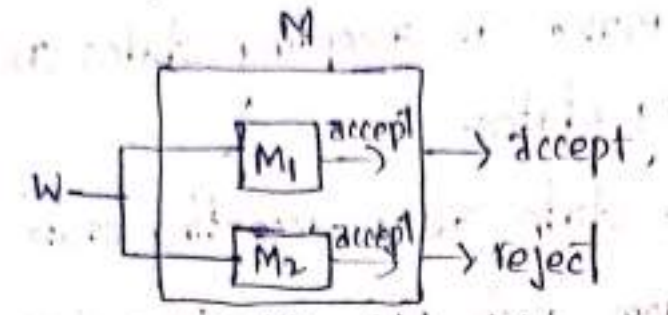
Theorem 2: If A and A' are recursively enumerable then A is recursive.

Proof: Take a Turing Machine M .

(Let A be recursive)

Let M_1 be the Turing machine for A and M_2 be the Turing machine for A' .

we can have a Turing machine 'M' by taking M_1 & M_2 as tapes.



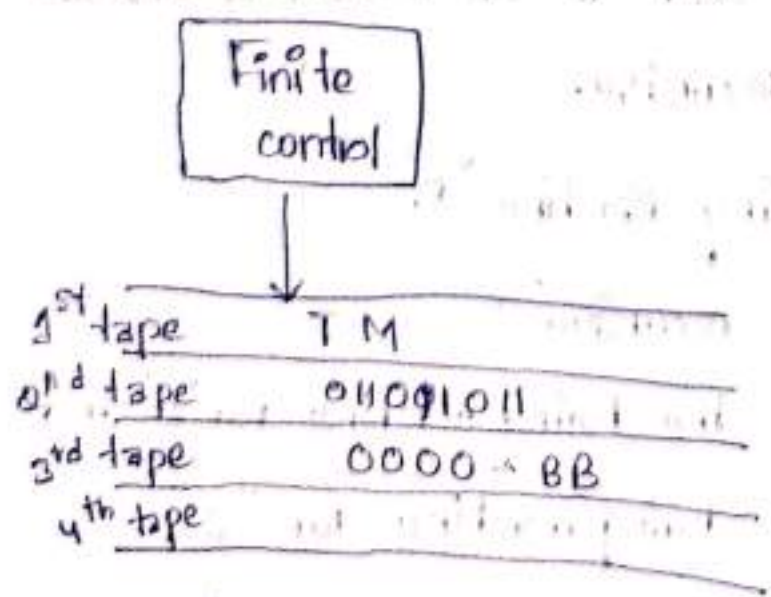
"M" acts same as "M"

Since "M" always halts

"L" is recursive

Universal language:- set of

It is a language of strings representing Turing machine and an input that is accepted by the Turing machine.



→ 1st tape holds TM, w

→ 2nd tape holds simulated code of TM

→ 3rd tape holds the state q_i in form.

→ 4th tape is used to scratch.

→ It is represented by the 'du'.

① Examine input to ensure that the code of Turing machine is legitimate.

② place input in 2nd tape in encoded form i.e., place 10 for each '0', 100 for each 1 and 1000 for each 'B'.

③ place start state (q_0) in tape 3, and also place tape head at 1st symbol of simulated in 2nd tape.

④ for each move of TM of the form

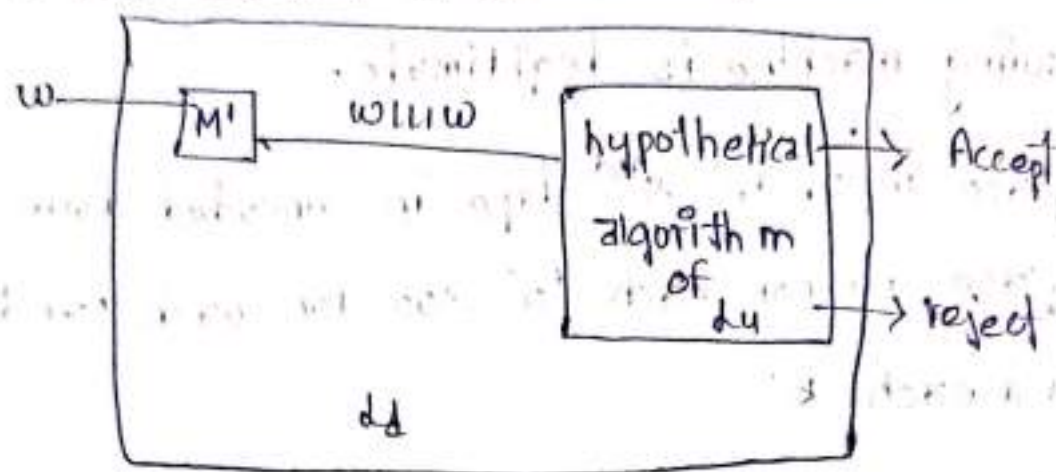
$$S(q_i, x_j) = (q_k, x_L, D_m)$$

a) replace the contents x_j with x_L

b) change the content of tape 3 by places q_k and make all previous q_i 's 10's.

Theorem: If L_u is RE, then L_u is not recursive

Proof: Let L_u is Recursively Enumerable, then there exists a Turing Machine M .



If there exist TM for L_u then TM can also possible for L_d .

As per the definition of L_d , there exist no TM for L_d .

So that there exist no TM for L_u i.e.

L_u is not recursively enumerable.

So that L_u is not recursive.

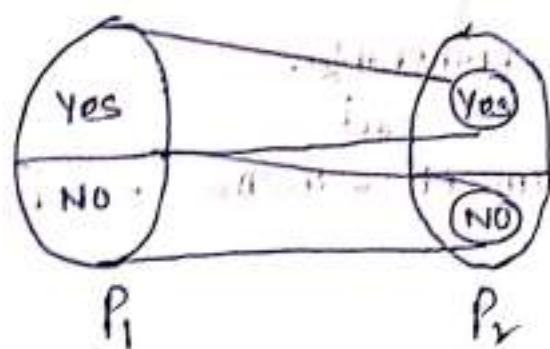
Undecidability about Turing machine:

Turing Reducability: It is a technique to convert one language into another language, when L_1 is recursive in L_2 .

→ In general an instance of, If there is an algorithm to convert P_1 to P_2 then we can say that an instance of P_1 is reduced to an instance of P_2 with the same answer.

→ If P_1 is (recursively enumerable) is not recursive then P_2 is not recursive.

→ If P_2 is not recursive then P_1 can't be recursive.



Reduction turns the positive.

Reduction from P_1 to P_2 is TM which take an instance of P_1 and accepts instance of P_2 .

In General, The reduction from P_1 to P_2 is program which takes instance of P_1 as Input.

Theorems- If there is reduction from P_1 to P_2

a) If P_1 is undecidable, then P_2 is also undecidable

Proof:- Let P_1 is undecidable and then exists an algorithm to solve P_2

So, P_2 is decidable, but P_1 is reduced to P_2

Since, there is reduction from P_1 to P_2 and P_1 is undecidable.

So that P_2 is also undecidable.

b) If P_1 is non-recursively enumerable then P_2 is also non-recursively enumerable.

Proof:- Let P_1 be form RE and there exists an algorithm for P_2

Since there is a reduction from P_1 to P_2 and P_1 is not RE P_2 is non RE.

Turing Machine for empty language:

$\rightarrow$ empty language

$\rightarrow$ non-empty language

$$L_e = \{ M / L(M) \neq \emptyset \}$$

$$L_{ne} = \{ M / L(M) = \emptyset \}$$

$\rightarrow$ L_e and L_{ne} are complement to each other.

→ L_n is recursively enumerable, and L_n is not recursive.
→ L_n is not recursively enumerable.

Rice Theorem and Property of RE languages:-

Property of RE languages:- It is a set of RE languages. There are two properties:

① Trivial Property of RE language: A property said to be Trivial, when it is has no strings or when it is empty or a set of all RE languages.

*② Non-Trivial Property of RE languages: A property said to be non-trivial, when it is a non-empty not a set of all RE languages.

Rice Theorem:-

For Every non-Trivial property of a recursively enumerable language is undecidable.

Implementation of Turing machine specifications:-

Applications of (Proof) Rice Theorem:-

① The Rice theorem is used to say the given problems are undecidable or not.

→ There are four problems

① whether the language accepted by TM is empty

② whether the language accepted by TM is finite.

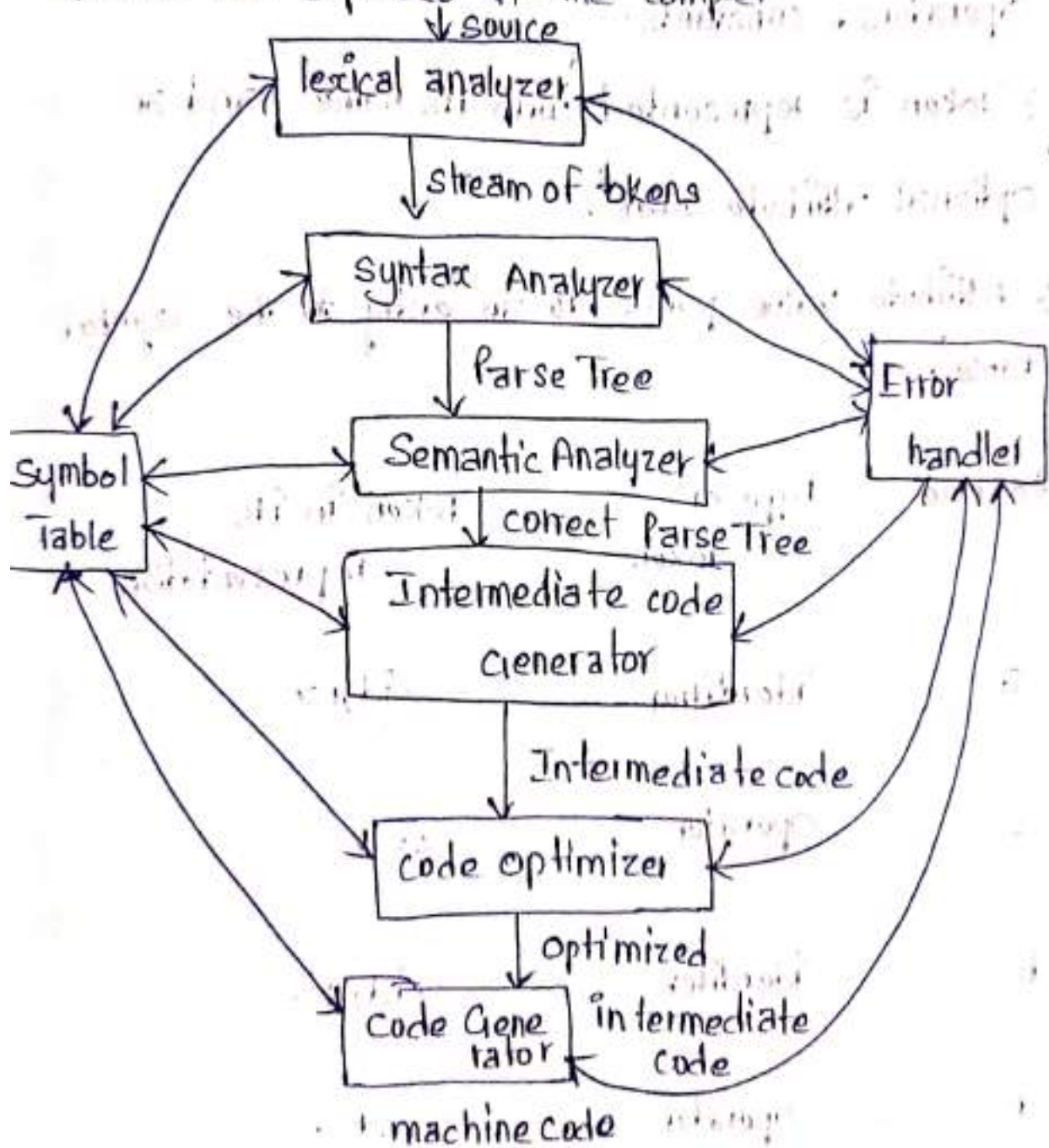
③ whether the language accepted by TM is regular

④ whether the language accepted by TM is context free Grammar or not.

Compiler: It is a system software to convert high level language to the machine code.

The structure of Compiler:-

→ There are 6 phases in the compiler.



① Lexical Analyzer: It reads the source program character by character and it groups the characters into a meaningful sequences called "lexemes". These are also called as "tokens".

Example: keywords, variable names (Identifiers), Operators, constants.

→ Token is represented "with its name", and an "Optional attribute value".

→ Attribute value points to an "entry" to the symbol table".

Lexeme	Type of token	Token in its representation
--------	---------------	-----------------------------

a	Identifier	$\langle id, 1 \rangle$
---	------------	-------------------------

=	Operator	$\langle = \rangle$
---	----------	---------------------

b	Identifier	$\langle id, 2 \rangle$
---	------------	-------------------------

+	Operator	$\langle + \rangle$
---	----------	---------------------

c	Identifier	$\langle id, 3 \rangle$
---	------------	-------------------------

*	operator	$\langle * \rangle$
---	----------	---------------------

60	Constant	60
----	----------	----

→ After converting it into tokens, it is represented

as $a = b + c * 60;$

" $\langle id1 \rangle = \langle id2 \rangle + \langle id3 \rangle * 60$ "

② Syntax Analyzer: It takes stream of tokens as the input and it gives parse tree as the output.

→ It checks the syntax of the program.

→ It is also called as "parser".

→ If there are any errors, it reports / informs to the same to the error handler.

→ The main task of the syntax analyzer is "Parse Tree" construction.

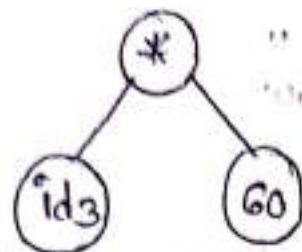
Stream of tokens $\Rightarrow \langle id1 \rangle = \langle id2 \rangle + \langle id3 \rangle * 60;$

→ Construction of parse tree will be based on the "precedence / priority of the operators".

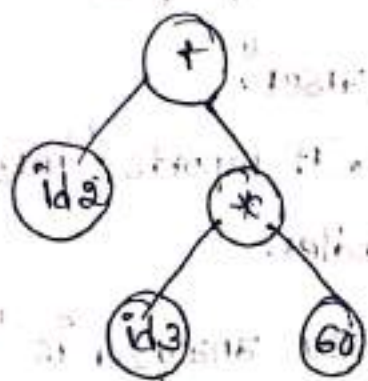
→ In parse tree take operators as "internal nodes" and "root" node.

→ It takes identifiers, constants as "leaf nodes".

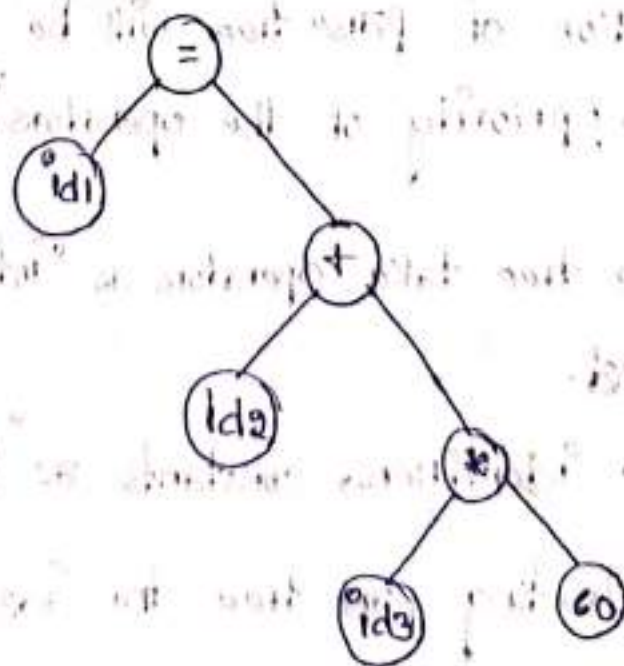
Step 1: constructing parse tree for $\langle id3 \rangle * 60$



Step2: Next we will construct the parse tree for the 'id2'.

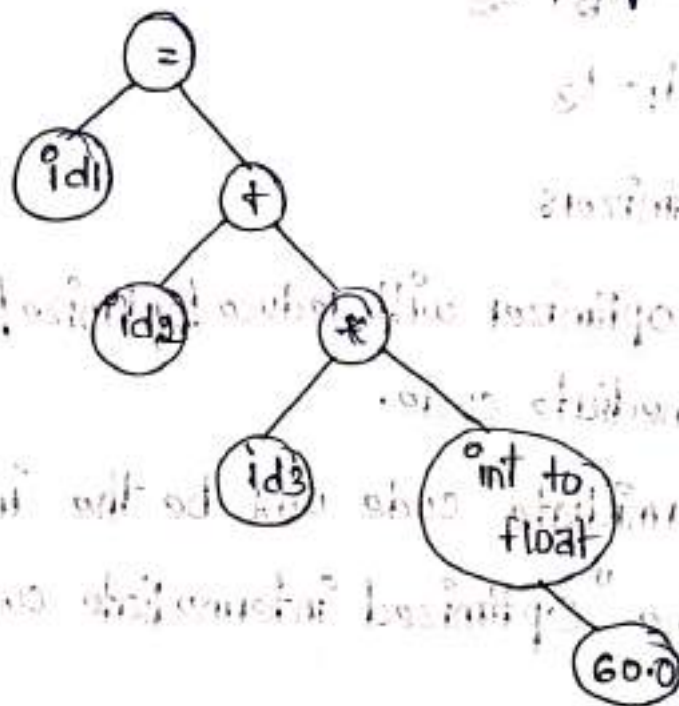


Step3:- Next we will construct the parse tree for the 'id1'.



→ It is constructed from top down to top.

- ③ Semantic Analyzer: It takes the parse tree as the input and generates the correct parse tree.
- semantic means "meaningful".
 - It uses the type checker to (make) check the type of identifiers and to make automatic conversion if there is any need.
 - let us assume a, b & c are floating point variables



- ④ Intermediate Code Generator: It takes the correct parse tree as input and generates the "intermediate code".
- To represent this, there are three address code.
 - ① Instruction has at most 3 operands.

② Assignment instruction has almost one operator on the right hand side.

③ Intermediate code generator uses temporary variables for storing intermediate results.

step 1: $t_1 = (\text{int to real}) 60$ t_1 : Temporary Variable

step 2: $t_2 = id_3 * t_1$

step 3: $t_3 = id_9 + t_2$

step 4: $id_1 = t_3$

⑤ Code Optimizers

→ The code optimizer will reduce/optimize/minimize the intermediate code.

→ The intermediate code will be the input and generates the "Optimized intermediate code".

step 1: $t_1 = id_3 * 60.0$

step 2: $id_3 = id_9 * t_1$

⑥ Code Generator: It takes the optimized intermediate code as the input and generates the "machine code".

→ The instructions for this is LDF, STF, ADDF, MULF because the example is in the "Floating point representation."

→ For LDF and STF they need 2 operands and ADDF, MULF needs 3 operands.

step 1:-
LDF R₂, id₃

step 2:-
MULF R₂, R₂, #60.0

step 3:-

LDF R₁, id₂

step 4:-

ADDF R₁, R₂

step 5:-

STF id₃, R₁

→ This is the machine code for the "a = btc * 60"

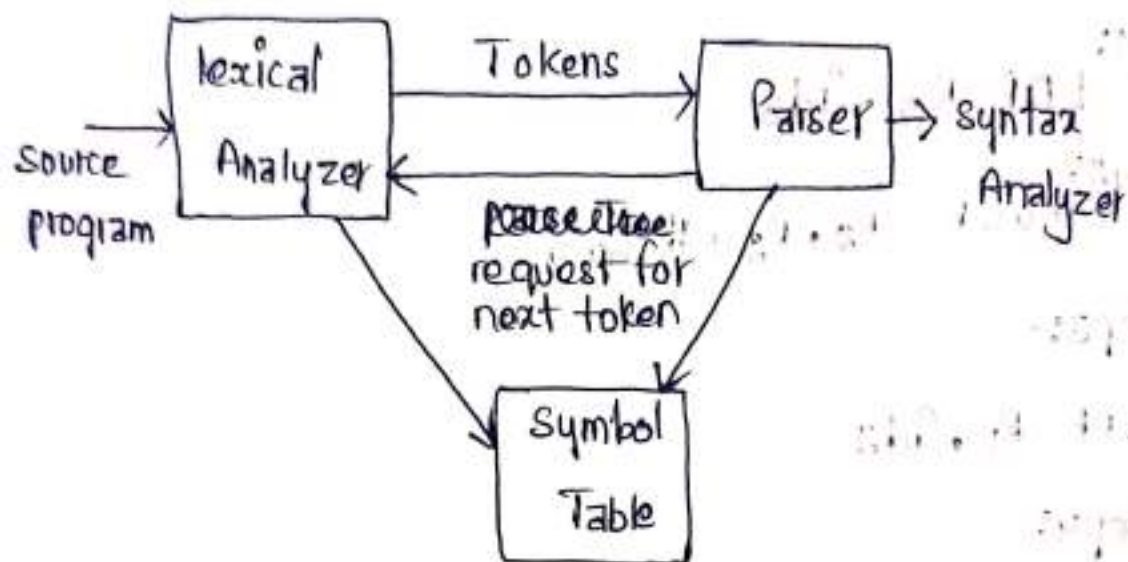
is "STF id₃, R₁"

The role of Lexical Analyzer:-

→ It generates the stream of tokens.

→ It eliminates the comments (For user understands) from the source program. It also eliminates the white space characters such as blank space, tab space and new line characters.

- It counts the no. of lines and it generates Symbol tables.
- It produces error messages with line numbers and column numbers.



Lexical Analysis versus parsing:- Benefits:-

- The task of designed ^{Compiler} is simplified. It is not easy to construct the syntax analyzer with comments.
- Efficiency of the (lexical analyzer) compiler is improved. Buffering techniques will be used to improve the lexical analyzer.
- Portability of compiler is enhanced. By restricting the abnormalities we can enhance the compiler.

Tokens, Patterns, lexemes:-

→ lexical analyzer meaningful sequence or group is called lexeme.

→ Token is a pair of "token name" and its attribute value or "optional attribute value".

→ Token names are there for keywords, constants, Identifiers, operators.

that we can make the
Pattern:- Description of the form of lexeme. to characters

lexeme:- Group of characters, meaningful sequence of, in the source program.

Token	Informal description	sample lexeme
if	characters i, f	if
else	characters e, l, s, e	else
comparison operations	>, <, >=, <=, !=	>, <, >=, <=, !=
id	letter followed by letters and digits	P1, A1
number	any numeric constants	0.635, 10.00
literals	any thing surrounded by " "	"lasya"

Attributes for the tokens:

$$e = mc^2$$

one * represents the single exponential

two ** represents the two exponential

$$e \Rightarrow \langle id, 1 \rangle$$

$$= \Rightarrow \langle = \rangle \text{ (or) } \langle assigned\ op \rangle$$

$$m \Rightarrow \langle id, 2 \rangle$$

$$* \Rightarrow \langle * \rangle \text{ (or) } \langle mul\ op \rangle$$

$$c \Rightarrow \langle id, 3 \rangle$$

$$\langle ** \rangle \text{ or } \langle exp_op \rangle$$

$$2 \Rightarrow \langle 2 \rangle \text{ number, Integer, value 2.}$$

Lexical Errors:

The errors that are recognized by the lexical analyzer is

- ① Spelling Errors [Ex: di not, present so error]
- ② unmatched strings. [if we forget the completion of the double quote].
- ③ appearance of the "illegal" characters. [if there are unwanted characters in the program].

- ① Delete one character from the matching part.
- ② Inserting a missing character to the matching input.
- ③ Replace a character by another character.
- ④ Transpose to adjacent characters.
(swap).

These are the error recovering options.

- ⑤ Delete sufficient element from the remaining input until lexical analyzer recognizes a will form tokens.

Input Buffering:- We use two buffers.

depends on the block size the buffer size will be produced.

- The normal size of buffer is 4096 bytes.
- using "One read" command we can read the "n" number of characters.
- EOF (End of file) is used to end the character where we need to stop the character at particular stage.
- In order to operate these buffers there are two pointers
 - ① lexeme Beginning: { Both start at the
 - ② Forward pointer: { beginning of the buffer.

Then it will move forward until the end of the file, But lexeme begin will move directly to the second buffer.

Sentinel: Some special characters, for reading the character there are two tests

- ① whether the Buffer is at the end
- ② What character read by the lexical Analyzer.

① when we add the special character to the Buffer at the end (sentinel) we can know that the buffer is at the end.

Recognition of Tokens

We need to write the regular expressions for the patterns. These patterns are drawn/designised as

the Transition diagrams. These are recognized as

→ recognize the keywords, constants, identifiers.

→ regular expression for digits - $"0-9"$.

→ regular expression for number - $\text{digit}(\cdot \text{digits})$

→ regular expression for white space (E(+ -) digits)?
is $w(\text{blank}|\text{Tab}|\text{newline})$.

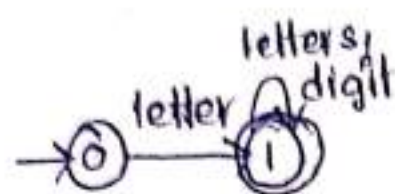
→ regular expression for letter $[A-z]$

→ regular expression for id $\rightarrow \text{letter}(\text{letter}|\text{digit})^*$

→ regular expression for if, else, then is same.
if, else, then.

→ regular expression for relop $\rightarrow <|>|=|>|=|<|$

Transition diagram for Identifiers:



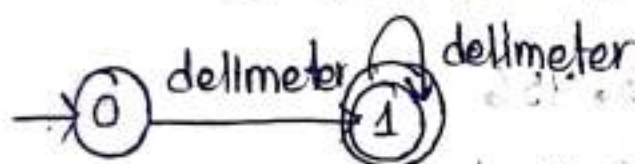
Ex:- abc12

Transition diagram for delimiters:-

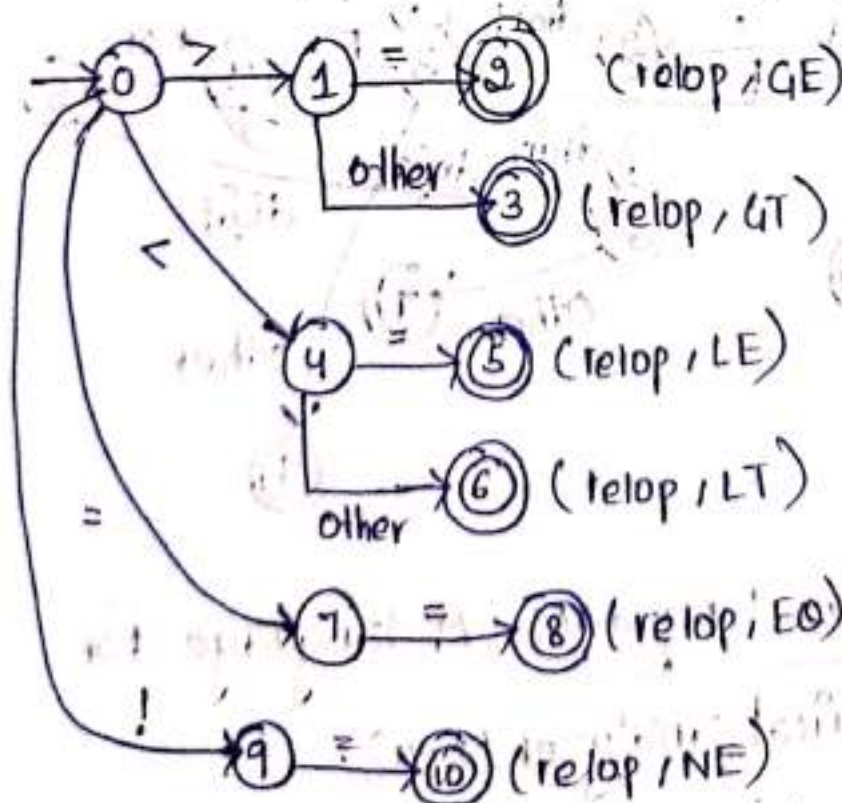
↳ It may be a blank space, newline character, tab space.

Regular expression:-

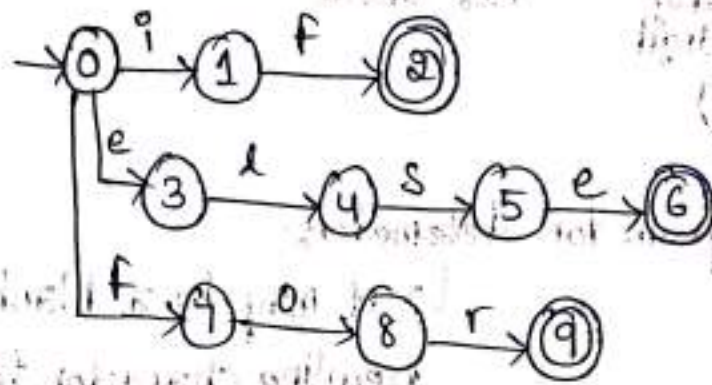
$ws \rightarrow \text{delimiter}(\text{delimiter})^*$



Transition diagram for Relational Operators:



Transition Diagram for the 'if, else, for':

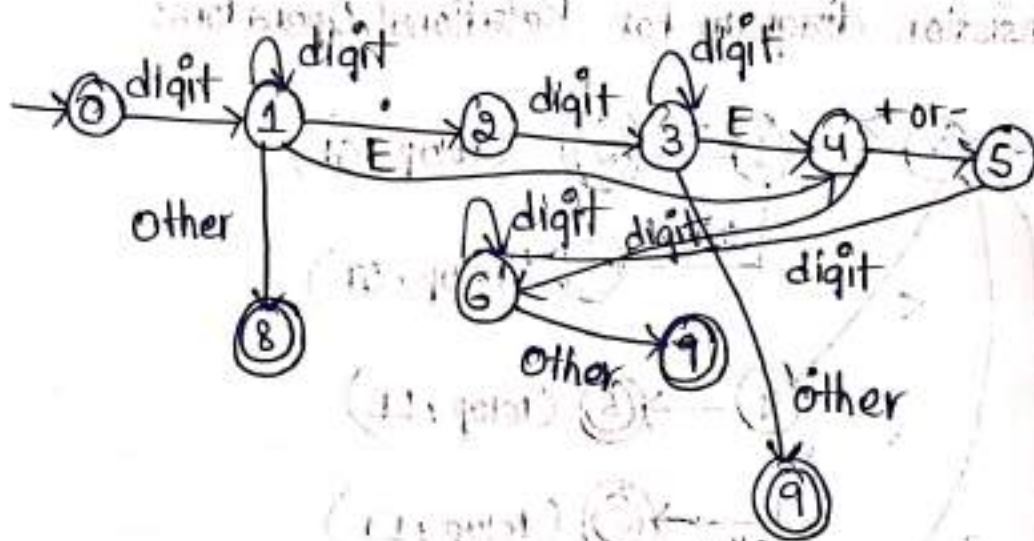


Regular expression for digits:

either digits are (int no, float no)

int digits $\rightarrow 123, 456$

float digits $\rightarrow 123.45 E 23$



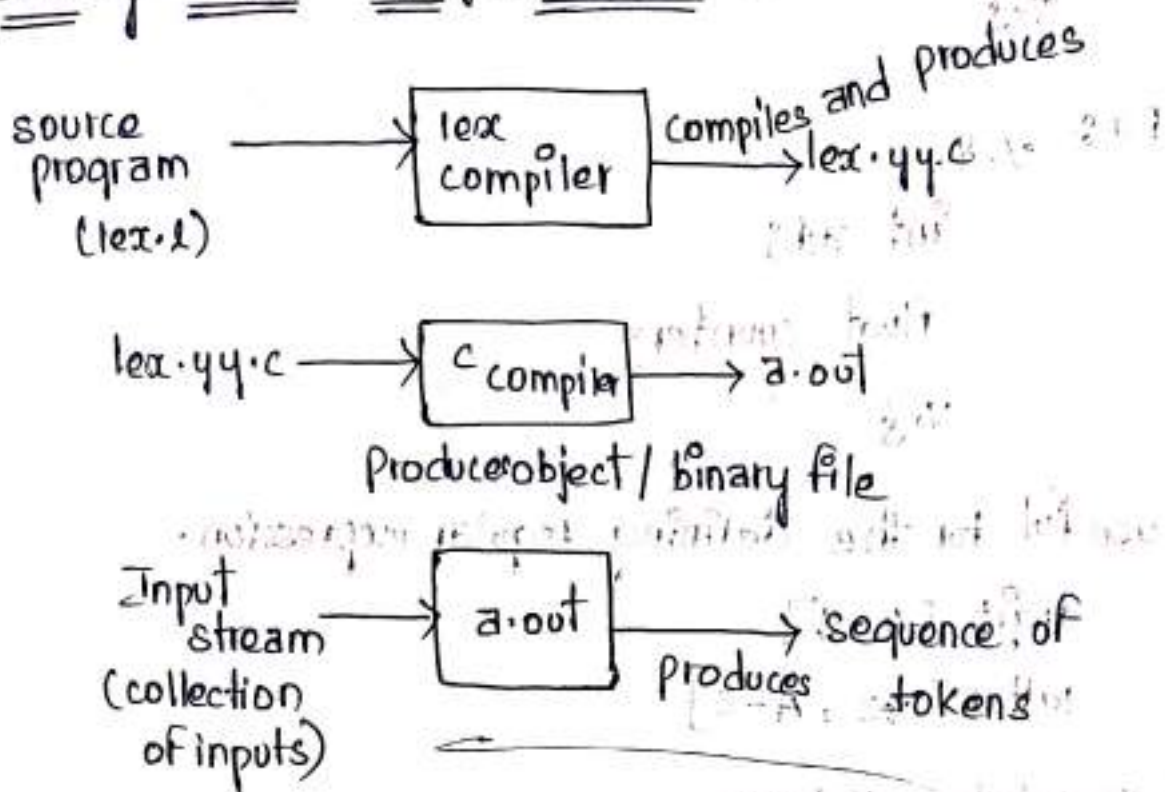
Lex Tool / Lex language: A language for specifying lexical analyzer (Lex).

(or)

The lexical analyzer generator (Lex).

- It is a tool or language which is used to generate the lexical analyzer.
- It specifies the RE, RE are used to represent the patterns for the tokens.

creating lexical Analyzer with lex:-



The Structure of lex program: It contains mainly three sections:

- ① Declarations
- ② Translation Rules
- ③ Auxiliary functions.

Declarations

% % [Translation rules are started]

Translation Rules

% % [Translation rules are overed]

Auxiliary functions

① Declarations:-

→ Used for to declare the c variables or constants.

%.&

variables, constants

%.&

Ex:-

%.&

int aib;

float counter;

%.&

→ useful for the defining regular expression.

digit [0-9]

letter [a-z, A-Z]

② Translation Rules:-

We can specify a rule in the form of pattern followed by an action.

pattern { Action }

↳ regular expression

↳ c language statement

If There are three rules:

pattern1 { Action1 }

pattern2 { Action2 }

pattern3 { Action3 }

→ These are used in order to specify the patterns.

③ Auxiliary Functions:

→ All the functions which are used they are defined in this section.

Lex program to recognize identifiers, keywords, relop, and numbers:

```
%lex
```

```
%&
```

```
/* lex program for recognizing tokens */
```

```
%&
```

```
letter [a-z, A-Z]
```

```
digits [0-9]
```

```
id {letter} ({letter} {digits})*
```

```
numbers {digits}+ ({digits}+)? ({E|+|-}? {digits}+)?
```

```
%&
```

```
{id} { printf ("%s is an identifier", yytext); }
```

```
if { printf ("%s is a keyword", yytext); }
```

```
else { printf ("%s is a keyword", yytext); }
```

```
"<" { printf ("%s is less than operator", yytext); }
```

[∴ same for the <=, >, >=, == and != operators].

{ numbers }

{ printf("%s is a number", ytext); }

Input:-

123 is 123 If else!

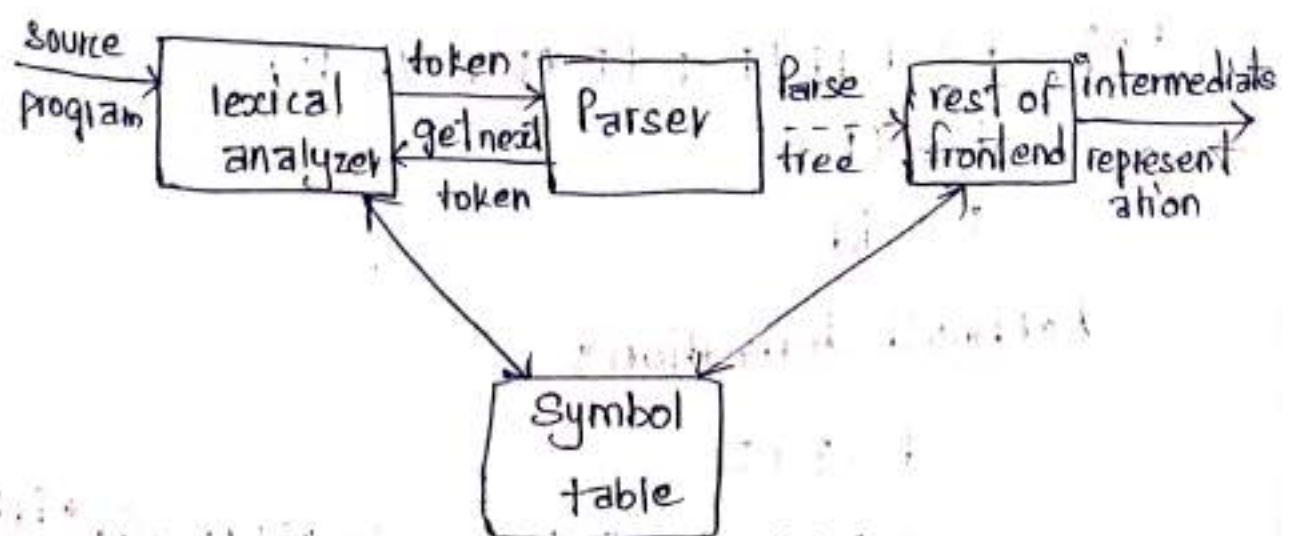
Output:

It generates the printf statements as the output.

Syntax Analysis:

Role of a Parser:

- Parser of a grammar that takes as input string of stream of tokens from the lexical analyzer and produces output as parse tree.
- The goal of a parser is to determine the syntactic validity of a source string. If valid, a tree is built for use by the subsequent phase of the computer.
- The tree reflects the sequence of derivations or reduction used during the parser.
- There are two ways of identifying an elementary parse tree:
 1. By deriving a string from a non-terminal
 2. By reducing a string of symbol to a non-terminal.



Context-Free Grammars:

→ Inherently recursive structures of a programming language are defined by a context-free Grammar.

→ In a context free Grammar, we have 4 Tuples

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

where, $V \rightarrow$ Finite set of ~~non~~ terminals

$T \rightarrow$ Finite set of non-terminals

$P \rightarrow$ Finite set of productions.

$S \rightarrow$ start symbol.

Derivations:

1. At each derivation step, we can choose any of the non-terminal in (each derivation step) the sentential form of G for the replacement.
2. If we always choose the left-most non-terminal in each derivation step, the derivation is called "left-most derivation".

Ex: $E \rightarrow E + E \mid E - E \mid E * E \mid E / E \mid -E$

$E \rightarrow (E)$

$E \rightarrow id$

Leftmost derivation:

$E \rightarrow E + E$

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

Given Grammar: $E \rightarrow E + E \mid E * E \mid (E) \mid \text{id}$
 sentence to be derived: $-(\text{id} + \text{id})$

LEFTMOST DERIVATION

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

RIGHTMOST DERIVATION

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

- String that appear in leftmost derivation are called "left sentinal forms".
- String that appear in rightmost derivations are called "right sentinal forms".

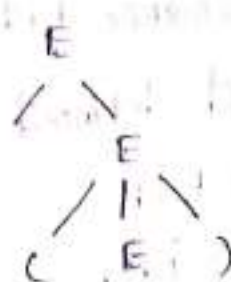
Parse Tree:

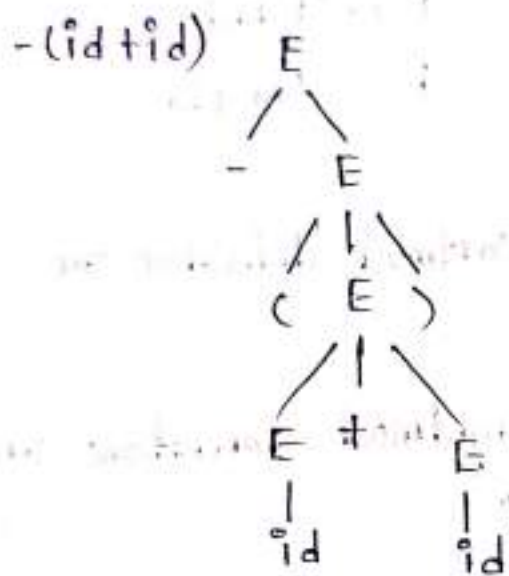
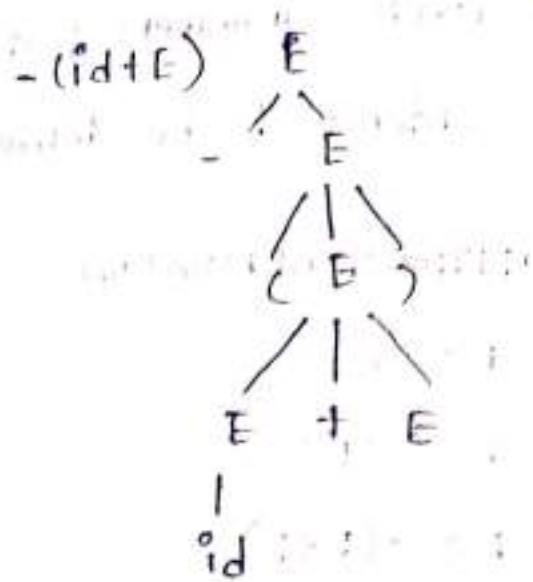
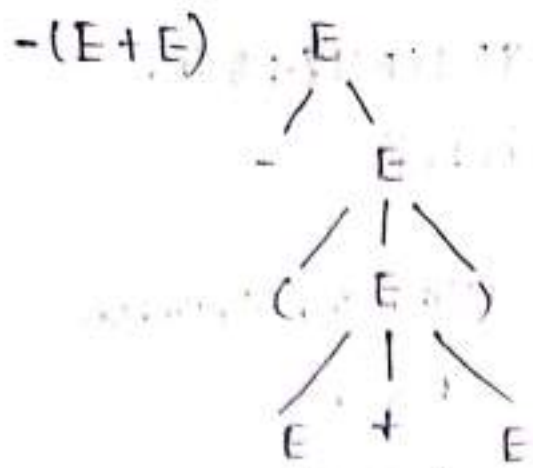
- Inner nodes of a parse tree are non-terminal symbols.
- The leaves of a parse tree are terminal symbols.
- A parse tree can be seen as a graphical representation of a derivation.

Ex:- $E \rightarrow -E$



$E \rightarrow -(E)$





Ambiguity:

A grammar that produces more than one parse for some sentence is called "ambiguous grammar."

Ex:- Given Grammar: $E \rightarrow E + E \mid E * E \mid (E) \mid -E \mid id$

The sentence $id + id * id$ has the following two distinct leftmost derivations:

$E \rightarrow E + E$

$E \rightarrow E * E$

$E \rightarrow id + E$

$E \rightarrow E + E * E$

$E \rightarrow id + E * E$

$E \rightarrow id + E * E$

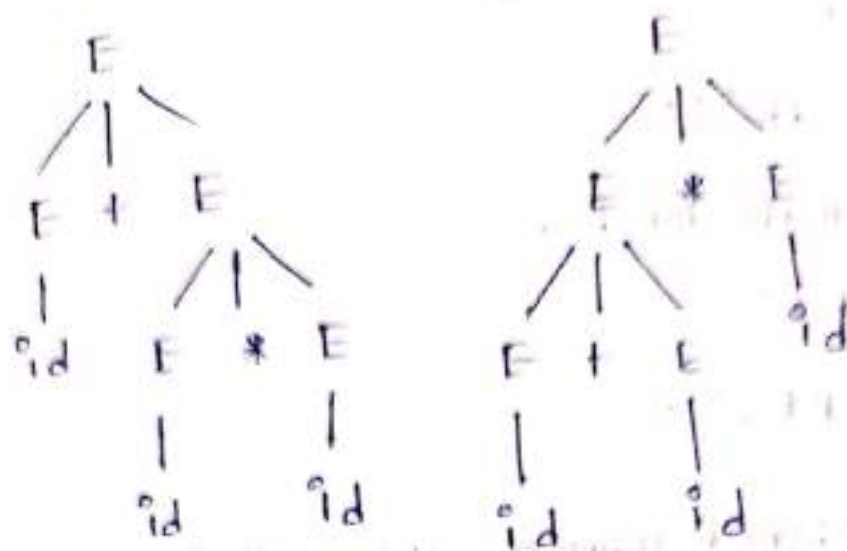
$E \rightarrow id + id * E$

$E \rightarrow id + id * E$

$E \rightarrow id + id * id$

$E \rightarrow id + id * id$

The two corresponding parse trees are:



Eliminating left Recursion:

A grammar is said to be left recursive if it has a nonterminal A such that there is a derivation $A \Rightarrow A\alpha$ for some string α . Top down parsing methods cannot handle the left-recursive grammars.

⇒ If there is a production $A \rightarrow A\alpha \mid \beta$ it can be replaced with a sequence of two productions

$$A \rightarrow \beta A'$$

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

Example:- Consider the following grammar for arithmetic expressions:

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

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

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

First eliminate the left recursion for E as

$$E \rightarrow TE'$$

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

Then eliminate for T as

$$T \rightarrow FT'$$

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

Thus the obtained grammar after eliminating left recursion is:

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

Algorithm to eliminate left recursion:

1. Arrange the non-terminals in some order $A_1, A_2, \dots, A_n$.

2. for $i := 1$ to n do begin

 for $j := 1$ to $i-1$ do begin

 replace each production of the form

$$A_i \rightarrow A_j \gamma$$

by the productions $A_i \rightarrow \delta_1 \gamma \mid \delta_2 \gamma \mid \dots$

$\delta_k \gamma$.

where $A_j \rightarrow S_1 | S_2 | \dots | S_k$ are the current A_j -productions;

end

eliminate the immediate left recursion among the A_j -productions

end.

Left Factoring:

Left Factoring is a grammar transformation that is useful for producing a grammar suitable for predictive parsing.

→ If there is any production $A \rightarrow \alpha\beta_1 | \alpha\beta_2$, it can be rewritten as

$$A \rightarrow \alpha A'$$

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

Consider the grammar, $G: S \rightarrow iEtS' | a$

$$E \rightarrow b$$

left factored, this grammar becomes:

$$S \rightarrow iEtSS' | a$$

$$S' \rightarrow eS | \epsilon$$

$$E \rightarrow b$$

Top down parsing:

It can be viewed as an attempt to find a left-most derivation of an input string or an attempt to construct a parse tree for the input starting from the root to the leaves.

→ There are two types of top-down parsing:

1. Recursive Descent Parsing:-

→ Recursive descent parsing is one of the top-down parsing techniques that use a set of recursive procedures to scan its input.

→ This parsing method may involve "backtracking", that makes repeated scans of the input.

Example:

A left-recursive grammar can cause a recursive descent parser to go into an infinite loop.

→ Hence, elimination of left-recursion must be done before parsing.

Consider the Grammar for arithmetic expressions:

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

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

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

After eliminating the left-recursion the grammar becomes:

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

Now we can write the procedure for grammar as follows:

Recursive Procedure:

Procedure E()

begin

T();

EPRIME();

End

Procedure EPRIME()

begin

If input_symbol = '+' then

ADVANCE();

T();

EPRIME();

end

Procedure T()

begin

F();

TPRIME();

End

procedure TPRIME()

begin

If input-symbol = '*' then

ADVANCE();

else if input-symbol = '(' then

ADVANCE();

E();

else if input-symbol = ')' then

ADVANCE();

end

else ERROR();

stack Implementation:

Procedure	Input String
E()	id + id * id
T()	id + id * id
F()	id + id * id
ADVANCE()	id + id * id
TPRIME()	id + id * id
EPRIME()	id + id * id
ADVANCE()	id + id * id
T()	id + id * id
F()	id + id * id
ADVANCE()	id + id * id
TPRIME()	id + id * id
ADVANCE()	id + id * id
F()	id + id * id
ADVANCE()	id + id * id
TPRIME()	id + id * id

2. Predictive Parsing:-

- predictive parsing is a special case of recursive descent parsing where no backtracking is required.
- The key problem of predictive parsing is to determine the production to be applied for a non-terminal in case of alternatives.

Predictive Parsing program:

The parser is controlled by a program that considers x , the symbol on top of stack, and a , the current input symbol. These two symbols determine the parser action. There are three possibilities:

1. If $x = a = \$$, the parser halts and announces successful completion of parsing.
2. If $x = a = \$$, the parser pops x off the stack and advances the input pointer to the next input symbol.
3. If x is a non-terminal, the program consults entry $M[x, a]$ of the parsing table M . This entry will either be an x -production of the grammar or an error entry.

If $M[x, a] = [x \rightarrow uvw]$, the parser replaces x on top of the stack by uvw .

If $M[x, a] = \text{error}$, the parser calls an error recovery routine.

Predictive parsing table Construction:

→ The construction of a predictive parser is aided by two functions associated with a grammar G :

1. FIRST
2. FOLLOW.

Rules for first():

1. If X is terminal, then $\text{FIRST}(X)$ is $\{X\}$.
2. If $X \rightarrow \epsilon$ is a production, then add ϵ to $\text{FIRST}(X)$.
3. If X is non-terminal and $X \rightarrow a\alpha$ is a production then add a to $\text{FIRST}(X)$.
4. If X is non-terminal and $X \rightarrow Y_1 Y_2 \dots Y_k$ is a production, then place a in $\text{FIRST}(X)$ if for some i , a is in $\text{FIRST}(Y_i)$, and ϵ is a $\forall$ of $\text{FIRST}(Y_1), \dots, \text{FIRST}(Y_{i-1})$; that is, $Y_1 \dots Y_{i-1} \Rightarrow \epsilon$. If ϵ is in $\text{FIRST}(Y_j)$ for all $j=1, 2, \dots, k$ then add ϵ to $\text{FIRST}(X)$.

Rules for follow():

1. If S is a start symbol, then $\text{FOLLOW}(S)$ contains $\$$.
2. If there is a production $A \rightarrow \alpha B \beta$, then everything in $\text{FIRST}(\beta)$ except ϵ is placed in $\text{follow}(B)$.

3. If there is a production $A \rightarrow \alpha B$, or a production

$A \rightarrow \alpha B \beta$ where $\text{FIRST}(\beta)$ contains ϵ , then everything in $\text{FOLLOW}(A)$ is in $\text{FOLLOW}(B)$.

Algorithm:-

Input: Grammar G

Output: Parsing table M .

Method:

1. For each production $A \rightarrow \alpha$ of the grammar, do steps 2 and 3.
2. For each terminal a in $\text{FIRST}(\alpha)$, add $A \rightarrow \alpha$ to $M[A, a]$.
3. If ϵ is in $\text{FIRST}(\alpha)$, add $A \rightarrow \alpha$ to $M[A, b]$ for each terminal b in $\text{FOLLOW}(A)$. If ϵ is in $\text{FIRST}(\alpha)$ and $\$$ is in $\text{FOLLOW}(A)$, add $A \rightarrow \alpha$ to $M[A, \$]$.
4. Make each undefined entry of M be error.

Examples:-

Consider the following Grammar:

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

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

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

After eliminating left-recursion the grammar is

$$E \rightarrow TE'$$

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

$$T \rightarrow FT'$$

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

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

First():

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

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

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

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

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

Follow():

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

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

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

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

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

Predictive parsing table:

Non terminal	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	$F \rightarrow id$			$F \rightarrow (E)$		

LL(1) Grammar:

The parsing table entries are single entries. So each location has not more than one entry.

→ This of grammar is called LL(1) grammar.

Consider this following grammar:

$$S \rightarrow iEtS \mid iEtSeS \mid a$$

$$E \rightarrow b$$

After eliminating left factoring, we have

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

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

$$E \rightarrow b$$

To construct a parsing table, we need $FIRST()$ and $FOLLOW()$ for all the non-terminals.

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

$$FIRST(S') = \{e, \epsilon\}$$

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

$$FOLLOW(S) = \{\$, \epsilon\}$$

$$FOLLOW(S') = \{\$, \epsilon\}$$

$$FOLLOW(E) = \{t\}$$

parsing table:-

Non-Terminal	a	b	e	i	t	\$
S	$S \rightarrow a$			$S \rightarrow iEtS'$		
S'			$S' \rightarrow eS$ $S' \rightarrow \epsilon$			$S' \rightarrow \epsilon$
E		$E \rightarrow b$				

Since, there are more than one production, the grammar is not LL(1) grammar.

Actions Performed in predictive parsing:

1. Shift
2. Reduce
3. Accept
4. Error

Implementation of Predictive parser:

1. Elimination of left recursion, left factoring and ambiguous grammar.
2. Construct $FIRST()$ and $FOLLOW()$ for all non-terminals.
3. Construct predictive parsing table.
4. parse the given input string using stack and parsing table.

Bottom-up Parsing:-

Constructing a parse tree for an input string beginning at the leaves and going towards the root is called "bottom-up Parsing".

→ A General type of bottom-up parser is a "shift-reduce parser".

Shift-Reduce Parsing:-

→ It is a type of bottom-up parsing that attempts to construct a parse tree for an input string beginning at the leaves (the bottom) and working up towards the root (the top).

Example:

Consider the grammar:

$S \rightarrow aABe$

$A \rightarrow Abc \mid b$

$B \rightarrow d$

The sentence to be recognized is "abcde".

REDUCTION (LEFTMOST)

abcde ($A \rightarrow b$)

aAbcde ($A \rightarrow Abc$)

aAde ($B \rightarrow d$)

aABe ($S \rightarrow aABe$)

RIGHTMOST DERIVATION

$S \rightarrow aABe$

$\rightarrow aAde$

$\rightarrow aAbcde$

$\rightarrow abcde$

The reductions trace out the right-most derivation in reverse.

Actions in shift-reduce parser:

- shift - The next input symbol is shifted onto the top of the stack.
- reduce - The parser replaces the handle within a stack with a non-terminal.
- accept - The parser announces successful completion of parsing.

- error - The parser discover that a syntax error has occurred and calls an error recovery routine.

Stack Implementation of shift-reduce parsing:

Stack	Input	Action
\$	$id_1 + id_2 * id_3 \$$	shift
$\$ id_1$	$+ id_2 * id_3 \$$	reduce by $E \rightarrow id$
$\$ E$	$+ id_2 * id_3 \$$	shift
$\$ E +$	$id_2 * id_3 \$$	shift
$\$ E + id_2$	$* id_3 \$$	reduce by $E \rightarrow id$
$\$ E + E$	$* id_3 \$$	shift
$\$ E + E *$	$id_3 \$$	shift
$\$ E + E * id_3$	$\$$	reduce by $E \rightarrow id$
$\$ E + E * E$	$\$$	reduce by $E \rightarrow E * E$
$\$ E + E$	$\$$	reduce by $E \rightarrow E + E$
$\$ E$	$\$$	accept

Conflicts in shift-reduce parsing:

There are two conflicts that occur in shift-reducing parser.

1. shift-reduce Conflict:

The parser cannot decide whether to shift or to reduce.

Examples-

Consider the Grammar:

$E \rightarrow E + E \mid E * E \mid id$ and input $id + id * id$.

Stack	Input	Action	Stack	Input	Action
\$E+E	*id\$	Reduce by $E \rightarrow E+E$	\$E+E	*id\$	shift
\$E	*id\$	shift	\$E+E*	id\$	shift
\$E*	id\$	shift	\$E+E*id	\$	Reduce by $E \rightarrow id$
\$E*id	\$	Reduce by $E \rightarrow id$	\$E+E*E	\$	Reduce by $E \rightarrow E*E$
\$E*E	\$	Reduce by $E \rightarrow E*E$	\$E+E	\$	Reduce by $E \rightarrow E*E$
\$E			\$E		

2. Reduce-reduce Conflict:

The parser cannot decide which of several reduction to make.

Consider the Grammar:

$$M \rightarrow R + R \mid R * c \mid R$$

$$R \rightarrow c$$

and input $c + c$

Stack	Input	Action	Stack	Input	Action
\$	$c + c \$$	shift	\$	$c + c \$$	shift
$\$c$	$+ c \$$	Reduce by $R \rightarrow c$	$\$c$	$+ c \$$	Reduce by $R \rightarrow c$
$\$R$	$+ c \$$	shift	$\$R$	$+ c \$$	shift
$\$R +$	$c \$$	shift	$\$R +$	$c \$$	shift
$\$R + c$	$\$$	Reduce by $R \rightarrow c$	$\$R + c$	$\$$	Reduce by $M \rightarrow R + c$
$\$R + R$	$\$$	Reduce by $M \rightarrow R + R$	$\$M$	$\$$	
$\$M$	$\$$				

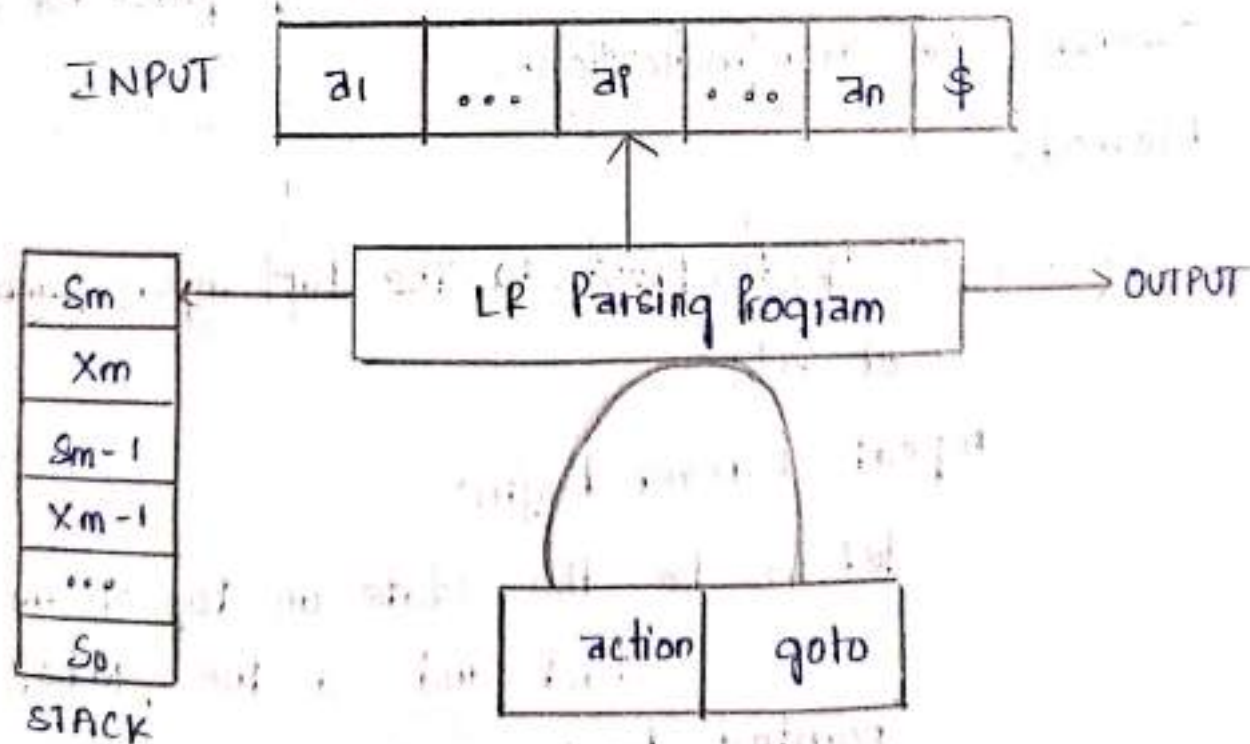
LR Parsers:-

→ An efficient bottom-up syntax analysis technique that can be used to parse a large class of CFG is called LR(k) parsing.

→ The 'L' is for left-to-right scanning of the input, the 'R' for constructing a rightmost derivation in reverse, and 'k' for the number of input symbols. when 'k' is omitted, it is assumed to be 1.

The LR parsing algorithm:

The schematic form of LR parser:



- The driver program is the same for all LR parsers.
- The parsing program reads characters from an input buffer one at a time.
- The program uses a stack to store a string of the form $S_0X_1S_1X_2S_2 \dots X_mS_m$.
- The parsing table consists of two parts: action and goto functions.

LR parsing algorithm:

Input: An input string w and an LR parsing table with functions action and goto for grammar G .

Output: If w is in $L(G)$, a bottom-up parse for w ; otherwise, an error indication.

Method:

set ip to point to the first input symbol of w ;

repeat forever begin

let s be the state on top of the stack and a the symbol pointed to by ip ;

if $action[s, a] = \text{shift } s$ then begin

push a then s' on top of the stack;

advance i to the next input symbol

end

else if $\text{action}[s, a] = \text{reduce } A \rightarrow \beta$ then begin

pop $|\beta|$ symbols off the stack;

let s' be the state now on top of stack;

Output the production $A \rightarrow \beta$

end

else if $\text{action}[s, a] = \text{accept}$ then

return

else error)

end

Constructing SLR(1) Parsing table:

To perform SLR parsing, take grammar as input and do the following:

1. Find LR(0) items.
2. Completing the closure
3. Compute $\text{goto}(I, x)$ where,
 I is set of items
 x is grammar symbol.

Algorithm for Construction of SLR parsing table:

Input: An augmented grammar G' .

Output: The SLR parsing table functions actions and goto for G' .

Method:

1. Construct $C = \{I_0, I_1, \dots, I_n\}$, the collection of sets of LR(0) items, for G' .

2. state i is constructed from I_i , The parsing functions for state i are determined as follows:

a) If $[A \rightarrow \alpha : a\beta]$ is in I_i . Then $\text{goto}(I_i, a) = I_j$, then set action $[i, a]$ to "shift j ". Here a must be terminal.

b) If $[A \rightarrow \alpha :]$ is in I_i , then set action $[i, a]$ to "reduce $A \rightarrow \alpha$ " for all a in $\text{Follow}(A)$.

c) If $[S' \rightarrow s :]$ is in I_i , then set action $[i, \$]$ to "accept".

3. The goto transitions for state i are connected for all non-terminals A using the rule:

If $\text{goto}(I_i, A) = I_j$, then $\text{goto}[i, A] = j$.

4. All entries not defined by rules (2) and (3) are made 'error'.

5. The initial state of the parser is the one constructed from the set of items containing $[s' \rightarrow \cdot s]$.

Example for SLR parsing:-

Construct SLR parsing for the following grammar:

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

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

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

The Given Grammar is:

$$E \rightarrow E + T - (1)$$

$$E \rightarrow T - (2)$$

$$T \rightarrow T * F - (3)$$

$$T \rightarrow F - (4)$$

$$F \rightarrow (E) - (5)$$

$$F \rightarrow id - (6)$$

Step 1: Convert given grammar into augmented grammar.

$$E' \rightarrow E$$

$$E \rightarrow E + T$$

$$E \rightarrow T$$

$$T \rightarrow T * F$$

$$T \rightarrow F$$

$$F \rightarrow (E)$$

$$F \rightarrow id$$

Step 2: Find LR(0) items

$$I_0: \cdot E' \rightarrow \cdot E$$

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

$$E \rightarrow \cdot T$$

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

$$T \rightarrow \cdot F$$

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

$$F \rightarrow \cdot id$$

$$GOTO(I_0, E)$$

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

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

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

$$GOTO(I_0, F)$$

$$I_3: T \rightarrow F \cdot$$

$$GOTO(I_0, ($$

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

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

$$E \rightarrow \cdot T$$

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

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

$$F \rightarrow \cdot id$$

$$T \rightarrow \cdot F$$

$$GOTO(I_0, id)$$

$$I_5: F \rightarrow id \cdot$$

$$GOTO(I_1, T)$$

$$I_6: E \rightarrow E + \cdot T$$

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

$$T \rightarrow \cdot F$$

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

$$F \rightarrow \cdot id$$

$$GOTO(I_2, *)$$

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

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

$$F \rightarrow \cdot id$$

$$GOTO(I_4, E)$$

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

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

$$GOTO(I_4, T)$$

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

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

$$GOTO(I_4, F)$$

$$I_3: T \rightarrow F \cdot$$

$$GOTO(I_6, T)$$

$$I_9: E \rightarrow E + T \cdot$$

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

$$GOTO(I_6, F)$$

$$I_3: T \rightarrow F \cdot$$

$$GOTO(I_6, ($$

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

$$GOTO(I_6, id)$$

$$I_5: F \rightarrow id \cdot$$

COTO($I_1, ()$)

$I_4: F \rightarrow (.E)$

$E \rightarrow .E + T$

$E \rightarrow .T$

$T \rightarrow .T * F$

$T \rightarrow .F$

$F \rightarrow .(E)$

$F \rightarrow .id$

COTO(I_7, id)

$I_5: F \rightarrow id.$

COTO($I_8, ()$)

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

COTO($I_8, ()$)

$I_6: E \rightarrow E + .T$

$T \rightarrow .T * F$

$T \rightarrow .F$

$F \rightarrow .(E)$

$F \rightarrow .id$

COTO($I_9, *$)

$I_7: T \rightarrow T * .F$

$F \rightarrow .(E)$

$F \rightarrow .id$

COTO($I_4, ()$)

$I_4: F \rightarrow (.E)$

$E \rightarrow .E + T$

$E \rightarrow .T$

$T \rightarrow .T * F$

$T \rightarrow .F$

$F \rightarrow .(E)$

$F \rightarrow id$

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

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

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

	ACTION						GOTO		
	id	+	*	(	)	\$	E.	T	F
I ₀	S ₅			S ₄			1	2	3
I ₁		S ₆				ACC			
I ₂		r ₂	S ₇		r ₂	r ₂			
I ₃		r ₄	r ₄		r ₄	r ₄			
I ₄	S ₅			S ₄			8	9	3
I ₅		r ₆	r ₆		r ₆	r ₆			
I ₆	S ₅			S ₄				9	3
I ₇	S ₅			S ₄					10
I ₈		S ₆			S ₁₁				
I ₉		r ₁	S ₇		r ₁	r ₁			
I ₁₀		r ₃	r ₃		r ₃	r ₃			
I ₁₁		r ₅	r ₅		r ₅	r ₅			

Stack Implementation:

STACK	INPUT	ACTION
0	$id + id * id \$$	$GOTO(I_0, id) = S5$; shift
0 id 5	$+ id * id \$$	$GOTO(I_5, +) = r6$; reduce by $F \rightarrow id$
0 F 3	$+ id * id \$$	$GOTO(I_0, F) = 3$ $GOTO(I_2, +) = r4$; reduce by $T \rightarrow F$
0 T 2	$+ id * id \$$	$GOTO(I_0, T) = 1$ $GOTO(I_1, +) = S6$; shift
0 E 1	$+ id * id \$$	$GOTO(I_6, id) = S5$; shift
0 E 1 + 6	$id * id \$$	$GOTO(I_5, *) = r6$; reduce by $F \rightarrow id$
0 E 1 + 6 id 5	$* id \$$	$GOTO(I_6, F) = 3$ $GOTO(I_3, *) = r4$; reduce by $T \rightarrow F$
0 E 1 + 6 F 3	$* id \$$	$GOTO(I_6, T) = 9$ $GOTO(I_9, *) = S7$; shift
0 E 1 + 6 T 9	$* id \$$	$GOTO(I_7, id) = S5$; shift
0 E 1 + 6 T 9 * 7	$id \$$	$GOTO(I_5, \$) = r6$; reduce by $F \rightarrow id$
0 E 1 + 6 T 9 * 7 id 5	$\$$	$GOTO(I_7, F) = 10$ $GOTO(I_{10}, \$) = r3$; reduce by $T \rightarrow T * F$
0 E 1 + 6 T 9 * 7 F 10	$\$$	$GOTO(I_7, F) = 10$ $GOTO(I_9, \$) = r3$; reduce by $T \rightarrow T * F$
0 E 1 + 6 T 9	$\$$	$GOTO(I_6, T) = 9$ $GOTO(I_9, \$) = r1$; reduce by $E \rightarrow E + T$
0 E 1	$\$$	$GOTO(I_0, E) = 1$ $GOTO(I_1, \$) = \text{accept}$

Syntax-directed definitions:

Synthesized Attributes:

Synthesized attribute for a nonterminal A at a parse-tree node N is defined by a semantic rule associated with the production N .

Inherited Attribute:

An inherited attribute for a nonterminal B at a parse-tree node N is defined by a semantic rule associated with the production at the parent of N .

Evaluate an SDD at the nodes of a parse tree:

To visualize the translation specified by an SDD, it helps to work with parse trees, even though a translator need not actually build a parse tree.

Example:

Production

Semantic Rules

$L \rightarrow E_n$

$L.Val = E.Val$

$E \rightarrow E_1 + T$

$E.Val = E_1.Val + T.Val$

$E \rightarrow T$

$E.Val = T.Val$

$T \rightarrow T_1 * F$

$T.Val = T_1.Val * F.Val$

$T \rightarrow F$

$T.Val = F.Val$

$F \rightarrow (E)$

$F.Val = E.Val$

$F \rightarrow \text{digit}$

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

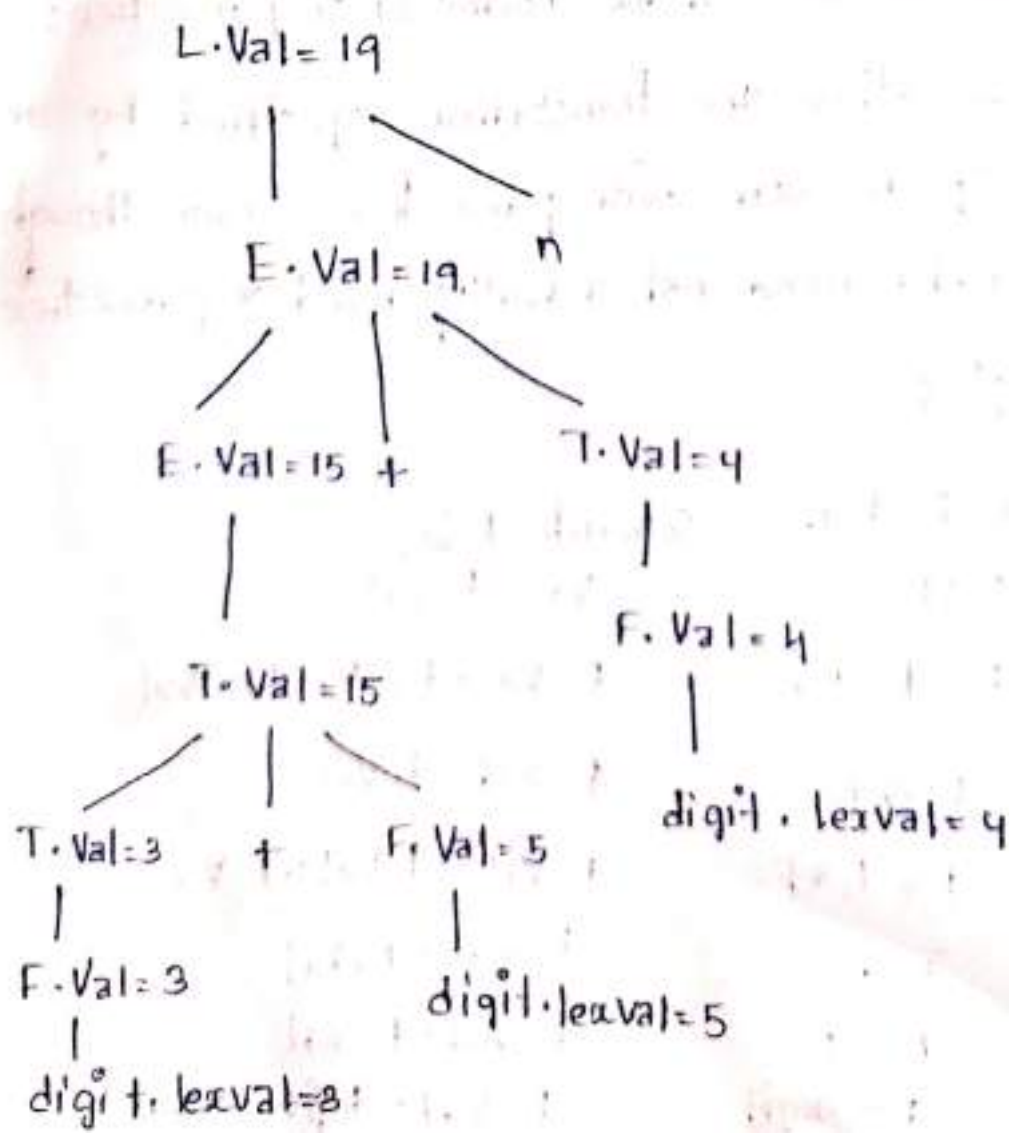
The SDD based for arithmetic expressions with operators "+" and "*". It evaluates expressions terminated by an endmarker ϵ .

→ In the SDD, each of the non-terminals has a single synthesized attribute, called "Val".

→ Suppose that the terminal digit has a synthesized attribute lexval, which is an integer value returned by the lexical analyzer.

Example:

$3 * 5 + 4 \epsilon$



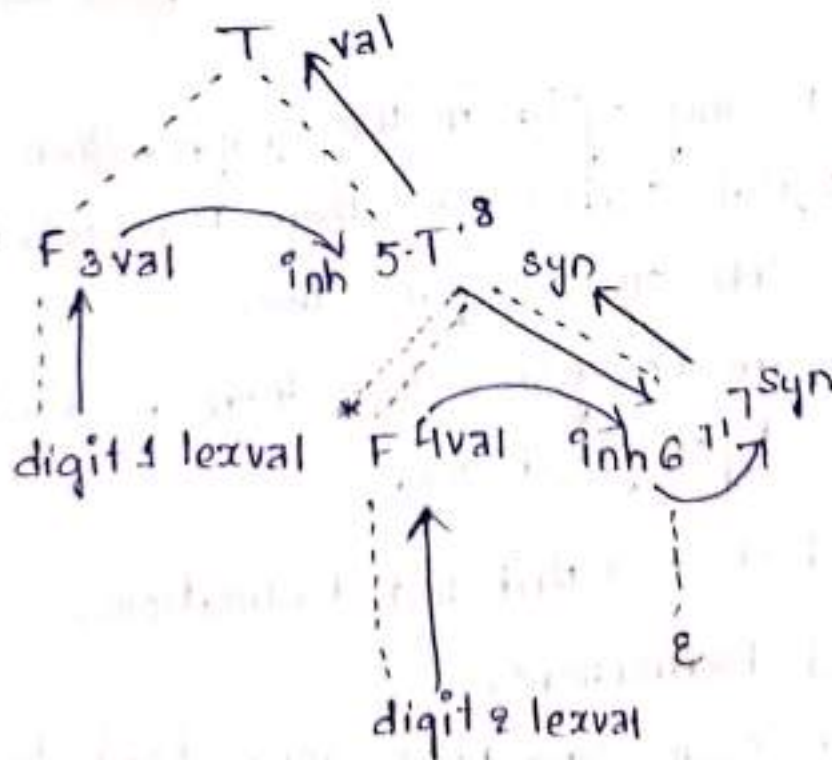
Evaluation Orders for spp's:

Dependency Graphs:

A dependency graph depicts the flow of information among the attribute instances in a particular parse tree: an edge from one attribute instance to another means that the value of the first is needed to compute the second. Edges express constraints implied by the semantic rules.

Example:

The nodes of the dependency graph, represented by the numbers 1 through 9, correspond to the attributes in the annotated parse tree:



Ordering the Evaluation of Attributes:

→ The dependency graph characterizes the possible orders in which we can evaluate the attributes at the various nodes of a parse tree.

→ If the dependency graph has an edge from node M to N , then the attribute corresponding to M must be evaluated before the attribute of N .

→ Thus, the allowable order of evaluation are those sequences of nodes $N_1, N_2, \dots, N_k$ such that if there is an edge of the dependency graph from N_i to N_j then $i < j$.

→ Such an ordering embeds a directed graph into a linear order, and is called a "topological sort" of the graph.

→ If there is any cycle in the graph, then there are no topological sorts, i.e.; there is no way to evaluate the SDD on this parse tree.

→ If there are no cycles, then there is always at least one topological sort.

3 - Attributed & L-Attributed Definitions:-

S-Attributed Definitions:

For a given SDD, it is very hard to tell whether there exists any parse trees whose

dependency graphs have cycles.

→ Translations can be implemented by using classes of SDD's that guarantee an evaluation order, because they did not permit dependency graphs with cycles.

→ There are two classes which are implementing efficiently in connection with top-down or bottom-up parsing.

d. Attributed Definitions:-

The second class of SDD's is called L-attributed definitions.

→ The class is defined that between the attributes associated with a production body, dependency-graph edges can go from left to right, but not from right to left.

→ Each attribute must be either:

1) synthesized

2) Inherited, but the rules are as follows:

Suppose, there is a production $A \rightarrow x_1 x_2 \dots x_n$.
there is an inherited attribute x_i computed by a rule associated with this production.

a) Inherited attributes associated with the head A

~~A~~

b) Either inherited or synthesized attributes associated with the occurrences of symbols $x_1, x_2, \dots, x_{i-1}$ located to the left of x .

Semantic Rules with Controlled Side Effects:

- A balance between attributes grammars and translation schemes.
- Attribute Grammars have no side effects and allow any evaluation order consistent with the dependency graph.
- Translation schemes impose left-to-right evaluation and allow semantic actions to contain any program fragment.

Side Effects in SDP's are one of following ways:

- 1) permit incidental side effects that do not constrain attribute evaluation.
- 2) Constrain the allowable evaluation orders, so that the same translation is produced for any allowable order. The constraints can be thought of as implicit edges added to the dependency graph.

Syntax-Directed Translation Schemes:

Postfix Translation Schemes:

- Simplest SDD implementation occurs when the grammar is passed bottom-up and the SDD is S-attributed.
- We can construct an SDT in which each action is placed at the end of the production and is executed along with the reduction of the body to the head of that production. SDT's with all actions at the right ends of the production bodies are called "postfix SDT's".

Parser-Stack Implementation of Postfix SDT's:

- The attributes can be put on a stack in a place where they can be found during the reduction.
- Best plan is to place the attributes along with the grammar symbol in records on the stack itself.

	x	y	z
	x.x	y.y	z.z

State/grammar Symbol

Synthesized attribute(s)

↑
top

- We can allow for more attributes, either by making the records large enough or by putting pointers to records on the stack.

→ If attributes are all synthesized, and actions occur at the ends of the productions, then we can compute the attributes for the head.

→ If we reduce by a production as $A \rightarrow XYZ$ then we have all attributes of X, Y, Z at known positions on the stack.

SOT's with Actions Inside productions:

→ An action may be placed at any position within the body of a production.

→ It is performed immediately after all symbols to its left are processed.

→ If the parse is bottom-up, then we perform action a as soon as this occurrence of a appears on the top of the parsing stack.

→ If the parse top-down, then we perform a just before we attempt to expand this occurrence.

→ SOT's can be implemented during parsing, include postfix SOT's and a class of SOT's that implements $\hookrightarrow$ -attribute definitions.

Eliminate Left Recursion from SDT's:-

- When transforming the grammar, treat the actions as if they were terminal symbols.
- The principle is based on the idea that the grammar transformation preserves the order of the terminals in the generated string.
- The actions are executed in the same order in any left-to-right parse, top-down or bottom-up.
- For eliminating left recursion is to take two productions

$$A \rightarrow A\alpha | \beta$$

that generates strings consisting of a β and any number of α 's and replace them by productions that generates the same using a new nonterminal R of the first production.

$$A \rightarrow \beta R$$

$$R \rightarrow \alpha R | \epsilon$$

Example:- Consider the following E-productions from an SDT for translating infix expression into postfix notation:

$$E \rightarrow E_1 + T \mid (E_1 + T);$$

$$E \rightarrow T$$

→ If we apply the standard transformation to E, the remainder of the left-recursive production is $\alpha = +T \{ \text{Print}(' + '); \}$ and β , the body of the production is T. If we introducing R for the remainder of E, we get of productions:

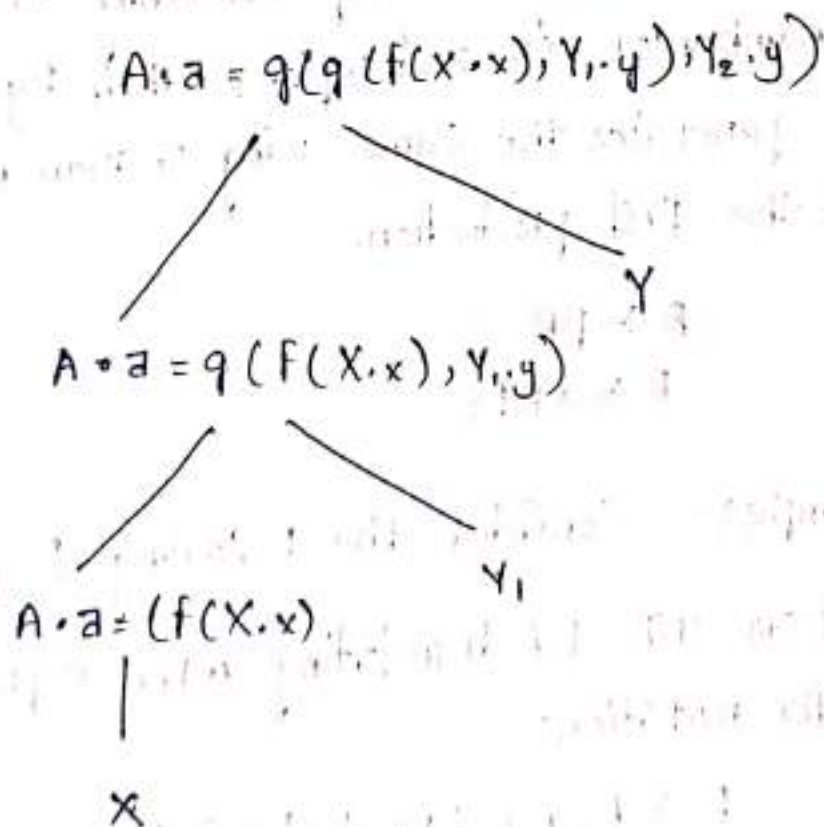
$$E \rightarrow TR$$

$$R \rightarrow +T \{ \text{Print}(' + '); \} R$$

$$R \rightarrow \epsilon$$

$$A \rightarrow XR$$

$$R \rightarrow YR \mid \epsilon$$



Implementing 2-Attributed SDN's:

Translation during Recursive-Descent Parsing:

→ A recursive-descent parser has a function A for each non-terminal A . We can extend the parser into a translator as follows:

- The arguments of function A are the inherited attributes of nonterminal A .

- The return-value of function A is the collection of synthesized attributes of nonterminal A .

→ In the body of function A , we need to both parse and handle attributes:

- 1) Decide upon the production used to expand A .
- 2) Check that each terminal appears on the input when it is required.
- 3) Preserve, in local variables, the values of all attributes needed to compute inherited attributes for nonterminals in the body or synthesized attributes for the head non-terminal.
- 4) Call functions corresponding to nonterminals in the body of the selected production, providing them with the proper arguments.

Example:-

Let us consider the SDD and SPT of example for while - statements. A pseudocode of the relevant parts of the function S appears in the below:

```
string S(table next){
```

```
    string Scode, Ccode;
```

```
    table L1, L2;
```

```
    if (current input == token while) {
```

```
        advance input;
```

```
        check '(' is next on the input,  
        and advance;
```

```
        L1 = new();
```

```
        L2 = new();
```

```
        C.code = C(next, L2);
```

```
        check ')' is next on the input,  
        and advance;
```

```
        S.code = S(L1);
```

```
        return ("label || L1 || C.code || "label"  
                || L2 || Scode);
```

```
    }  
    else
```

```
}
```

On-The-Fly Code Generation:-

The elements we need to make this technique work are:

1) For one or more non-terminals, a main attribute.

2) The main attributes are synthesized.

3) The rules that evaluate the main attribute(s) ensure that -

a) The main attribute is the concatenation of main attributes of non-terminals appearing in the body of the production, perhaps, with other elements that are not main attributes, such as the string label or the values of labels L_1 and L_2 .

b) The main attributes of non-terminals appear in the rule in the same order as the nonterminals themselves appear in the production body.

→ The main attribute can be constructed by emitting the (non-terminals) non-main-attribute elements of the concatenation. We can rely on the recursive calls to the translation.

Example:-

We can modify the function of diagram to emit elements of the main translation S. code instead of saving them for concatenation into a return value of S. code. The revised function S appears in the following:

```
void S(label next){
```

```
    label L1, L2;
```

```
    if (current input == tokenwhile){
```

```
        advance input;
```

```
        check ' ' is next on the input, and
```

```
        advance;
```

```
        L1 = new();
```

```
        L2 = new();
```

```
        Print ("label", L1);
```

```
        C (next L2);
```

```
        check ' ' is next on the input,  
        and advance;
```

```
        print ("label", L2);
```

```
        S(L1);
```

```
    }
```

```
    else /* other statement types */
```

```
    }
```

d-attributed sdp's and LR parsing:

→ suppose an d-attributed sdp is based on LL-grammar & that we have converted it to an sdp with actions embedded in the productions.

→ we can then perform the translation during LL parsing by extending the parser stack to hold actions & certain data items needed for attribute evaluation.

→ The data items are copies of attributes:

The two principles to manage attributes on the stack:

1) The inherited attributes of nonterminal A are placed in the stack record that represents the nonterminal.

2) The synthesized attributes for a nonterminal A are placed in a separate synthesize-record that is immediately below the record for A on the stack.

→ We had separated data using for different purposes into different records.

→ Action-records contain pointers to code to be executed.

The principle that makes all this copying of attributes work is:

- All copying takes place among the records that are created during one expansion of one nonterminal.

Bottom-up parsing of k -attributed SDP's:-

It has three parts:-

1) start with the SDP construction, which places embedded actions before each nonterminal to compute its ~~into~~ inherited attributes and an action at the end of the production to compute synthesized attributes.

2) Introduce into the grammar a marker nonterminal in place of each embedded action. Each such space gets a marker that is on production for marker M . i.e. $M \rightarrow \epsilon$.

3) Modify the action a if marker nonterminal M replaces it in some production $A \rightarrow \alpha \{a\} \beta$, and associates with $M \rightarrow \epsilon$ an action a' that

a) copies, as inherited attributes (in the same way ~~way~~) of M , any attributes of A or symbols of α that action a needs.

b) Computes, attributes in the same way as a , but makes those attributes be synthesized attributes of M .

Variants of Syntax Trees:

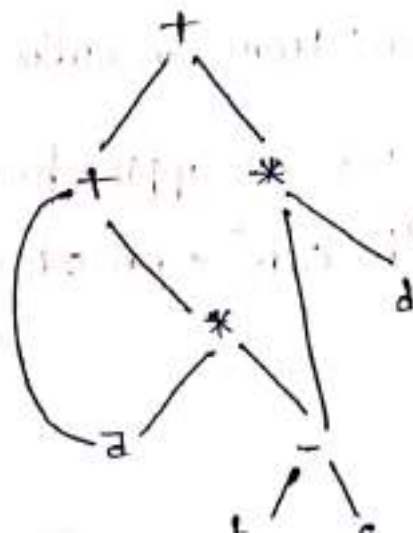
Directed Acyclic Graphs for Expressions:

- Syntax tree for an expression is like that of a DAG has leaves corresponding to atomic operands & interior nodes corresponding to operators.
- Thus, DAG not only represents expressions more sufficiently, it gives the compiler important clues regarding the generation of efficient code to evaluate the expressions.
- The difference is that a node N in a DAG has more than one parent if N represents a common subexpression.
- In a syntax tree, the tree for the common subexpression would be replicated as many times as the subexpression appears in the original expression.

Example:

The below shows the DAG for the expression:

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



Value-Number Method for constructing DAGs:-

Algorithm:-

The value-number method for constructing the nodes of a DAG.

Input:

Label op , node l , and node r

Output:

The value number of a node in the array with signature $\langle op, l, r \rangle$.

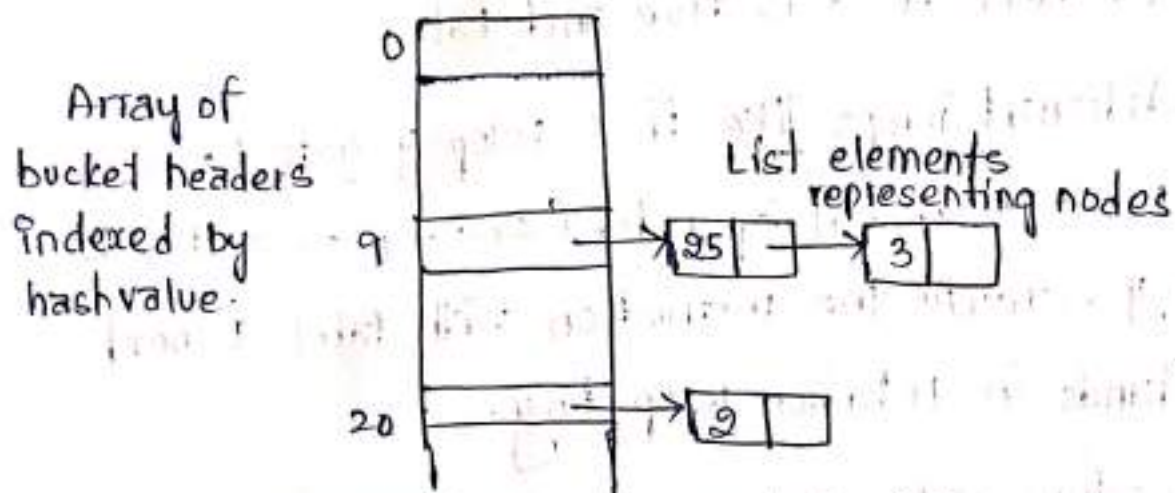
Method:

Search the array for a node M with label op , left child l and right child r . If there is such a node, return the value number of M . If not, create in the array a new node N with label op , left child l , right child r , and return its value number.

With this algorithm yields the desired output, searching the entire array every time we are asked to locate one node is expensive, especially if the array holds expressions from an entire program.

To avoid this, we use an approach, i.e; hash table, nodes are kept into "buckets" each of have a few nodes.

To construct a hash table, for the nodes of a DAG, we need a hash function h that computes the index of the bucket for a signature $\langle op, i, r \rangle$ in a way that distributes the signatures across buckets



Three Address Code:-

The list of common three-address instruction form are as follows:

1. Assignment instructions of the form $x = y \text{ op } z$, where op is a binary arithmetic or logical operation, x, y, z are addresses.
2. Assignment of the form $x = op y$, where op is an unary operation. Unary operation includes unary minus, logical negation, conversion operators.
3. Copy instructions of the form $x = y$, where x is assigned the value of y .

4. An Unconditional jump goto L. The three address instruction with label L is the next to be executed.
5. Conditional jumps of the form if x goto L and if false x goto L. These instructions execute the instruction with label L, next if x is true and false.
6. Conditional jumps like if x relop y goto L, where relop is a relational operator ($<$, $=$, $>$, etc) to x and y, execute the instruction with label L next if x stands in relation relop to y.
7. procedure calls & returns are implemented using the instructions: param x for parameters; call p, n and y = call p; n for procedure & function calls, returns y, where y represents a returned value is optional. Their typical use as the sequence of three-address instructions:

param x_1

param x_2

.....

param x_n

call p, n

8. Indexed copy instructions of the form $x = y[i^1]$ and $x[i] = y$.

$x = y[i]$ sets x to the value in the location i units memory beyond location y .

$x[i] = y$ sets the contents of the location i units beyond x to the value of y .

q. Address & pointer assignments of the form:

$x = \&y \Rightarrow$ sets the r-value of x to be the location (l-value) of y .

$x = *y \Rightarrow$ x is a pointer name or a temporary whose r-value is a location.

$*x = y \Rightarrow$ r-value of the object pointed to by x to the r-value of y .

Example:-

Consider the statement:

do $i = i + 1$; while ($a[i] < v$);

There are two possible translations for the statement.

In (a): uses a symbolic label L , attached to the first instruction.

(b): shows position number for the instructions, starting at position 100.

L : $t_1 = i + 1$

$i = t_1$

$t_2 = i * 8$

if $t_3 < v$ goto L

(a) symbolic labels

100 : $t_1 = i + 1$

101 : $i = t_1$

102 : $t_2 = i * 8$

103 : $t_3 = a[t_2]$

104 : if $t_3 < v$ goto 100

(b) Position Numbers

Quadruples:

These instructions can be implemented as objects or as records with fields for the operator and operands.

→ A quadruple (or just "quad") has four fields, $op, arg_1, arg_2, result$.

where op is an internal code for the operator.

Example:

Three address code for the assignment

$a = b * -c + b * -c;$

The three-address code and quadruples are as follows:

$$t_1 = \text{minus } c$$

$$t_2 = b * t_1$$

$$t_3 = \text{minus } c$$

$$t_4 = b * t_3$$

$$t_5 = t_2 + t_4$$

$$a = t_5$$

a) Three address code.

	op	arg1	arg2	result
0	minus	c		t ₁
1	*	b	t ₁	t ₂
2	minus	c		t ₃
3	*	b	t ₃	t ₄
4	+	t ₂	t ₄	t ₅
5	=	t ₅		a

Triples:

→ It has only three fields i.e.; op, arg1 and arg2.

→ Using triples, we refer to the result of an operation,

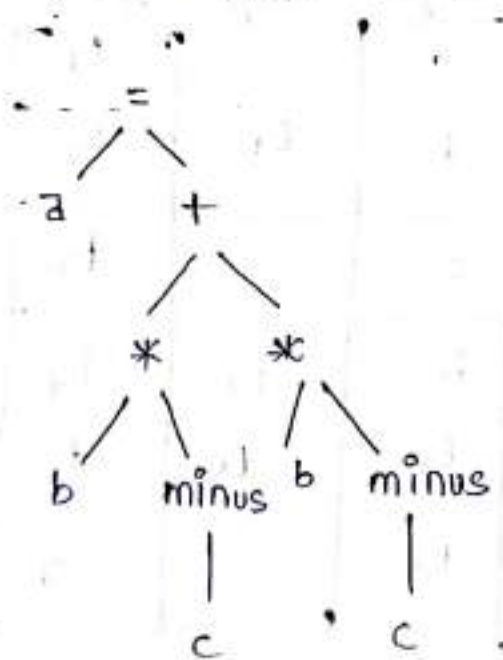
x op y by its position.

→ A triple representation refer to position (0) parenthe sized numbers represent pointers into the triple structure itself.

→ Triples are equivalent to signatures. Hence, the DAG & triple representations of expressions are equivalent.

→ Since syntax-tree variants and three-address code represents ^{control} flow quite differently.

Example: The syntax tree and triples representation of $a = b * -c + b * c;$



(a) syntax tree

	op	arg1	arg2
0	minus	c	
1	*	b	(0)
2	minus	c	
3	*	b	(2)
4	+	(1)	(3)
5	=	(a)	(4)

Static Single - Assignment:

- SSA is an intermediate representation that facilitates certain code optimizations.
- Two distinctive aspects distinguish SSA from three-address code.
- The first is that all assignments in SSA are to variables with distinct names; hence the term SSA.
- The below diagram shows the representation of TAC and in SSA.

Note: The subscripts distinguish each differentiation of variables p and q in the SSA representation.

$$\begin{aligned}
 p &= a + b \\
 q &= p - c \\
 p &= q * d \\
 p &= e - p \\
 q &= p + q
 \end{aligned}$$

(a) TAC

$$\begin{aligned}
 p_1 &= a + b \\
 q_1 &= p_1 - c \\
 p_2 &= q_1 * d \\
 p_3 &= e - p_2 \\
 p_{q_2} &= p_3 + q_1
 \end{aligned}$$

(b) SSA

Intermediate Program in TAC and SSA.

Stack Allocation of Space:

Activation Trees:

stack Allocation would not be feasible if procedure calls or activation of procedures, did not nest in time. The following example illustrates nesting of procedure calls.

Example:

The below shows the sketch of a program that reads nine integers into an array a and sorts them using the recursive quicksort algorithm:

```
int a[10];  
void read Array() { /* Reads 9 integers into  
                    a[1], ..., a[9] */  
    int i;  
    ...  
}  
int partition (int m, int n) {  
    /* Returns p;  
    ...  
}  
void quicksort (int m, int n)  
{  
    int i;  
    if (n > m) {  
        ...  
    }  
}
```

```

    i = partition(m, n);
    quicksort(m, i-1);
    quicksort(i+1, n);
}
}
main()
{
    read Array();
    a[10] = -9999;
    a[10] = 9999;
    quicksort(1, 9);
}

```

There are three common cases:

- 1) The activation of q terminates normally. Then, in any language, control resumes just after the point of P at which the call to q was made.
- 2) The activation of q , or some procedure q called, either directly or indirectly aborts. In that case, P ends simultaneously with q .
- 3) The activation of q terminates because of an exception that q cannot handle. Procedure P may handle the exception in which the activation of P continues, although not necessarily from the point at which the call to q was made.

Activation Records:

Actual parameters
Returned values
Control link
Access link
Saved Machine status
Local data
Temporaries

- 1) Temporary values, such as those arising from the evaluation of expressions, in cases where those temporaries cannot be held in registers.
- 2) Local data belonging to the procedure whose activation record this is.
- 3) A saved machine status, with information about the state of the machine before the call to procedure. This information includes the return address and the contents of registers that were used by the calling procedure and that must be restored when the return occurs.

- 4) An "access link" is needed to locate data needed by the called procedure but found elsewhere.
- 5) A control link, pointing to the activation record of the caller.
- 6) Space for the return value of the called function, if any. Again, not all called procedures return a value, and, if one does, we may prefer to place that value in a register for efficiency.
- 7) The actual parameters used by the calling procedure. Commonly, these values are not placed in the activation record but registers for greater efficiency.

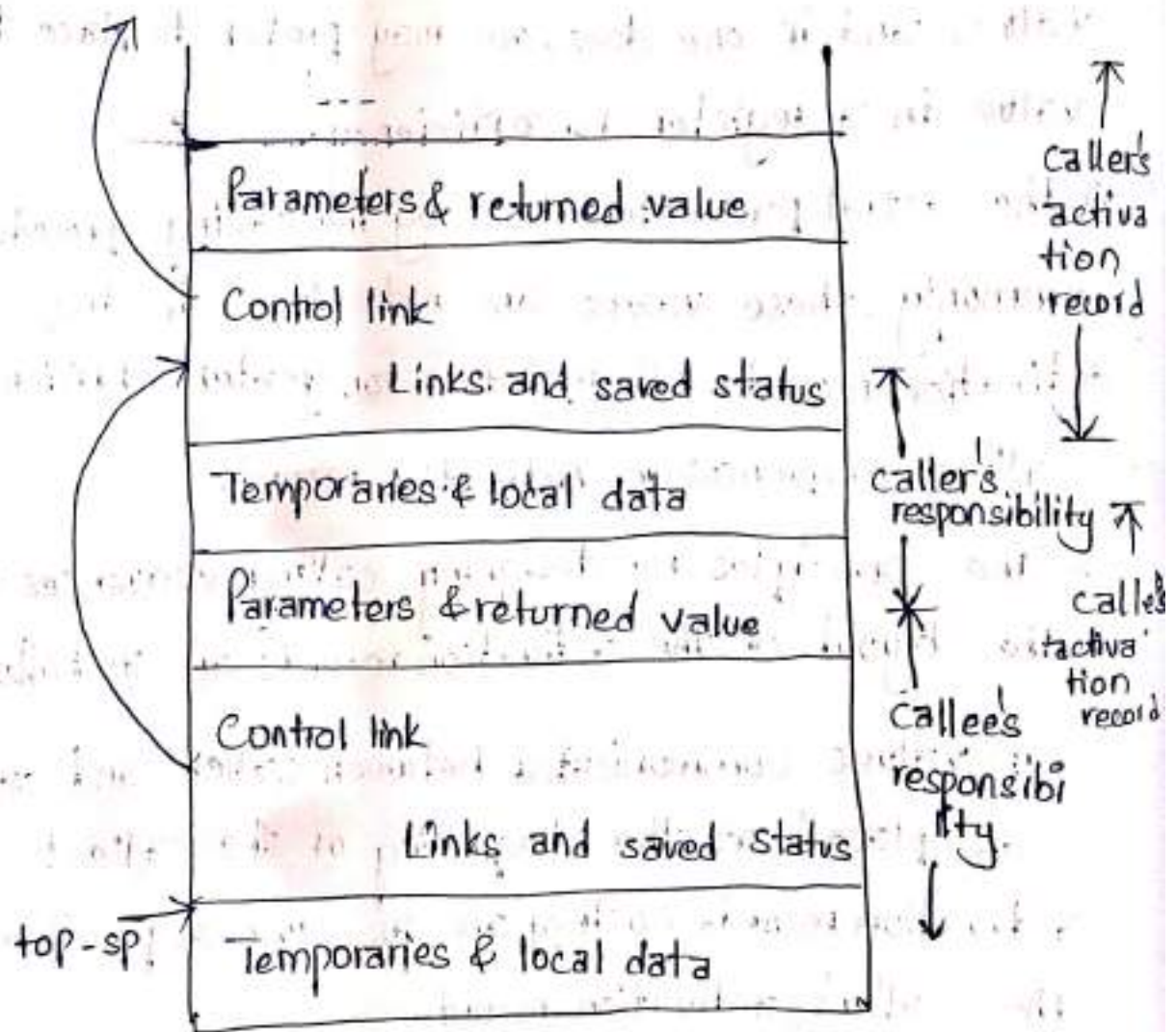
Calling Sequences in Activation records

→ The principles for designing calling sequences and the layout of the activation records are as follows:

- 1) Values communicated between caller and callee are placed at the beginning of the callee's activation record, so they are as close as possible to the caller's activation record.

- 2) Fixed-length items are placed in the middle. If exactly the same components of the machine status are saved for each call, then the same code can do the saving & restoring for each.

- 3) Items whose size may not be known early enough are placed at the end of the activation record.
- 4) We must locate the top-of-stack pointer. A common approach is to have it point to the end of the fixed-length fields in the activation record.



Access to Nonlocal Data on the stack:

Nesting depth for procedures:

let us give nesting depth 1 to procedures that are not nested within any other procedure.

Example:

All C functions are at nesting depth 1. If a procedure p is defined immediately within a procedure at nesting depth i , then give p the nesting depth $i+1$.

Example:

The below contains a sketch in ML of quicksort example.

```
fun sort (inputFile, outputFile) =  
  vallet a = array (1,10); let  
    fun readArray (inputFile) = a ... a ...;  
    fun exchange (i,j) = a ... a ...;  
    fun quicksort (m,n) =  
      vallet v = ...;  
      fun partition (y,z) =  
        ... a ... v ... exchange ...  
        in  
          ... a ... v ... partition ... quicksort  
        end  
    in
```

```
...a...:read Array ...quicksort...
```

```
end
```

→ In the code, sort are several functions:

readArray, exchange each, access the array a. Function partition is defined and it is accessed the array a and the pivot value v, and also called the function exchange. Since partition is defined because to accesses both the array a and the pivot value v, and also calls the function exchange.

Access Links:-

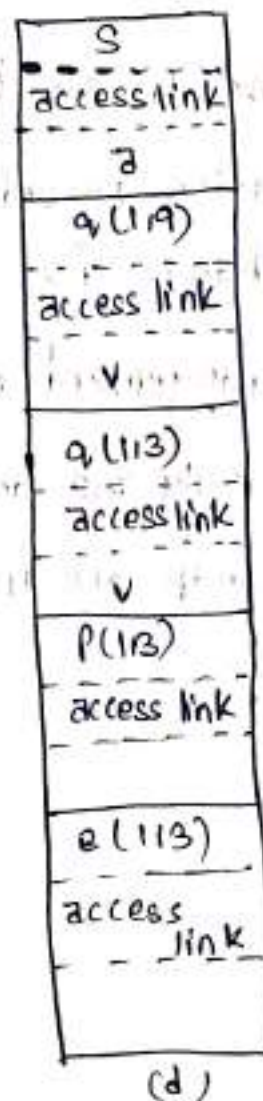
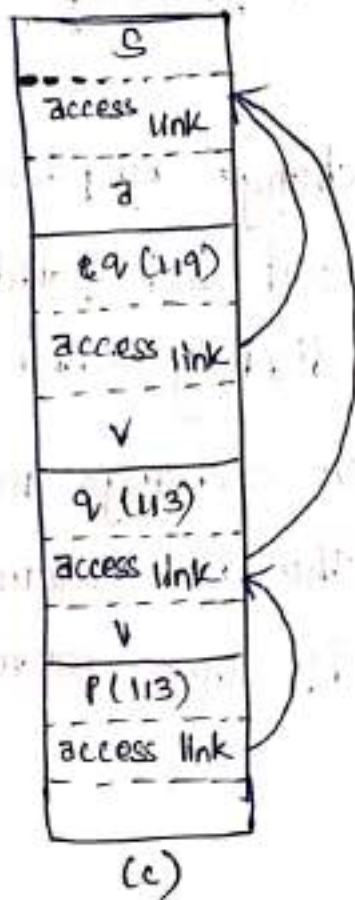
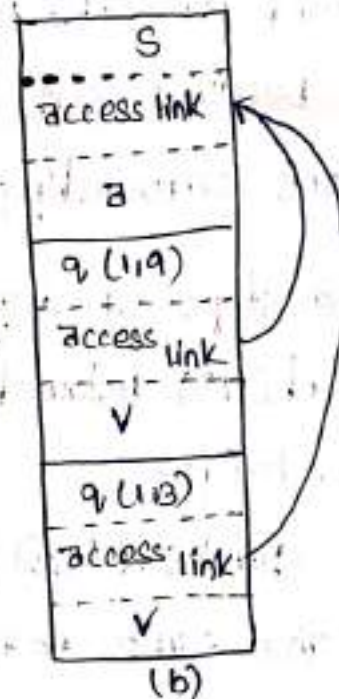
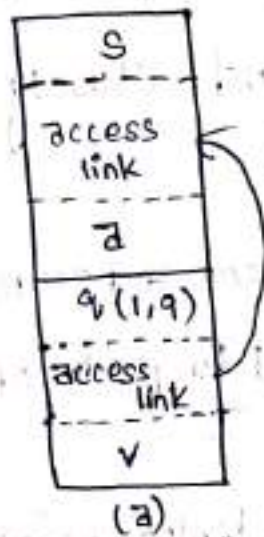
→ A direct implementation of the normal static scope rule for nested functions is obtained by adding a pointer called the access links to each activation record.

→ If the procedure p is nested immediately within procedure q in the source code, then the access link in any activation of p points to the most recent activation of q.

→ Nesting depth of p must be exactly one less than the nesting depth of q.

Example:

The below diagram shows a sequence of stacks that might result from execution of the function sort.



In (a): We see the situation after sort has called readArray to load input to the array a and then called quicksort(1,9) to sort the array.

→ The access link from quicksort(1,9) points to the activation record for sort, not because sort called quicksort, but because sort is the most closely nested functions surrounding quicksort.

→ In successive steps of diagram, we see a recursive call to quicksort(1,3) followed by a call to partition, which calls exchange.

→ notice that quicksort(1,3) is access links points to sort, for the same reason that quicksort(1,9) does.

In (d): The access link for exchange by passes the activation records for quicksort and partition, since exchange is nested immediately within sort.

→ That arrangement is fine, since exchange needs to access only the array a , and the two elements it must swap are indicated by its own parameters i and j .

Access links for procedure parameters:-

When a procedure p is passed to another procedure q as a parameter, and q then calls its parameter, it is possible that q does not know the context in which p appears in the program.

→ It is impossible for q to know the ~~the~~ how to set the access link for p .

→ ~~When~~ when procedures are used as parameters, the caller needs to pass, along with the name of the procedure-parameter the proper access link for that parameter.

→ The caller always knows the link since if p is passed by procedure r as an parameter, then p must be a name accessible to r , and r can determine the access link for p exactly if p were being called by r directly.

Example:-

The below ML function a has functions b and c nested within it. Function b has a function valued parameter f , which calls it. Function c defines with a function d and c then calls b with actual parameter c .

fun a(x) =

let

fun b(f) =

... f ...;

fun c(y) =

let

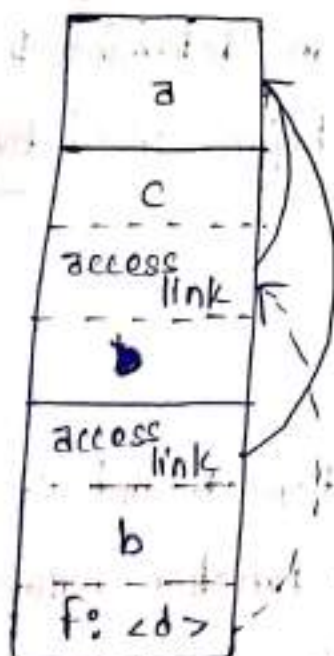
fun d(z) = ...

in ... b(d)...

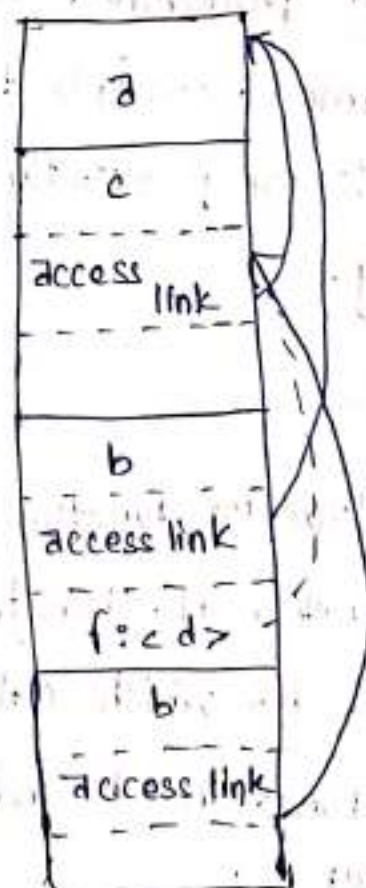
end;

in ... c(f)...

end;



(a)



(b)

Heap Management:-

Memory Manager:

- The subsystem that allocates and deallocates space within the heap.
- It keeps track of all the free space in heap storage at all times.
- Serves as an interface between application programs and the operating systems.
- It has two functions:

1. Allocation:

- It produces a chunk of contiguous heap memory of the requested size.
- If no chunk of the needed size is available, it increases the heap storage by getting consecutive bytes of virtual memory from the OS.

2. Deallocation:

- It returns deallocated space to the pool of free space (so it can reuse the space).
- Typically, does not return memory to the OS, even if heap usage drops.

The properties of memory managers are:

1. Space Efficiency:

A memory manager should minimize the total heap space needed by a program. space efficiency is achieved by minimizing "fragmentation".

2. Program Efficiency:

A memory manager should make use of the memory subsystem to allow the programs to run faster. The time taken to execute an instruction can vary widely depending on where objects are placed in memory.

3. Low Overhead:

Because memory allocation and deallocation are operations in many programs. These operations are important to be as efficient as possible. The overhead is minimized because the fraction of execution time spent performing allocation and deallocation.

Note:-

Cost of allocation is dominated by small requests; the overhead of managing large objects is less important because it is amortized over a large amount of computation.

Memory hierarchy of a Program:

- All modern computers arrange their storage as a memory hierarchy.
- The memory hierarchy follows, which consists of a series of storage elements, with the smaller faster ones, "closer" to the processor, and the larger slower ones further away.

Typical sizes

Typical Access Times

>2GB

Virtual Memory (Disk)

3-15ms

256MB-2GB

Physical Memory

100-150ns

128KB-4MB

2nd-Level Cache

40-60ns

16-64KB

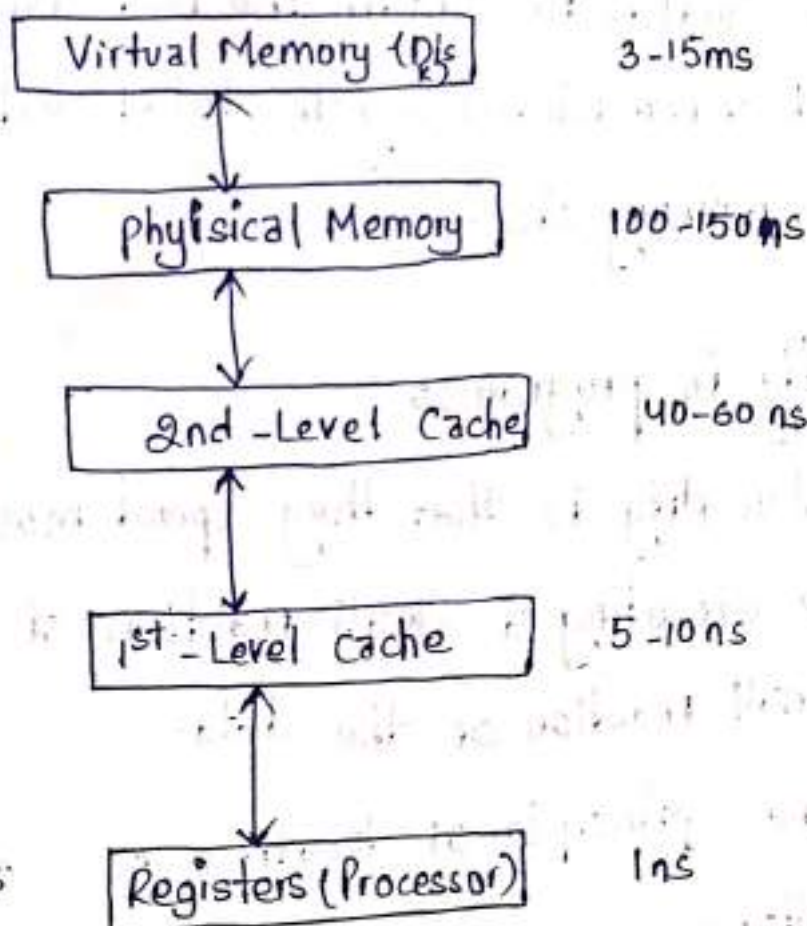
1st-Level cache

5-10ns

32 words

Registers (Processor)

1ns



- Data transferred as blocks of contiguous storage.
- To amortize the cost of access, larger blocks are used with the slower level of the hierarchy.
- Between main memory and cache, data is transferred in blocks called as cache lines, which are typically from 32 to 256 bytes long.
- Between virtual memory (disk) and main memory, data is transferred in blocks known as pages, typically between 4K and 64K bytes in size.
- The goal is to obtain the best average access speed while minimizing the total cost of the entire memory size.

locality in programs:

- Locality is that they spend most of their time executing a small fraction of code and only a small fraction of the data.
- The principle of locality is known as "locality of reference".
- The phenomenon of the same value or related storage locations being frequently accessed.

→ "90/10 Rule" comes from the observation as "A program spends 90% of its time in 10% of its code".

→ There are two types of access locality;

1. Temporal Locality:

The memory locations the program accesses are likely to be accessed again within a short period of time.

Example:

Instructions in the body of the inner loops.

2. Spatial Locality:

Memory locations close to the location accessed are likely also to be accessed within a short period of time.

Example:

Traversing the elements in a 1D array.

In this, programs spend 90% of their time executing 10% of the code because:

1) program often contain many instructions that are never executed.

2) Only a small fraction of the code that could be invoked is actually executed in a typical run of the program.

3) The program spends most of its time executing innermost loops and tight recursive cycles in a program.

Reducing Fragmentation:

→ At the beginning execution, the heap is one contiguous unit of free space.

Main memory is divided into two parts:

1) Resident operating system, usually held in low memory with interrupt vector.

2) User processor then hold in high memory.

→ As the program allocates and deallocates memory this space is broken into free and used chunks of memory and free chunks need not reside in a contiguous area of the heap.

- Free chunks of memory as holes. Holes of various size are scattered throughout memory.
- When a process arrives, it is allocated memory from a hole large enough to accommodate it.
- We reduce fragmentation by controlling how the memory manager places new objects in the heap.
- It has been found that a good strategy for minimizing fragmentation for real-life programs is to allocate the requested memory in the smallest available hole that is large enough.